



**VNiVERSIDAD
D SALAMANCA**

FACULTAD D CIENCIAS • GRADO EN ESTADÍSTICA

TRABAJO FIN DE GRADO

Análisis de Series Temporales de Índices Bursátiles

Time Series Analysis of Stock Indexes

Realizado por
Emma Santos Kofod

Tutorizado por
Rosa Amanda Sepúlveda Correa

Salamanca, 2024



**VNiVERSIDAD
D SALAMANCA**

FACULTAD D CIENCIAS • GRADº EN ESTADÍSTICA

TRABAJO FIN DE GRADO

Análisis de Series Temporales de Índices Bursátiles

Time Series Analysis of Stock Indexes

Firmado por SANTOS KOFOD EMMA –
***2346** el día 02/07/2024 con un
certificado emitido por AC FNMT

Firma de la autora
Emma Santos Kofod

Firmado por SEPULVEDA CORREA ROSA AMANDA
- ***6011** el día 02/07/2024 con un certificado emitido
por AC FNMT Usuarios

Firma de la tutora
Rosa Amanda Sepúlveda Correa

Certificado del/los tutor/es TFG

D./Dña. ROSA AMANDA SEPÚLVEDA CORREA profesor/a del
Departamento de ESTADÍSTICA de la Universidad de Salamanca,

HACE/N CONSTAR:

Que el trabajo titulado “ANÁLISIS DE SERIES TEMPORALES DE ÍNDICES BURSÁTILES”, que se presenta, ha sido realizado por EMMA SANTOS KOFOOD, con DNI ****3467G y constituye la memoria del trabajo realizado para la superación de la asignatura Trabajo de Fin de Grado en ESTADÍSTICA en esta Universidad.

Salamanca, 2 de JULIO de 2024

Firmado por SEPULVEDA
CORREA ROSA AMANDA -
***6011** el día 02/07/2024 con
un certificado emitido por AC

Fdo.: _____

Índice general

1	Introducción	1
1.1.	Estado del arte	2
1.2.	Justificación	4
2	Objetivos	5
3	Marco teórico	6
3.1.	Bolsa de valores	6
3.2.	Acciones	6
3.3.	Índices bursátiles	7
3.4.	Análisis bursátil	8
3.5.	Series temporales	10
3.6.	Inteligencia Artificial y Aprendizaje Automático	12
4	Modelización de Series Temporales	14
4.1.	Modelos clásicos	14
4.1.1.	Modelos ARIMA	15
4.1.2.	Metodología de Box-Jenkins	19
4.2.	Aprendizaje Automático	26
4.2.1.	Gradient Boosting	28
4.2.2.	Support Vector Regression (SVR)	30
5	Aplicación práctica	33
5.1.	Software	33
5.2.	Datos reales	33
5.3.	Caso práctico	35
6	Discusión	45
7	Conclusiones	47
	Referencias	48

Índice de figuras

4.1. Correlograma de un proceso $AR(1)$	17
4.2. Correlograma de un proceso $MA(2)$	18
5.1. Series logarítmicas diferenciadas divididas en conjunto de entrenamiento y prueba	36
5.2. Correlograma de la serie transformada del índice IBEX 35	37
5.3. Predicciones sobre el conjunto de prueba de los modelos ARIMA	38
5.4. Predicciones sobre el conjunto de prueba de los modelos <i>Gradient Boosting</i>	40
5.5. Predicciones sobre el conjunto de prueba de los modelos SVR	42

Índice de tablas

4.1. Comparación de FAC y FACP para modelos AR, MA y ARMA	21
5.1. Métricas de precisión de los modelos ARIMA	38
5.2. Posibles valores de hiperparámetros en <i>Gradient Boosting</i>	39
5.3. Métricas de precisión de los modelos <i>Gradient Boosting</i> manuales . .	39
5.4. Métricas de precisión de los modelos <i>Gradient Boosting</i> automáticos .	40
5.5. Posibles valores de hiperparámetros en SVR	41
5.6. Métricas de precisión de los modelos SVR manuales	41
5.7. Métricas de precisión de los modelos SVR automáticos	41
5.8. Métricas de precisión de los mejores modelos para el IBEX 35	43
5.9. Métricas de precisión de los mejores modelos para el S&P 500	43
5.10. Métricas de precisión de los mejores modelos para el Nikkei 225 . . .	43
5.11. Métricas de precisión de los mejores modelos para el IBOVESPA . . .	44

1. Introducción

En la era contemporánea, el sistema financiero y el mercado de valores desempeñan papeles esenciales para la dinámica económica global, ejerciendo una influencia significativa sobre la sociedad y la toma de decisiones económicas. De hecho, numerosos estudios relacionan, desde hace tiempo, el desarrollo de los mercados financieros con el crecimiento económico, pues apoyan que el primero es condición necesaria para impulsar el segundo (Brugger y Ortiz, 2012; Lanteri, 2011; Walker, 1998).

El mercado de capital tiene como función canalizar el ahorro hacia la inversión, la cual se lleva a cabo a través del mercado de valores (Rosero, 2010). Por tanto, el mercado de valores representa un componente fundamental del sistema financiero, pues no solo facilita la compra y venta de instrumentos financieros de compañías que cotizan en bolsa, sino que también proporciona a las empresas la oportunidad de obtener financiamiento y permite a los inversores participar en el crecimiento y rendimiento de las compañías (Hernández, 2000). Esta oportunidad de obtener capital posibilita la expansión de las empresas, y fomenta la innovación y la generación de empleo, lo que a su vez impulsa el crecimiento económico. Además, las inversiones contribuyen a la creación de riqueza a largo plazo y desempeñan un papel esencial en la planificación de la jubilación.

Ahora bien, a pesar de que la bolsa de valores ofrece la posibilidad de obtener rendimientos significativos, también conlleva riesgos inherentes, por lo que los inversores deben evaluar detenidamente sus objetivos y tolerancia al riesgo antes de participar en este mercado. En particular, es importante entender el comportamiento del rendimiento y la volatilidad de los precios ante distintos sucesos, ya que esto está directamente ligado con la gestión de riesgos, la asignación de carteras de inversión y las decisiones de políticas públicas (Romero-Meza, Coronado, y Ibañez-Veizaga, 2021).

La amplia variedad de activos financieros comprendida en los mercados bursátiles puede dificultar el análisis y la comprensión del mercado en su conjunto. Es por ello que surgen los índices bursátiles, para simplificar el comportamiento y la tendencia que siguen un conjunto representativo de acciones al proporcionar una visión global y fácilmente interpretable del rendimiento del mercado (Jiménez Almaraz, 2010). Estos son conjuntos de valores cotizados en Bolsa que

ofrecen una visión agregada de cómo ha evolucionado un determinado valor, un conjunto de valores o un mercado bursátil (Nicolás Martínez y Devesa Santamaria, 2021).

Su seguimiento y análisis son fundamentales para comprender la dinámica del mercado de valores, identificar tendencias y evaluar la volatilidad. La creciente complejidad de los mercados financieros, marcada por la globalización, la rapidez de las transacciones y el dinamismo de los mercados, requiere herramientas analíticas más avanzadas y sofisticadas. Es precisamente en este escenario donde emerge el análisis de series temporales como un elemento crucial, ya que permite a los inversores evaluar y anticipar los riesgos asociados con los movimientos del mercado. Al analizar datos históricos y tendencias de precios, el análisis de series temporales ayuda a identificar patrones y comportamientos del mercado que pueden indicar la probabilidad de ciertos eventos (Hernández Rodríguez, 2020).

Concretamente, el análisis de series temporales de índices bursátiles es fundamental en la comprensión y anticipación de sus movimientos a lo largo del tiempo. Al identificar patrones recurrentes y tendencias históricas, se facilita la detección de posibles escenarios futuros, una información valiosa para inversores y analistas en la toma de decisiones financieras. Este enfoque también contribuye significativamente a la gestión de riesgos al identificar períodos de alta volatilidad y posibles cambios de tendencia, permitiendo ajustes estratégicos para disminuir riesgos potenciales (Tsay, 2005). Además, el análisis exhaustivo de series temporales proporciona una base sólida para el desarrollo de estrategias de inversión más efectivas, la optimización de carteras y la toma de decisiones fundadas.

1.1. Estado del arte

A lo largo de la historia, se han ido sucediendo importantes hitos en el desarrollo del sistema financiero, pero el verdadero auge se ha dado en la era actual. El número de personas y entidades interesadas en participar en el mercado ha experimentado un crecimiento exponencial (Nieves Moina, 2019), lo que ha generado la necesidad de desarrollar métodos de predicción de los valores futuros de los activos que participan en él.

En este ámbito, surge la necesidad de abordar el desafío de la toma de decisiones, que implica seleccionar una opción entre varias alternativas posibles (González Casimiro, 2009). A la hora de tomar una decisión es esencial disponer de

una visión de lo que va a suceder en el futuro, no se debe tomar una decisión sin considerar cómo evolucionarán todos aquellos eventos que puedan influir en ella. Por tanto, en los últimos años se ha destacado la importancia de mejorar el proceso de toma de decisiones, y es aquí donde cobra relevancia la idea de predicción.

Según [Navales Cardona \(2013\)](#) y [Acevedo Arango \(2012\)](#), si bien la literatura no establece una técnica dominante para el pronóstico de índices de mercado, los principales modelos matemáticos se pueden clasificar en métodos estadísticos paramétricos (ARIMA, GARCH, ARCH, regresión lineal múltiple, entre otros), métodos no paramétricos (árboles de decisión, entre otros) y computación blanda (redes neuronales, lógica difusa y algoritmos genéticos).

Los métodos más comúnmente usados se basan en los modelos ARIMA (*AutoRegressive Integrated Moving Average*), que combinan la autocorrelación de la serie temporal con la diferenciación para modelar y predecir patrones en los datos financieros, y en los modelos GARCH (*Generalized Autoregressive Conditional Heteroskedasticity*), o variaciones, que se utilizan para modelar y predecir la volatilidad en los mercados financieros al ayudar a comprender cómo varía el mercado a lo largo del tiempo y cómo puede impactar en la toma de decisiones de inversión. Estos métodos no son una novedad, pues forman parte de la literatura desde hace un considerable período de tiempo ([Bollerslev, 1986](#); [Engle, 1982](#); [Glosten, Jagannathan, y Runkle, 1993](#)). Sin embargo, debido a su amplia utilidad, estudios recientes continúan utilizando estos métodos ([Shen, Wan, y Leatham, 2021](#)) y los consideran como los más efectivos ([Osorio-Barreto, Mejía-Rubio, y Mora-Mora, 2022](#)).

Además, también es habitual emplear modelos de regresión lineal y no lineal ([Kendory, Kareem, y Humadi, 2020](#); [Valova, Gueorguieva, Aayushi, Nikitha, y Mohamed, 2023](#)), pues permiten identificar relaciones entre variables económicas y el rendimiento del mercado de valores, y pueden ayudar a entender cómo diferentes factores afectan al comportamiento de los índices bursátiles.

Ahora bien, el avance tecnológico ha facilitado la recopilación y disponibilidad de datos financieros a gran escala, proporcionando una fuerte base para el análisis de series temporales. Además, el aumento en la capacidad computacional y el desarrollo de algoritmos más sofisticados han permitido realizar análisis más complejos y detallados de las series temporales. Por tanto, se observa un creciente interés por el uso del aprendizaje automático y la inteligencia artificial para mejorar la precisión en la predicción de movimientos futuros en los índices bursátiles ([Dash y Dash, 2016](#)).

En este contexto, los modelos de redes neuronales son cada vez más populares en el análisis financiero debido a su capacidad para identificar patrones complejos en los datos, pues emplean el historial de datos y otros factores relevantes para predecir el rendimiento futuro del mercado. Además de las redes neuronales, se recurre también a menudo a otros métodos de aprendizaje automático, como árboles de decisión, *Support Vector Machine* (SVM) y modelos de *Gradient Boosting*, para realizar análisis predictivos de series temporales en el mercado de valores (Fu y Wang, 2021; Jaramillo, Velásquez, y Franco, 2017).

1.2. Justificación

Dada la alta incertidumbre asociada con los cambios en los precios de las acciones e índices (Atsalakis y Valavanis, 2009), y considerando que las series temporales tienden a ser ruidosas, no estacionarias y determinísticamente caóticas (Cao y Tay, 2001), desarrollar la capacidad de comprender, anticipar y responder a las dinámicas cambiantes del mercado se ha convertido en una necesidad.

La elección de centrar el análisis en las series temporales de índices bursátiles se fundamenta en su relevancia en el entorno financiero actual y en su amplia variedad de aplicaciones prácticas, como la predicción de precios futuros, la optimización de carteras de inversión, la detección de anomalías en los mercados y la evaluación de estrategias comerciales.

Adicionalmente, es esencial explorar cómo las metodologías y técnicas del análisis de series temporales pueden ser aplicadas al contexto de los índices bursátiles, no solo para mantenerse al día con los últimos avances en el campo, sino también para profundizar en la comprensión de los mercados financieros. Al llevar a cabo este trabajo, se pretende afianzar conocimientos sobre el tema, contribuir al progreso de áreas como la econometría, la teoría financiera y la gestión de riesgos, y ofrecer aportaciones relevantes a la literatura académica.

2. Objetivos

El presente trabajo tiene como objetivo principal comparar el desempeño predictivo de los modelos ARIMA con aquellos derivados de técnicas de aprendizaje automático, tales como *Gradient Boosting* o *Support Vector Regression* (SVR), en el análisis de series temporales de índices bursátiles. El trabajo se estructura en torno a los siguientes objetivos específicos:

- Identificar, mediante una revisión exhaustiva, el estado actual del análisis de series temporales aplicado a índices bursátiles.
- Conocer, en líneas generales, el comportamiento de los índices bursátiles, utilizando herramientas de análisis estadístico y software especializado.
- Comparar la adecuación de las funciones automatizadas de búsqueda de hiperparámetros (*auto_arima* para modelos ARIMA y *GridSearchCV* para modelos *Gradient Boosting* y SVR) con la elección no automatizada.
- Evaluar el desempeño de los modelos y la precisión de las predicciones.

3. Marco teórico

3.1. Bolsa de valores

Entre la multitud de instituciones y plataformas involucradas en la compraventa de valores financieros englobadas por el mercado de valores, se destacan las bolsas de valores. Desde esta perspectiva económica, se entiende la bolsa de valores como una institución o entidad dentro del mercado financiero que facilita la interacción entre quienes ofrecen y demandan activos financieros (Villanueva Gonzales, 2007).

Las bolsas de valores son mercados secundarios de valores, pues en ellas se negocian activos financieros emitidos previamente en el mercado primario (Gitman y Joehnk, 2009). De hecho, en contraste con el mercado primario, ni las bolsas de valores ni ningún mercado secundario implica la participación directa de la empresa emisora de los valores. En cambio, están compuestas por diferentes instituciones, como casas de bolsa, corredores o la propia bolsa de valores, que ejercen de intermediarios y gestionan las negociaciones financieras entre inversionistas y compañías.

3.2. Acciones

En las bolsas de valores se negocian muchos activos y/o productos financieros, como títulos de renta fija (bonos y obligaciones) o fondos de inversión; pero las acciones son la gran protagonista. En el mercado financiero, una acción representa una porción de propiedad de una empresa (Egea Rosas, 2020). Al adquirir una parte del capital de una sociedad, se obtiene la condición de socio de la empresa, además de un conjunto de derechos económicos y políticos (Hom Serra, 2021).

En el contexto de las inversiones, una acción implica una inversión en renta variable, ya que no garantiza un retorno fijo por contrato, sino que depende del desempeño de la empresa y de su decisión al repartir beneficios (Comisión Nacional del Mercado de Valores, 2020). Su mayor riesgo es la falta de garantía en cuanto a los rendimientos esperados. El precio por el que se intercambian las acciones de una empresa, la cotización, fluctúa en comparación con su valor de adquisición, lo que

hace posible no alcanzar los beneficios previstos o incluso perder la totalidad del capital invertido. Por este motivo, la bolsa representa un desafío complejo y azaroso.

Aunque la relación entre la oferta y la demanda influye significativamente en los precios, existen otros aspectos que pueden influir sobre el valor de las acciones, como el desempeño de la empresa o las condiciones del mercado. Por tanto, en cierto modo, el rendimiento de una compañía puede estudiarse a partir de sus cotizaciones en bolsa.

3.3. Índices bursátiles

«Un índice bursátil es un indicador, usualmente adimensional, que registra la cotización de los títulos de una muestra representativa de empresas listadas en una bolsa.» (Mendoza-Rivera, Lozano-Díez, y Venegas-Martínez, 2020, pp. 132).

En estadística, un número índice se utiliza para contrastar una o varias magnitudes en dos situaciones distintas, ya sea en diferentes contextos temporales o espaciales, tomando una de ellas como referencia. Es una herramienta que se utiliza para indicar de manera relativa el cambio porcentual experimentado por una variable en relación a un punto de referencia específico, que suele corresponder a un período de tiempo determinado (Rey y Ramil, 2007).

Por tanto, el índice bursátil es el principal indicador de la actividad bursátil, pues permite obtener información sobre la evolución de los valores de un activo financiero específico, una industria, un mercado o una zona geográfica (Canales, 2017). Además, los índices ofrecen un parámetro de evaluación del éxito de medios de inversión, como las carteras de acciones.

Estos índices no sirven solo como referencias para inversionistas y analistas, sino también como indicadores económicos que reflejan la salud general de la economía, contribuyendo así a la toma de decisiones y a la gestión de riesgos de manera más precisa. Además, al invertir en varias acciones en lugar de concentrarse en una única empresa, los inversores pueden reducir el impacto negativo de eventos específicos en su cartera, pues se distribuye el riesgo en los distintos activos. Por tanto, la diversificación que proporcionan los posiciona también como instrumentos valiosos para reducir el riesgo de la inversión (Lillo León, 2016).

Habitualmente, los principales índices de un país agrupan las compañías con mayor liquidez y capitalización bursátil, como el IBEX 35 en España, el Nikkei 225

en Japón y el S&P 500 en Estados Unidos. Sin embargo, también pueden agrupar empresas de un sector específico, como el Nasdaq, que considera empresas tecnológicas y biotecnológicas. Por tanto, según los datos que resuman, pueden clasificarse en índices geográficos e índices sectoriales (Rossi, Arbocco, Vacas, Maco, y Posadas, 2012).

De acuerdo con Martín (1994), aunque los índices bursátiles sobre precios de acciones pueden formarse empleando promedios geométricos de las cotizaciones, el uso de medias aritméticas es el más frecuente. Además, los números índices pueden clasificarse en índices simples o complejos, según analicen la evolución de una o más variables.

Los índices complejos son los que mayor relevancia tienen en el ámbito bursátil debido a la gran cantidad de valores que existen en el mercado. Analizar únicamente un título no es suficiente para determinar la dirección del mercado. En función de la importancia relativa de los distintos títulos, estos se clasifican, a su vez, en dos categorías: los índices complejos no ponderados, cuyos componentes tienen todos el mismo peso independientemente de su tamaño o importancia relativa en el mercado; y los índices complejos ponderados, a cuyos componentes se les asignan pesos proporcionales a su relevancia (López Rial, 2016).

Es común enfocarse en los índices ponderados, pues considerarlos todos de manera equitativa podría deformar la percepción del estado del mercado. En este punto, aunque Fisher (1922), y Paasche (1874), han propuesto metodologías alternativas para su cálculo, la teoría formulada por Laspeyres (1871) es ampliamente utilizada para los índices bursátiles españoles. Siendo P_{it} los precios correspondientes al momento t , P_{i0} los precios iniciales y Q_{i0} los factores de ponderación en el momento base 0, el cálculo del índice de Laspeyres (I) sigue la siguiente expresión:

$$I = \frac{\sum_{i=1}^n P_{it} \times Q_{i0}}{\sum_{i=1}^n P_{i0} \times Q_{i0}} \quad (3.1)$$

3.4. Análisis bursátil

Afirma Cardona Ochoa (2023) que el análisis del comportamiento de los precios de las acciones es un área de especial interés en finanzas, pues debido a la multitud de riesgos inherentes asociados a las inversiones, se busca comprender las variaciones de los precios en los mercados de valores.

El análisis bursátil consiste en el estudio detallado y exhaustivo de los aspectos financieros y bursátiles de los títulos-valores, así como de los factores técnicos y económicos que afectan a una inversión en el mercado de valores (Rodríguez Aranday, 2020). En el ámbito del análisis bursátil, se pueden distinguir dos enfoques: el análisis técnico y el análisis fundamental.

El análisis técnico sostiene que el precio, junto con el volumen de negociación, son las variables más importantes, pues el precio refleja toda la información relevante, mientras que el volumen determina la fuerza del movimiento del precio (Sánchez García, 2015). Este método se basa en el supuesto de que la historia tiende a repetirse y que los precios de los activos tienden a seguir tendencias discernibles. Siguiendo las palabras de Pring (1991), el objetivo del análisis técnico es anticipar con suficiente margen un cambio en la dirección del mercado para adaptarse a él hasta que la evidencia demuestre que la tendencia ha vuelto a cambiar. Se centra en estudiar cómo evolucionará el precio de un activo financiero en el futuro, mediante el uso de gráficos y herramientas técnicas, como promedios móviles, osciladores y niveles de soporte y resistencia, para identificar patrones y señales que sugieren posibles movimientos futuros de precios.

El análisis de series temporales es una herramienta clave dentro del análisis técnico en el mercado financiero, ya que proporciona una base sólida para la toma de decisiones fundamentadas en el mercado financiero. Al examinar la evolución histórica, los inversores pueden formar perspectivas informadas sobre las condiciones actuales y futuras del mercado, así como evaluar y gestionar riesgos en sus inversiones, pues identificar períodos de alta volatilidad o posibles cambios de tendencia ayuda a los inversores a ajustar estrategias y mitigar riesgos potenciales (Aguirre Ariza y Castañeda Rueda, 2010).

Por el contrario, el análisis fundamental se mantiene en que el precio de cotización de un valor no refleja necesariamente el verdadero valor de una compañía. Esta metodología se basa en realizar un estudio pormenorizado de los factores económicos, financieros y cualitativos que afectan a una empresa, con el fin de evaluar con mayor precisión el valor intrínseco de un activo financiero (Moreno Gómez, 2020). Los analistas fundamentales utilizan técnicas como el análisis de ratios financieros, el descuento de flujos de caja y la valoración comparativa para estimar si el precio de un activo está sobrevalorado, indicando la posible necesidad de venta, o infravalorado, lo que sugiere la oportunidad de compra, en relación con su precio de mercado (Scherk, 2011).

3.5. Series temporales

Comprender los conceptos fundamentales de las series temporales es esencial para poder analizar e interpretar adecuadamente la información contenida en estos conjuntos de datos. Por tanto, se abordarán estos conceptos a partir de un enfoque fundamentado en las contribuciones de [Gujarati y Porter](#) (2009), cuyo libro proporciona una sólida base teórica y práctica para el análisis de series temporales.

Concepto 3.5.1. Un *proceso estocástico* es una secuencia de variables aleatorias ordenadas en el tiempo y definidas sobre el mismo espacio de probabilidad. Si Y es una variable continua, se denota por $Y(t)$, pero si se trata de una variable discreta, como es habitual en economía, se expresa como Y_t .

Concepto 3.5.2. Una *serie temporal* es una realización particular de un proceso estocástico. Se trata de una variable estadística cuyos valores varían a lo largo del tiempo y son recopilados en intervalos de tiempo regulares de manera cronológica ([Ramón Escolano y López Espín](#) 2016).

En un modelo de series temporales univariante, la variable respuesta Y se descompone en dos elementos principales: la parte sistemática PS y la innovación a_t . La parte sistemática captura los patrones regulares y predecibles identificados en la serie temporal, utilizando esta información para predecir futuras observaciones. Por otro lado, la innovación representa el componente puramente aleatorio del modelo, cuyos valores son independientes entre sí y de la parte sistemática, funcionando como ruido o variaciones inexplicables dentro del modelo.

$$Y_t = PS_t + a_t \quad \forall t = 1, 2, \dots \quad (3.2)$$

Concepto 3.5.3. Un proceso estocástico, Y_t , es *fuertemente estacionario* cuando cualquier colección finita $\{Y_1, Y_2, \dots, Y_t\}$ con $t \geq 1$ tiene el mismo comportamiento probabilístico que $\{Y_{t_1+k}, Y_{t_2+k}, \dots, Y_{t_n+k}\}$ para cualquier entero $k = \pm 1, \pm 2, \dots$. Es decir, sea F la función de distribución:

$$F(Y_1, Y_2, \dots, Y_t) = F(Y_{t_1+k}, Y_{t_2+k}, \dots, Y_{t_n+k}) \quad \forall t, k \quad (3.3)$$

Concepto 3.5.4. Un proceso estocástico es *débilmente estacionario* o estacionario en covarianza, si su media y su varianza son constantes en el tiempo y si el valor de la covarianza entre dos periodos depende solo de la distancia o rezago entre estos dos periodos, y no del tiempo en el cual se calculó la covarianza. Esto es, si una serie

temporal estocástica, Y_t , cumple:

$$E(Y_t) = \mu < \infty \quad \forall t \quad (3.4)$$

$$\text{Var}(Y_t) = E(Y_t - \mu)^2 = \sigma^2 < \infty \quad \forall t \quad (3.5)$$

$$\text{Cov}(Y_t, Y_{t+k}) = E[(Y_t - \mu)(Y_{t+k} - \mu)] = \gamma_k < \infty \quad \forall t \quad (3.6)$$

Si únicamente se cumple la condición (3.5), se dice que la serie es *estacionaria en media*, mientras que si solo se satisface la (3.6), la serie será *estacionaria en varianza*.

Concepto 3.5.5. Se dice que un proceso estocástico es *puramente aleatorio o de ruido blanco* si tiene una media igual a cero, una varianza constante σ^2 y no está serialmente correlacionado. Denotado habitualmente por a_t , debe verificar:

$$E(a_t) = 0 \quad \forall t \quad (3.7)$$

$$\text{Var}(a_t) = E(a_t^2) = \sigma_a^2 < \infty \quad \forall t \quad (3.8)$$

$$\text{Cov}(a_i, a_j) = 0 \quad \forall i \neq j \quad (3.9)$$

Por naturaleza, un proceso de ruido blanco es estacionario.

Concepto 3.5.6. Un *proceso estocástico no (débilmente) estacionario* es aquel que no cumple con las condiciones mencionadas en la definición 3.5.4, lo que implica que la media y/o la varianza experimentarán cambios a lo largo del tiempo, mostrando una clara tendencia o una dispersión no constante.

La ausencia de estacionariedad representa una desventaja significativa, ya que restringe el estudio del comportamiento de la serie temporal al período específico considerado, impidiendo la generalización para otros períodos. Sin embargo, aunque pocas series económicas presentan estacionariedad, en la mayoría de los casos es posible transformarlas para que sean estacionarias. Estas series se conocen como *procesos estocásticos integrados*, en los que al diferenciar la serie, se logra verificar la estacionariedad.

Concepto 3.5.7. En el estudio de las series temporales, la *autocovarianza de orden j* , representada como γ_j , describe la relación entre los valores de la serie temporal que están separados por j periodos. Matemáticamente, esto se puede definir como:

$$\gamma_j = \text{Cov}(Y_t, Y_{t-j}) = E[(Y_t - \mu)(Y_{t-j} - \mu)] \quad \forall j = 0, 1, 2, \dots \quad (3.10)$$

donde μ es la media de la serie temporal. Nótese que cuando $j = 0$, se trata de la

varianza:

$$\gamma_0 = \text{Cov}(Y_t, Y_t) = \text{Var}(Y_t) \quad (3.11)$$

El conjunto de autocovarianzas calculadas para los diferentes intervalos j forman la *función de autocovarianza*, que en los procesos estacionarios, verifica la propiedad de simetría: $\gamma_j = \gamma_{-j}$.

Ahora bien, la autocovarianza puede ser difícil de interpretar directamente debido a que su magnitud depende de las unidades de las variables involucradas. Por esta razón, es más común y apropiado utilizar la función de autocorrelación.

Concepto 3.5.8. La *función de autocorrelación (FAC)* se obtiene al dividir cada valor de autocovarianza por la varianza, siendo su expresión matemática la siguiente:

$$\rho_j = \text{Corr}(Y_t, Y_{t-j}) = \frac{\gamma_j}{\gamma_0} \quad \forall j = 0, 1, 2, \dots \quad (3.12)$$

Así, se produce una medida adimensional que oscila entre -1 y 1, independientemente de las unidades de medida de las variables.

No obstante, a pesar de la condición de simetría que garantiza la estacionariedad, los coeficientes de autocorrelación simple proporcionan una visión incompleta de la estructura de dependencia temporal. Ante esta situación, la *función de autocorrelación parcial (FACP)* proporciona la información adicional necesaria al eliminar el efecto de las observaciones intermedias, $Y_t, Y_{t-1}, \dots, Y_{t-j+1}$ (Ramón Escolano y López Espín, 2016).

En este contexto, las funciones de autocorrelación simple y parcial son esenciales para identificar la estructura de dependencia temporal. Esto se logra mediante el uso de un correlograma, una representación gráfica de las autocorrelaciones estimadas de una serie de datos a diferentes retrasos (Cáceres Hernández, Martín Rodríguez, y Martín Álvarez, 2008).

3.6. Inteligencia Artificial y Aprendizaje Automático

La inteligencia artificial (IA) abarca un ámbito de investigación interdisciplinario orientado a desarrollar sistemas que emulen la inteligencia humana. Estos sistemas se diseñan con la capacidad de adquirir datos y perfeccionar su funcionamiento mediante el aprendizaje, permitiéndoles llevar a cabo diversas tareas sin requerir una programación detallada para cada una de ellas (Martín González, 2023).

El aprendizaje automático (ML, por sus siglas en inglés), es una rama de la IA cuyo objetivo es crear sistemas o algoritmos capaces de aprender sin ser programados explícitamente. En lugar de depender de los programadores para establecer reglas de procesamiento de datos, las máquinas utilizan algoritmos para identificar patrones recurrentes complejos, permitiéndoles predecir comportamientos futuros y generar respuestas a partir de nuevos datos (Janiesch, Zschech, y Heinrich, 2021). Estos sistemas se mejoran de forma autónoma con el tiempo, sin intervención humana, lo que significa que pueden adaptarse y optimizar su rendimiento automáticamente (Fierro, 2020).

Los algoritmos de aprendizaje automático (ML) se clasifican principalmente en dos categorías: supervisado y no supervisado. En el aprendizaje supervisado, el algoritmo se entrena utilizando un conjunto de datos etiquetados, donde las respuestas son conocidas de antemano. Posteriormente, el modelo entrenado se evalúa con un conjunto de datos diferente, cuyas respuestas son desconocidas para la máquina pero conocidas por el usuario, permitiendo así evaluar la precisión del modelo. Por otro lado, el aprendizaje no supervisado se enfrenta a datos sin etiquetar y debe identificar estructuras y patrones subyacentes por sí mismo (Pineda Pertuz, 2022).

Dentro del aprendizaje supervisado, se destacan dos técnicas principales: los algoritmos de clasificación, que se aplican a datos con respuestas categóricas, y los algoritmos de regresión, que se emplean para predecir valores continuos (Salvador Maceira, 2019).

Los algoritmos de aprendizaje automático se basan en tres pilares fundamentales: una función de pérdida, un criterio de optimización y la optimización utilizando el conjunto de datos disponible (Díez García, 2020). La función de pérdida, o función de costo, evalúa la precisión del modelo al comparar las predicciones con los valores reales del conjunto de entrenamiento. Por otro lado, el criterio de optimización es la métrica que el algoritmo busca minimizar o maximizar durante el proceso de ajuste de los parámetros del modelo. La elección de estos elementos depende del tipo de problema y del objetivo del modelo. Finalmente, la optimización con los datos implica iterar sobre las observaciones y ajustar la función óptima utilizando una tasa de aprendizaje.

4. Modelización de Series Temporales

4.1. Modelos clásicos

Los modelos clásicos de series temporales son una parte fundamental del análisis de datos temporales en estadística y se utilizan para prever comportamientos futuros basados en datos históricos. Estos modelos son predominantemente lineales, lo que significa que asumen que las relaciones entre los términos en la serie son constantes a lo largo del tiempo y se pueden describir mediante ecuaciones lineales. Considerando a_t ruido blanco, su expresión general es la siguiente:

$$Y_t = \pi_1 Y_{t-1} + \pi_2 Y_{t-2} + \cdots + a_t \quad \forall t = 1, 2, \dots \quad (4.1)$$

En términos generales, los procesos lineales estacionarios deben satisfacer dos criterios fundamentales para ser considerados válidos en el análisis de series temporales (Gujarati y Porter, 2009): ser predictivos únicamente a partir de información pasada, y ser invertibles. El criterio de no anticipación establece que el valor de Y en el momento t no debe estar influenciado por valores ni innovaciones futuras. La invertibilidad, por otro lado, implica que la influencia de un momento anterior Y_{t-k} en el estado actual Y_t debe atenuarse progresivamente con el paso del tiempo, lo que significa que el estado actual del proceso puede expresarse de forma convergente en función de su propio historial pasado. Esta última condición se verifica si:

$$\sum_{i=1}^{\infty} \pi_i^2 < \infty \quad (4.2)$$

La evolución hacia modelos estocásticos lineales permite abordar series temporales con variaciones aleatorias subyacentes. Estos modelos asumen estacionariedad, pero frente a series no estacionarias, Box y Jenkins (1976) propusieron los modelos Autorregresivos Integrados de Medias Móviles (ARIMA), centrados en transformar la serie a una forma estacionaria (Cáceres Hernández y cols., 2008).

4.1.1. Modelos ARIMA

Los modelos ARIMA son una técnica econométrica que incorpora memoria de eventos pasados, representándose a través de una combinación lineal de observaciones previas y errores históricos (Gutiérrez, Puche, y Hernández, 2020; Ortiz Martín, 2022). Este método aplica diferenciación y técnicas de regresión para identificar patrones y realizar proyecciones futuras.

Estos modelos constan de tres elementos fundamentales: el componente Autorregresivo, AR, el componente Integrado, I, y el componente de Medias Móviles, MA (Gutiérrez Ramos, 2022). Se representa habitualmente mediante la notación $ARIMA(p, d, q)$, donde p , d y q son enteros no negativos que definen el orden de los componentes autorregresivo, integrado y de medias móviles, respectivamente. De esta forma, el modelo ARIMA se presenta como una extensión de los modelos estacionarios $AR(p)$, $MA(q)$ y $ARMA(p, q)$, permitiendo incorporar características no estacionarias mediante la integración y combinación de sus componentes básicos (nótese que un modelo $AR(p)$ puede ser visto como un $ARIMA(p, 0, 0)$, un $MA(q)$ como un $ARIMA(0, 0, q)$ y un $ARMA(p, q)$ como un $ARIMA(p, 0, q)$).

En otras palabras, una serie Y_t sigue un modelo $ARIMA(p, d, q)$ si no es estacionaria y, al integrarla d veces aplicando el operador de diferencia $\Delta^d Y_t$, se obtiene un proceso $ARMA(p, q)$ invertible y estacionario (Box y Jenkins, 1976). El operador de diferencia Δ se define como $\Delta Y_t = Y_t - Y_{t-1}$.

Procesos Autorregresivos

Los modelos autorregresivos operan sobre la premisa de que los valores anteriores de una serie temporal ejercen influencia directa sobre sus valores futuros. En esencia, estos modelos implementan una regresión de la serie contra sus propias observaciones pasadas, sugiriendo que el conjunto de datos históricos tiene un sólido poder predictivo sobre los futuros valores de la serie (Andrade Flores y Flores Carvajal, 2023).

Afirma Podadera Molero (2022) que en un modelo autorregresivo de orden p , denominado $AR(p)$, se postula que el valor actual de una serie temporal, Y_t , puede ser representado por una suma ponderada de sus p valores anteriores más un término de error aleatorio, a_t , que representa la innovación en el tiempo contemporáneo t . Sabiendo, por tanto, que el parámetro p indica la cantidad de términos pasados de la serie que se incluirán en el modelo, que los coeficientes

$\phi_1, \phi_2, \dots, \phi_p$ determinan el grado de influencia que cada uno de estos p valores rezagados tiene sobre el valor actual de la serie, que δ es el término constante, y que a_t es un ruido blanco, el modelo puede definirse de la siguiente forma:

$$Y_t = \delta + \phi_1 Y_{t-1} + \phi_2 Y_{t-2} + \dots + \phi_p Y_{t-p} + a_t \quad \forall t = 1, 2, \dots \quad (4.3)$$

Este modelo puede expresarse también de la siguiente forma abreviada:

$$\Phi_p(L)Y_t = \delta + a_t \quad (4.4)$$

donde $\Phi_p(L)$ es el polinomio autorregresivo de retardos, cuya expresión se muestra en la ecuación 4.5. En este polinomio, cada término incluye un coeficiente ϕ y una potencia de L , que representa el número de periodos hacia atrás que se consideran para la variable aplicada (Bonilla Ruiz, 2019).

$$\Phi_p(L) = 1 - \phi_1 L - \phi_2 L^2 - \dots - \phi_p L^p \quad (4.5)$$

Un proceso autorregresivo $AR(p)$ alcanza la estacionariedad bajo condiciones específicas que involucran las raíces de su polinomio autorregresivo $\Phi_p(L)$. Para ser estacionario, es esencial que las p raíces de dicho polinomio tengan un módulo superior a la unidad, $|L| > 1$, lo que equivale a que los parámetros autorregresivos cumplan con $|\phi_i| < 1$. Por otro lado, un proceso $AR(p)$ es, por definición, no anticipante, ya que su comportamiento depende exclusivamente de valores pasados y de la innovación en el momento actual. Asimismo, el proceso autorregresivo también es naturalmente invertible al mantener la suma de los cuadrados de los coeficientes ϕ_i dentro de un límite finito, facilitado por el carácter finito del orden p (Cáceres Hernández y cols., 2008).

La función de autocorrelación muestra un patrón de disminución gradual manteniendo valores distintos de cero a medida que los rezagos aumentan, mientras que las autocorrelaciones parciales se anulan para rezagos mayores al orden del modelo, lo que se traduce en un gráfico con p picos significativos distintos de cero (véase la figura 4.1), según Miranda Chinlli (2021).

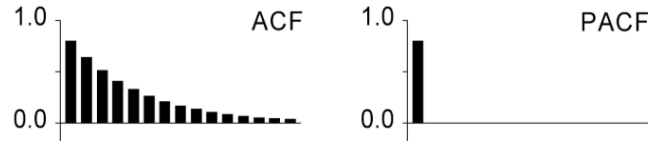


Figura 4.1: Correlograma de un proceso $AR(1)$

Fuente: [Mauricio \(2017\)](#)

Procesos de Medias Móviles

En los modelos de medias móviles, se sugiere que el valor actual de la serie está influenciado por los errores o residuos pasados, que capturan variaciones o perturbaciones aleatorias no explicadas por el modelo hasta ese momento ([Podadera Molero, 2022](#)).

Frente a los modelos autorregresivos, que se basan en los valores anteriores para predecir los siguientes, los modelos de medias móviles de orden q , $MA(q)$, hacen sus pronósticos mediante una combinación lineal de innovaciones pasadas. Es decir, el valor actual, Y_t , se construye como una combinación lineal de la innovación propia de dicha observación, sumada a las innovaciones de las q observaciones inmediatamente anteriores ([El Amrani Joutey, 2020](#)).

Así, de acuerdo con [Marcos Rubio \(2023\)](#), sea q el parámetro que determina la cantidad de innovaciones pasadas incluidas en el modelo que configuran el valor presente de la serie temporal, $\theta_1, \theta_2, \dots, \theta_q$ los coeficientes del modelo que determinan la influencia que cada residuo anterior q tiene sobre la predicción actual, δ el término constante, y a_t las variables aleatorias que representan el ruido blanco en el tiempo, el modelo viene dado por la siguiente expresión:

$$Y_t = \delta + a_t - \theta_1 a_{t-1} - \theta_2 a_{t-2} - \dots - \theta_q a_{t-q} \quad \forall t = 1, 2, \dots \quad (4.6)$$

Siguiendo el enfoque anterior, este proceso también se puede representar de forma abreviada:

$$Y_t = \delta + \Theta_q(L) a_t \quad (4.7)$$

siendo $\Theta_q(L)$ el polinomio de medias móviles de retardos:

$$\Theta_q(L) = 1 - \theta_1 L - \theta_2 L^2 - \dots - \theta_q L^q \quad (4.8)$$

Según [Cáceres Hernández y cols. \(2008\)](#), un proceso de medias móviles $MA(q)$

se define como estacionario cuando la suma de los cuadrados de sus coeficientes θ_i es finita, una condición asegurada por el carácter finito del orden q . Adicionalmente, similar a los modelos autorregresivos, los modelos $MA(q)$ también son inherentemente no anticipativos al enfocarse únicamente en innovaciones actuales e inmediatamente anteriores sin considerar información futura. Por otro lado, la invertibilidad de un modelo de medias móviles requiere que las q raíces de su polinomio $\Theta_q(L)$ se localicen fuera del círculo unidad, es decir, $|L| > 1$. De nuevo, esto es análogo a que los coeficientes del modelo tengan magnitudes menores que 1, $|\theta_i| < 1$.

Como destaca [Ramón Escolano y López Espín \(2016\)](#), la condición de estacionariedad en un modelo autorregresivo es formalmente similar a la condición de invertibilidad en un modelo de medias móviles. Mientras que un modelo $MA(q)$ es inherentemente estacionario, requiere ser evaluado en cuanto a su invertibilidad. Por otro lado, un modelo $AR(p)$ es fundamentalmente invertible, pero es necesario verificar su estacionariedad.

El comportamiento de los correlogramas de los modelos $MA(q)$ es opuesto al de los modelos $AR(p)$. En este caso, la función de autocorrelación se cancela para rezagos que exceden el orden q del modelo, mientras que en el gráfico de la función de autocorrelación parcial se aprecia una tendencia amortiguada, sin llegar a anularse ([Miranda Chinlli, 2021](#)). Véase la figura [4.2](#).

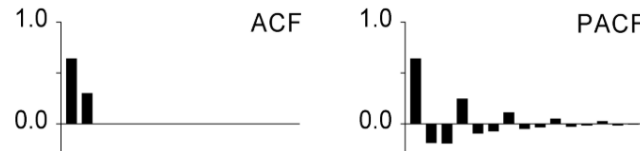


Figura 4.2: Correlograma de un proceso $MA(2)$

Fuente: [Mauricio \(2017\)](#)

Procesos Autorregresivos de Medias Móviles

Los procesos autorregresivos de medias móviles, conocidos como $ARMA(p, q)$, integran características de los modelos autorregresivos de orden p y de medias móviles de orden q discutidos anteriormente, siendo su expresión general la siguiente:

$$Y_t = \delta + \phi_1 Y_{t-1} + \cdots + \phi_p Y_{t-p} - \theta_1 a_{t-1} - \cdots - \theta_q a_{t-q} + a_t \quad \forall t = 1, 2, \dots \quad (4.9)$$

Alternativamente, esta relación también se puede describir utilizando operadores de retardos:

$$\Phi_p(L)Y_t = \delta + \Theta_q(L)a_t \quad (4.10)$$

donde $\Phi_p(L)$ y $\Theta_q(L)$ son, respectivamente, el polinomio autorregresivo y el polinomio de medias móviles basados en los retardos.

Así, el comportamiento temporal se ve afectado por mediciones pasadas de la misma variable, que se integran en el modelo a través de componentes autorregresivos, junto con efectos aleatorios o perturbaciones, tanto pasadas como presentes, que se reflejan en los términos de medias móviles (Kinjang Li Ye, Aguilar Game, y Muñoz Franco, 2023). Esto es, el valor actual de la serie, Y_t , se deriva de una combinación lineal de los p valores pasadas, las q innovaciones anteriores, y la innovación actual. Por tanto, se dice que los modelos $ARMA(p, q)$ son capaces de describir la relación de una serie temporal con el ruido blanco (a través de la componente de media móvil) y con sus propios valores pasados (mediante la componente autorregresiva), según Podadera Molero (2022).

Dado que el modelo ARMA combina elementos tanto autorregresivos (AR) como de media móvil (MA), la parte AR del modelo deberá satisfacer las condiciones de estacionariedad, mientras que la parte MA deberá cumplir con las condiciones de invertibilidad (Ramón Escolano y López Espín, 2016). En otras palabras, el modelo ARMA será estacionario e invertible cuando las p raíces del polinomio autorregresivo, $\Phi_p(L)$ y las q raíces del polinomio de medias móviles $\Theta_q(L)$, respectivamente, estén fuera del círculo unidad. Esto se corresponde con que los módulos de los coeficientes del modelo, $|\phi_i|$ y $|\theta_i|$, sean menores que 1.

Del mismo modo, el componente de medias móviles define la función de autocorrelación, la cual decrece a partir del rezago q , mientras que el componente autorregresivo determina el rezago p a partir del cual la función de autocorrelación parcial comienza a decaer. Sin embargo, en los modelos ARMA, ni la función de autocorrelación ni la función de autocorrelación parcial se anulan de manera tan clara como en los modelos AR y MA puros, debido a la combinación compleja de ambos componentes (Gutiérrez Ramos, 2022).

4.1.2. Metodología de Box-Jenkins

Box y Jenkins (1976) desarrollaron un enfoque sistemático diseñado para construir modelos que se ajusten de manera óptima a una serie temporal específica

con el objetivo de realizar pronósticos precisos. Se asume que la serie temporal es una manifestación de un proceso estocástico específico y que puede ser modelado.

Esta metodología se centra principalmente en la identificación, estimación y verificación de modelos ARIMA para series temporales, seguido de su aplicación para predicción (Gutiérrez y cols., 2020).

Identificación

La primera fase, denominada identificación, busca determinar el modelo ARIMA más adecuado para la serie temporal. Esto implica encontrar los órdenes p , d y q que ofrezcan el mejor ajuste, además de evaluar si es necesario incluir un término constante en el modelo (Andrade Flores y Flores Carvajal, 2023).

El primer paso en esta etapa es realizar un análisis de estacionariedad. Para que una serie temporal sea considerada estacionaria debe carecer de tendencias y su varianza debe permanecer constante a lo largo del tiempo. Este análisis se puede llevar a cabo mediante la inspección visual del gráfico de series temporales, que muestra los valores de la serie temporal a lo largo del tiempo (Miranda Chinlli, 2021). Un gráfico que revele una tendencia clara o una varianza no constante puede indicar que la serie no es estacionaria en media o en varianza, respectivamente.

Para complementar la inspección visual, es aconsejable realizar pruebas estadísticas que verifiquen la estacionariedad de la serie temporal. Una de las pruebas más comunes y recomendadas para este propósito es la prueba de raíz unitaria de Dickey-Fuller Aumentada (ADF), que evalúa la presencia de una raíz unitaria en una muestra de una serie temporal, lo que sería indicativo de no estacionariedad en media (Gómez Rueda, 2016; Osorio y Ángel, 2019).

Frente a la no estacionariedad en varianza se pueden aplicar diversas transformaciones estabilizadoras de la varianza, siendo una de las más destacadas la transformación de Box y Cox (1964), cuya fórmula es la siguiente:

$$Z(\lambda) = \begin{cases} \frac{y^\lambda - 1}{\lambda} & \text{si } \lambda \neq 0 \\ \log(y) & \text{si } \lambda = 0 \end{cases} \quad (4.11)$$

El valor óptimo de λ se elige de manera que la transformación de Box y Cox maximice la función de verosimilitud de la serie transformada, o alternatively, minimice la desviación estándar de los errores ajustados. Ahora bien, generalmente, la transformación logarítmica suele ser suficiente para alcanzar la estacionariedad

en varianza (Cáceres Hernández y cols., 2008).

Para abordar la ausencia de estacionariedad en media, el método más comúnmente utilizado es la diferenciación. Este proceso implica restar cada valor de la serie temporal con el valor inmediatamente anterior, lo cual ayuda a eliminar tendencias y convertir la serie en estacionaria. La diferenciación puede aplicarse tantas veces como sea necesario para lograr la estacionariedad, y el número de veces que se aplica es representado por el parámetro d en el modelo ARIMA (Miranda Chinlli, 2021). Es decir:

$$Z_t = (1 - L)^d Y_t \quad (4.12)$$

Generalmente, una serie que requiere diferenciación una vez suele presentar una tendencia lineal, mientras que una que requiere diferenciación dos veces podría tener una tendencia cuadrática.

Una vez conseguida la serie estacionaria, el siguiente paso consiste en identificar los órdenes p y q del modelo ARMA o ARIMA. Como se ha mencionado, esto se realiza mediante el análisis de las funciones de autocorrelación (FAC) y autocorrelación parcial (FACP), que muestran cómo los valores actuales de la serie temporal están correlacionados con sus valores pasados, y sus patrones pueden sugerir los órdenes adecuados para el modelo (Gujarati y Porter, 2009). La tabla 4.1 resume el comportamiento teórico de estas funciones para los diferentes modelos.

	FAC	FACP
AR(p)	Decrece sin anularse	Se anula para $j > p$
MA(q)	Se anula para $j > q$	Decrece sin anularse
ARMA(p,q)	Decrece sin anularse	Decrece sin anularse

Tabla 4.1: Comparación de FAC y FACP para modelos AR, MA y ARMA

Cáceres Hernández y cols. (2008) destacan el principio de parsimonia que sustenta esta metodología, el cual establece que deben priorizarse los modelos más simples de entre las opciones disponibles. Se busca una aproximación adecuada utilizando el menor número posible de parámetros.

Finalmente, se evalúa la necesidad de incluir un término constante (δ) en el modelo ARIMA. Si la media de una serie estacionaria es significativamente diferente de cero, agregar un término constante en el modelo puede ser relevante. En contraste, si la media es cero, implicaría que el proceso tiene una media nula por naturaleza, lo que sugiere que un término constante en el modelo ARIMA no

es necesario.

Estimación

Tras identificar el proceso estocástico que ha generado la serie, la fase de estimación se centra en cuantificar los parámetros del modelo (Gómez Rueda, 2016). Esto implica utilizar el conjunto de entrenamiento para estimar los valores de los coeficientes de los términos autorregresivos (ϕ_i) y de medias móviles (θ_i) determinados en la fase anterior, así como el término constante (δ), si se ha decidido su inclusión.

Según Cáceres Hernández y cols. (2008), los métodos más comunes utilizados en el proceso de estimación de parámetros son el método de Máxima Verosimilitud (MV) y el método de Mínimos Cuadrados Ordinarios (MCO). El método de MV, implementado en la librería statsmodels de Python, maximiza la función de verosimilitud, que es la probabilidad de observar los datos dados ciertos valores de los parámetros del modelo (Andrade Flores y Flores Carvajal, 2023). En otras palabras, el MV encuentra los valores de los parámetros que hacen que los datos observados sean los más probables.

Validación

La fase de validación consiste en asegurar que el modelo estimado se ajusta adecuadamente al conjunto de datos. Esta etapa incluye un análisis de los coeficientes estimados para verificar su significatividad individual, así como las condiciones de estacionariedad e invertibilidad para garantizar la estabilidad del modelo. Además, se realiza un análisis de los residuos para confirmar que se comportan como un ruido blanco (Ortiz Martín, 2022).

En el análisis de los coeficientes estimados, el primer paso es verificar si todos ellos contribuyen significativamente al modelo. Esto se realiza mediante los contrastes habituales de significatividad individual para los parámetros AR y MA, ϕ_i y θ_i , así como para el término constante, δ , si está incluido (Ferrari, 2023). Por tanto, se plantean hipótesis nulas que postulan que los parámetros son nulos frente a hipótesis alternativas que sostienen que los coeficientes son significativamente distintos de cero. Esto se contrasta utilizando el estadístico t habitual, el cual sigue una distribución t de Student.

Para verificar las condiciones de estacionariedad e invertibilidad, se calculan las raíces del polinomio autorregresivo y del polinomio de medias móviles. Como

se ha señalado previamente, raíces que estén próximas a la unidad pueden sugerir una posible ausencia de estacionariedad y/o invertibilidad en el modelo (Miranda Chinlli, 2021).

Por otro lado, para comprobar si los residuos son ruido blanco, se evalúa la normalidad y se verifica que tengan media nula, varianza constante y ausencia de correlación serial (García Soler, 2018). El estadístico más utilizado para contrastar la normalidad de los residuos en modelos econométricos y de series temporales es el test propuesto por Jarque y Bera (1980), que combina la información sobre la asimetría y la curtosis de los residuos para determinar si se desvían significativamente de la distribución normal.

Para contrastar la nulidad de la media de los residuos frente a la hipótesis alternativa de no nulidad, se emplea el estadístico clásico t . De manera visual, también se puede recurrir al gráfico de residuos en función del tiempo, donde lo ideal es observar que los valores fluctúan alrededor del cero (Ramón Escolano y López Espín, 2016).

En cuanto a la varianza constante, además de utilizar el gráfico de los residuos para observar si la varianza permanece constante, se destacan la prueba ARCH y el test de Goldfeld-Quandt, que permiten contrastar formalmente la homocedasticidad. El test de heterocedasticidad ARCH es una prueba combinada del multiplicador de Lagrange (LM) para errores autorregresivos condicionales heterocedásticos (ARCH) en modelos vectoriales autorregresivos (VAR) para evaluar si existe heterocedasticidad condicional en los residuos, es decir, si la varianza de los residuos depende de sus valores pasados (Catani y Ahlgren, 2017). Por otro lado, Goldfeld y Quandt (1965) presentan un método de prueba que analiza la heterocedasticidad en un modelo de regresión al examinar si la suma de cuadrados del primer tercio de la muestra es significativamente diferente de la suma de cuadrados del último tercio de la muestra. Aunque estas pruebas fueron diseñadas originalmente para modelos VAR y de regresión, respectivamente, también se aplican a modelos univariados ARIMA.

Por último, para probar la incorrelación serial en los residuos del modelo, es común utilizar la prueba de Ljung y Box (1978), una generalización de la prueba de Box y Pierce. Esta prueba calcula un estadístico basado en los cuadrados de las autocorrelaciones de los residuos y lo compara con la distribución chi-cuadrado para determinar si hay autocorrelación significativa en la serie temporal (Gujarati y Porter, 2009).

En caso de contar con más de un modelo potencial, se utilizan criterios de

información para determinar cuál es el más adecuado. Estos criterios evalúan la calidad de ajuste del modelo a la vez que penalizan la complejidad para evitar el sobreajuste. Entre los criterios de información más comunes se encuentra el Criterio de Información de Akaike (1974), AIC, que examina tanto la complejidad de un modelo, en términos del número de parámetros, como su capacidad para ajustarse a los datos. Un modelo con más parámetros puede ajustarse muy bien a los datos actuales, pero corre el riesgo de sobreajustarse y tener un rendimiento pobre en datos futuros, mientras que un modelo con menos parámetros, aunque podría no ajustarse tan precisamente a los datos actuales, tiene una mayor capacidad para generalizar a nuevos datos al ser menos susceptible al sobreajuste (Rodríguez-Marín, 2024).

Predicción

Una vez validado y confirmada su adecuación con el conjunto de entrenamiento, el modelo se utiliza para prever datos futuros, específicamente sobre el conjunto de datos de prueba que no se ha utilizado previamente (Cáceres Hernández y cols., 2008).

Roque García (2021) y Mauricio (2017) señalan que en un modelo ARIMA, considerando T como el origen de predicción (el último valor disponible) y l como el horizonte de predicción (el número de períodos hacia adelante para los que se realiza la predicción), la predicción puntual de Y_{T+l} , denotada como $Y_T(l)$, viene dada por su valor esperado condicionado a toda la información observada hasta el momento T :

$$\hat{Y}_{T+l} = Y_T(l) = E[Y_{T+l}/Y_1, Y_2, \dots, Y_T] = E_T[Y_{T+l}] \quad (4.13)$$

Esto implica que la predicción para Y_{T+l} se basa en toda la información disponible hasta el momento T . Considerando, por tanto, un modelo general $ARMA(p, q)$, las predicciones siguen la siguiente expresión:

$$Y_{T+l} = \delta + \phi_1 Y_{T+l-1} + \dots + \phi_p Y_{T+l-p} - \theta_1 a_{T+l-1} - \dots - \theta_q a_{T+l-q} + a_{T+l} \quad (4.14)$$

De esta forma, considerando como predicción óptima aquella que minimiza el error cuadrático medio de predicción (ECMP), se busca una predicción a horizonte l , $Y_T(l)$, tal que la esperanza del cuadrado del error de predicción sea mínima (de la Fuente Fernández, 2013).

Según Hyndman y Athanasopoulos (2021), en esta línea, se identifican dos

enfoques principales para la generación de predicciones futuras, delineados por la elección del parámetro del horizonte de predicción, l : las predicciones unietapa (*one-step-ahead predictions*) y las predicciones multietapa (*multi-step-ahead predictions*).

En el enfoque unietapa ($l = 1$), las predicciones se limitan al período inmediatamente posterior al origen de predicción, T , actualizando la información disponible con la observación más reciente cada vez que se realiza una predicción. En contraste, el enfoque multietapa ($l > 1$) se centra en las predicciones a períodos más distantes, por lo que la información disponible no se actualiza con las observaciones más recientes. Dentro de las predicciones multietapa, se distinguen dos estrategias: la estrategia recursiva, donde se predice un solo paso hacia adelante y esta predicción se incorpora como información para la siguiente, repitiéndose este proceso hasta alcanzar el horizonte de predicción deseado; y la estrategia directa, que implica la creación de l modelos, uno para cada etapa de predicción, sin considerar las predicciones anteriores (Ben Taieb, Bontempi, Atiya, y Sorjamaa, 2012). Elegir la estrategia de predicciones multietapa recursivas puede ser beneficioso en escenarios donde se requiere prever múltiples pasos hacia el futuro. Al actualizar el modelo con las predicciones más recientes, esta estrategia permite simular un entorno realista donde las predicciones se hacen secuencialmente y se utilizan como entrada para los pasos siguientes.

Ahora bien, tan importante como realizar estas predicciones es evaluar la capacidad predictiva del modelo. Por ello, se utilizan los datos de prueba en esta etapa, permitiendo comparar los pronósticos con los valores reales y verificar la precisión del modelo.

En toda predicción habrá un margen de error, definido como la discrepancia entre el valor observado Y_{T+l} y la predicción realizada $Y_T(l)$:

$$e_T(l) = Y_{T+l} - Y_T(l) \quad (4.15)$$

Estos errores de predicción se utilizan en diversas métricas para cuantificar la precisión de las predicciones, de forma que valores cercanos a cero sugieren una alta capacidad predictiva del modelo (González Casimiro, 2009). Se clasifican en dos categorías según se basen en errores absolutos, $|\hat{e}_T|$, o cuadráticos, $\hat{e}_T^2$, y se destacan métricas que promedian estos errores, como el Error Absoluto Medio y la Raíz del Error Cuadrático Medio:

- Error Absoluto Medio (MAE):

$$MAE = \frac{1}{L} \sum_{l=1}^L |\hat{e}_T(l)| \quad (4.16)$$

- Raíz del Error Cuadrático Medio (RMSE):

$$RMSE = \sqrt{\frac{1}{L} \sum_{l=1}^L (\hat{e}_T(l))^2} \quad (4.17)$$

Es importante destacar que estas métricas son sensibles a la unidad de medida de la serie temporal. Por lo tanto, si se intentara comparar series temporales con unidades diferentes, estas métricas no serían adecuadas y sería necesario recurrir a alternativas adimensionales (Martí Sanahuja, 2021), como el Error Porcentual Absoluto Medio (MAPE):

$$MAPE = \frac{1}{L} \sum_{l=1}^L \left| \frac{\hat{e}_T(l)}{Y_{T+l}} \right| \times 100 \quad (4.18)$$

4.2. Aprendizaje Automático

El análisis de series temporales de índices bursátiles se integra dentro del aprendizaje automático como una aplicación específica de los algoritmos de regresión propios del aprendizaje supervisado. Este análisis se centra en la predicción de valores futuros basándose en datos históricos estableciendo relaciones entre las variables temporales (Nieto Benito, 2023).

A menudo, se dispone únicamente de una variable predictora: la propia serie temporal. Sin embargo, es posible enriquecer el modelo considerando los retardos (o *lags*) de la serie temporal como variables predictoras adicionales. Este enfoque permite capturar la dependencia temporal y mejorar la capacidad predictiva del modelo. Para lograr esto, se utilizan diversos métodos, entre los que destaca el método de ventanas deslizantes (*sliding window*). Este método implica crear una ventana fija de tamaño n que se desplaza a lo largo de la serie temporal, generando subconjuntos de datos para entrenamiento y prueba. Cada ventana incluye un conjunto de valores pasados que se utilizan como predictores para el valor futuro (Ayala Castrejón y Bucio Pacheco, 2020).

El parámetro `look.back` es un hiperparámetro adicional fundamental en

modelos de series temporales que utilizan ventanas deslizantes. Este parámetro especifica la cantidad de pasos temporales anteriores que se consideran para realizar predicciones. En esencia, determina cuántos valores históricos de la serie temporal se emplean como variables predictoras para estimar los valores futuros. Al ajustar este parámetro, se puede influir en la precisión y el rendimiento del modelo, ya que define la longitud de la ventana de observaciones pasadas utilizadas para hacer las predicciones (Hyndman y Athanasopoulos, 2021). Valores más bajos pueden ser útiles para capturar patrones a corto plazo, donde las relaciones entre los datos más recientes son más relevantes. En cambio, valores más altos permiten al modelo capturar patrones a largo plazo y tendencias más amplias en los datos, lo cual puede ser crucial si los eventos pasados tienen un impacto significativo en el presente.

Similar a los modelos clásicos ARIMA, donde se optimizan los parámetros autorregresivos y de media móvil para lograr el mejor ajuste, en los algoritmos de aprendizaje automático (ML) también se busca la combinación óptima de hiperparámetros para obtener el máximo rendimiento del modelo. Los hiperparámetros son configuraciones que se establecen antes del proceso de entrenamiento del modelo y no se ajustan durante el entrenamiento (Sotaquirá, 2023). Ejemplos incluyen la profundidad de los árboles en los bosques aleatorios, la tasa de aprendizaje en los métodos de *boosting*, y el parámetro de regularización en los algoritmos *Support Vector Regression* (SVR).

El desarrollo de un algoritmo de aprendizaje automático sigue una serie de pasos estructurados que aseguran la efectividad y la precisión del modelo. Una vez identificado el problema, seleccionado el algoritmo de aprendizaje más adecuado para su resolución, y dividido el *dataset* en conjunto de entrenamiento y de prueba, se entrena el modelo utilizando el primer conjunto. Durante el entrenamiento, se ajustan los hiperparámetros para optimizar el rendimiento del modelo, utilizando técnicas como la validación cruzada si es necesario. Una vez entrenado, el modelo se evalúa con el conjunto de prueba para medir su rendimiento en datos nuevos, utilizando métricas adecuadas. Según Díez García (2020), en problemas de regresión, es común evaluar los errores de predicción y métricas derivadas, como la Raíz del Error Cuadrático Medio (RMSE).

La correcta selección y ajuste de los hiperparámetros se realiza mediante técnicas de optimización, siendo la búsqueda en cuadrícula (*Grid Search*) una de las más destacadas. Esta técnica implica definir un espacio de búsqueda discreto para los hiperparámetros y evaluar exhaustivamente todas las combinaciones posibles (Mesafint Belete y Huchaiah, 2021).

A diferencia de los métodos tradicionales de series temporales, que requieren estacionariedad para funcionar correctamente, los algoritmos de aprendizaje supervisado pueden detectar relaciones no lineales y complejas en los datos sin necesidad de preprocesar la serie para hacerla estacionaria. Estos algoritmos parten de la idea de que existe una función subyacente compleja que genera los datos, y a través del entrenamiento se intenta aproximar esta función. Por tanto, no es estrictamente necesario convertir la serie temporal en estacionaria antes de modelar. Al transformar la serie, es posible que se pierdan detalles importantes de los datos originales. Esto puede reducir la capacidad del modelo para capturar patrones complejos y variaciones en los datos.

En este contexto, se eligen para la tarea de predicción los algoritmos de *Gradient Boosting* y *Support Vector Machine*, aplicados a problemas de regresión.

4.2.1. Gradient Boosting

El *Gradient Boosting* es una técnica perteneciente al conjunto de métodos de ensamblaje conocidos como *boosting*, los cuales combinan varios modelos débiles (*weak learners*), comúnmente árboles de decisión, para formar un modelo más robusto (*strong learner*), con el propósito de mejorar el rendimiento y la precisión de las predicciones. La idea central radica en entrenar los modelos débiles de manera secuencial, de modo que cada nuevo modelo trate de corregir los errores cometidos por los modelos anteriores, buscando minimizar la suma de los errores cuadráticos acumulados hasta la iteración anterior (Camacho, Ramallo, y Ruiz Marín, 2021).

Los modelos de *Gradient Boosting* comienzan creando un único nodo hoja y construyendo árboles de regresión de manera iterativa. Un árbol de regresión es un tipo de árbol de decisión diseñado para estimar una función continua de valores reales en lugar de un clasificador. En cada iteración, el árbol de regresión se construye dividiendo los datos en nodos o ramas, segmentando los datos en grupos más pequeños para minimizar el error de predicción (Yoon, 2021).

Presentado por Friedman (2001), el *Gradient Boosting* es un algoritmo para regresión que utiliza el método de optimización numérica *Functional Gradient Descent* para iterativamente buscar el mínimo de la función de pérdida. En cada iteración, se determina la dirección y magnitud de la actualización que maximiza la reducción de la función de pérdida, encontrando el *weak learner* que se acerque más al negativo del gradiente de la función de pérdida y ajustando el tamaño de la actualización para optimizar la disminución de la función de pérdida en esa dirección.

Esto es, dado un conjunto de datos de entrenamiento $D = \{x_i, y_i\}_{i=1}^N$, el objetivo es encontrar una aproximación $\hat{F}(x)$ de la función $F^*(x)$, que mapea las instancias x a sus valores de salida y , minimizando el valor esperado de una función de pérdida dada, $L(y, F(x))$.

Según [Bentéjac, Csörgő, y Martínez-Muñoz \(2021\)](#), el proceso iterativo comienza obteniendo una aproximación constante de $F^*(x)$ como:

$$F_0(x) = \arg \min_{\alpha} \sum_{i=1}^N L(y_i, \alpha) \quad (4.19)$$

donde L es la función de pérdida, y_i son los valores observados y α es la constante que minimiza la pérdida. Se espera que los modelos subsecuentes minimicen:

$$(\rho_m, h_m(x)) = \arg \min_{\rho, h} \sum_{i=1}^N L(y_i, F_{m-1}(x_i) + \rho h(x_i)) \quad (4.20)$$

Sin embargo, en lugar de resolver esto directamente, cada h_m se entrena en un nuevo conjunto de datos $D = \{x_i, r_{mi}\}_{i=1}^N$, donde los pseudo-residuos se calculan como:

$$r_{mi} = \left[\frac{\partial L(y_i, F(x))}{\partial F(x)} \right]_{F(x)=F_{m-1}(x)} \quad (4.21)$$

El *Gradient Boosting* construye una aproximación aditiva de $F^*(x)$ como una suma ponderada, siendo $F_{m-1}(x)$ el modelo acumulado hasta la iteración $m - 1$, $h_m(x)$ el nuevo modelo débil añadido en la iteración m y ρ_m la tasa de aprendizaje que controla la contribución de este nuevo modelo:

$$F_m(x) = F_{m-1}(x) + \rho_m h_m(x) \quad (4.22)$$

Además de la tasa de aprendizaje (`learning_rate`), los principales hiperparámetros a considerar en los modelos *Gradient Boosting* son la tasa de muestreo (`subsample`), que controla la proporción de muestras utilizadas en cada árbol para reducir el sobreajuste; el número de árboles (`n_estimators`), donde más árboles generalmente mejoran el rendimiento, pero pueden aumentar el riesgo de sobreajuste; la función de pérdida (`loss`), que evalúa la calidad de las predicciones; y la profundidad máxima del árbol (`max_depth`), que ayuda a controlar el sobreajuste limitando la complejidad de cada árbol, aunque valores más altos permiten capturar relaciones más complejas ([Jain, 2022](#)).

Respecto al submuestreo, un valor más bajo puede reducir el sobreajuste al introducir variabilidad en los árboles individuales, pero también puede aumentar el ruido (Pérez García, 2018). En cuanto al número de estimadores, un mayor número de árboles generalmente mejora el rendimiento del modelo, pero también incrementa el tiempo de entrenamiento y el riesgo de sobreajuste. En lo que se refiere a la tasa de aprendizaje, una tasa más baja puede requerir más estimadores para alcanzar un rendimiento óptimo, pero puede mejorar la generalización, mientras que tasas más altas pueden acelerar el entrenamiento, aunque conllevan un mayor riesgo de sobreajuste (Amat Rodrigo, 2023).

En términos de la función de pérdida, la elección depende del tipo de problema y la robustez deseada frente a valores atípicos en los datos. El error cuadrático es sensible a grandes errores pero vulnerable a valores atípicos extremos debido a su penalización cuadrática, mientras que el error absoluto es menos sensible a valores atípicos, ya que penaliza los errores de manera lineal, pero menos eficiente en términos de convergencia, lo que puede requerir más iteraciones o una tasa de aprendizaje más baja (Pérez García, 2018). Por su parte, la *Huber Loss* combina ambas características para ofrecer robustez ajustable mediante el parámetro δ , que controla la transición entre la pérdida cuadrática y absoluta. No obstante, la selección de este parámetro puede requerir ajuste y no siempre es trivial determinar el valor óptimo para diferentes conjuntos de datos y problemas.

Adicionalmente, se fija una semilla aleatoria para asegurar la reproducibilidad de los resultados, lo que permite realizar comparaciones consistentes entre diferentes configuraciones de modelos y evitar sesgos debido a variaciones aleatorias.

4.2.2. Support Vector Regression (SVR)

Desarrollado por Drucker, Burges, Kaufman, Smola, y Vapnik (1996) y Vapnik (1998), el *Support Vector Regression* (SVR) es una variante específica de los algoritmos de *Support Vector Machine* (SVM), comúnmente utilizados para tareas de clasificación, adaptada para problemas de regresión.

Para comprender el funcionamiento de SVR, es esencial entender primero cómo opera la clasificación SVM. En la clasificación SVM, se trata un conjunto de datos de entrenamiento, donde cada muestra etiquetada se representa como un punto en un espacio de características multidimensional. El objetivo es encontrar un hiperplano en este espacio que separe correctamente las clases de muestras de entrenamiento. Posteriormente, las nuevas muestras se clasifican según la región

del espacio multidimensional en la que se encuentran con respecto al hiperplano. La búsqueda de este hiperplano óptimo implica maximizar el margen entre los vectores de soporte, que son los puntos de datos más cercanos al hiperplano (Pisner y Schnyer, 2019).

Para la regresión, en lugar de buscar un hiperplano que pueda separar las clases de manera óptima, SVR introduce una función de pérdida insensible a los errores para calcular un hiperplano de tal manera que las respuestas predichas de las muestras de entrenamiento se desvíen como máximo en una cantidad ϵ de sus valores observados. La optimización se lleva a cabo minimizando una banda insensible a ϵ para que sea lo más plana y estrecha posible, y contenga la mayoría de las muestras de entrenamiento (Zhang y O'Donnell, 2019).

En este contexto, el hiperplano se define en términos de unos pocos vectores de soporte que representan las muestras de entrenamiento que se encuentran fuera del límite de la banda insensible a ϵ . Estos puntos determinan la posición y la orientación de la función de regresión.

Matemáticamente, en SVR, se busca una función $f(x) = w^T x + b$ que minimice:

$$\frac{1}{2} ||w||^2 + C \sum_{i=1}^n L_{\epsilon}(y_i, f(x_i)) \quad (4.23)$$

donde L_{ϵ} es la función de pérdida ϵ -insensible, definida como:

$$L_{\epsilon}(y, f(x)) = \begin{cases} 0 & \text{si } |y - f(x)| \leq \epsilon \\ |y - f(x)| - \epsilon & \text{en caso contrario} \end{cases} \quad (4.24)$$

Como espacio de hiperparámetros, se presta especial atención a los parámetros kernel y C . El kernel define la función de núcleo que transforma los datos de entrada en un espacio de características de mayor dimensión, facilitando la captura de relaciones no lineales entre las variables de entrada y la variable objetivo. La elección del kernel es fundamental porque determina cómo se transforman los datos y, en última instancia, la capacidad del modelo para capturar patrones complejos (Bishop, 2006). Las opciones de funciones de núcleo incluyen el kernel lineal, polinómico, de función de base radial (RBF) y sigmoidal.

El kernel lineal es simple y efectivo para datos linealmente separables, pero carece de la capacidad para manejar complejidades no lineales. En contraste, el kernel polinómico amplía esta capacidad modelando interacciones más complejas a través de un espacio dimensional superior, aunque con un coste computacional

mayor y un riesgo elevado de sobreajuste si se elige un grado polinómico demasiado alto. Por otro lado, el kernel radial es altamente flexible y adecuado para la mayoría de los problemas, ya que puede manejar relaciones no lineales complejas, aunque requiere un ajuste cuidadoso del parámetro gamma para evitar sobreajuste. Por su parte, el kernel sigmooidal, similar a una función de activación en redes neuronales, puede modelar relaciones no lineales, pero es menos utilizado debido a su inestabilidad y necesidad de ajuste preciso de parámetros (Martín Guareño, 2016).

Por otro lado, el parámetro C controla el *trade-off* entre la minimización del error en el entrenamiento y la simplicidad del modelo (regularización). Un valor bajo permite un margen más amplio, dando lugar a modelos más simples que pueden generalizar mejor. En contraste, valores altos reducen el margen, haciendo que el modelo ajuste el entrenamiento, pero con un mayor riesgo de sobreajuste (James, Witten, Hastie, y Tibshirani, 2013).

A pesar de que en los modelos de *Support Vector Regression* (SVR) en Python no existe un hiperparámetro específico para fijar la semilla aleatoria, es posible asegurar la reproducibilidad de las funciones que dependen del estado aleatorio de scikit-learn fijando manualmente la semilla. Esto ayuda a obtener resultados consistentes en cada ejecución del código.

5. Aplicación práctica

5.1. Software

El análisis de series temporales de índices bursátiles puede ser abordado mediante diversos lenguajes de programación y software especializados, destacando entre ellos Python, R y Gretl.

Gretl es un software estadístico eficaz para análisis econométricos y modelos de regresión. Sin embargo, debido a su simplicidad, puede resultar limitado en comparación con lenguajes más versátiles como Python o R. Estos lenguajes ofrecen amplias librerías y herramientas que facilitan la descarga, manipulación de datos y la implementación de técnicas estadísticas avanzadas.

Como experiencia académica, el uso de R ha sido predominante, lo cual ha proporcionado una sólida base en el análisis estadístico y en el manejo de grandes volúmenes de datos. No obstante, en un esfuerzo por expandir habilidades y enfrentar nuevos desafíos profesionales, se ha optado por utilizar Python. Este lenguaje, a pesar de ser menos familiar durante la formación académica, representa una oportunidad valiosa para fortalecer competencias y habilidades estadísticas en un contexto diferente.

5.2. Datos reales

Para llevar a cabo un análisis exhaustivo del comportamiento de los mercados financieros, se ha elegido un grupo representativo de índices bursátiles:

- Nikkei 225 de Japón,
- IBEX 35 de España,
- S&P 500 de Estados Unidos,
- IBOVESPA de Brasil.

Estos índices son esenciales para comprender la dinámica de las principales economías globales y con esta elección se busca obtener una visión completa y

variada sobre las tendencias y el rendimiento de las economías más relevantes a nivel mundial.

La selección del periodo de estudio, que abarca desde el año 2000 hasta el 2023, está justificada por varias razones estratégicas y analíticas. En primer lugar, disponer de más de dos décadas de información proporciona un volumen de datos considerable, esencial para llevar a cabo evaluaciones estadísticas detalladas y modelar series temporales de manera efectiva. En segundo lugar, este intervalo incluye eventos económicos importantes que han tenido un impacto significativo en los mercados financieros globales, como la crisis financiera mundial de 2007-2008 y la pandemia de COVID-19 en 2020, lo que hace que este período sea particularmente relevante y atractivo para el análisis.

Para la recopilación de los datos se ha utilizado *Yahoo Finance*, una fuente ampliamente utilizada debido a su extensa base de datos y la facilidad de acceso a la información financiera. La plataforma ofrece una manera práctica de descargar estos datos a través de su Interfaz de Programación de Aplicaciones (API, por sus siglas en inglés), permitiendo una recolección eficiente y organizada de datos históricos de índices bursátiles. La información histórica de precios que ofrece Yahoo Finance no solo incluye precios de cierre, sino también los precios máximos y mínimos alcanzados durante cada sesión de negociación, así como los volúmenes de acciones negociadas, entre otros, proporcionando una imagen completa de la actividad del mercado. Así, mediante la función `download` de la biblioteca `yfinance`, se obtienen las siguientes variables:

- **Open (Apertura):** Precio de apertura del primer día de negociación del mes. Representa el primer precio registrado para ese mes, proporcionando una indicación inicial del valor de mercado del índice al comienzo del período.
- **High (Máximo):** Precio más alto alcanzado por el índice durante el mes. Muestra el punto máximo de optimismo o disposición a pagar entre los inversores durante ese período mensual.
- **Low (Mínimo):** Precio más bajo registrado durante el mes. Refleja el punto más bajo de valoración que los inversores estaban dispuestos a aceptar, indicando los niveles de soporte o el mínimo interés por el índice durante el período.
- **Close (Cierre):** Precio de cierre del último día de negociación del mes. Es crucial porque refleja el consenso final de valor entre los compradores y vendedores al cierre del mes y es a menudo utilizado como un punto de referencia para el análisis técnico y la valoración del desempeño mensual.

- **Adj Close (Cierre Ajustado):** Precio de cierre ajustado ajustado para reflejar cualquier cambio debido a dividendos pagados, derechos de suscripción, y divisiones de acciones durante el mes.
- **Volume (Volumen):** Total de unidades de fondos negociados durante el mes. El volumen puede dar una indicación de la actividad o liquidez del mercado durante ese periodo. En el caso de los índices, este número puede ser representativo de los movimientos subyacentes en las acciones que forman parte del índice.

De estas, el cierre ajustado es la variable preferida y seleccionada, ya que proporciona una imagen más precisa del valor de los índices y permite evaluar el rendimiento a largo plazo al reflejar la realidad del valor ajustado a lo largo del tiempo.

5.3. Caso práctico

Una vez descargado el conjunto de datos y definidas algunas funciones útiles para el análisis de series temporales (véase el anexo I), se han aplicado las tres metodologías descritas en el capítulo anterior. Esto ha incluido tanto la búsqueda manual de las combinaciones óptimas de parámetros como la utilización de funciones automatizadas disponibles. El objetivo es evaluar inicialmente el rendimiento de las funciones automáticas y luego comparar entre las distintas metodologías para determinar el modelo óptimo final.

Los tres métodos comparten aspectos comunes como la división del conjunto de datos en entrenamiento y prueba, lo cual es fundamental para validar los modelos de series temporales. El conjunto de entrenamiento se emplea para ajustar los parámetros del modelo, permitiendo que este aprenda las características y patrones subyacentes en los datos históricos. Posteriormente, el conjunto de prueba se utiliza para evaluar el rendimiento del modelo, asegurando que este no solo se ajuste a los datos históricos, sino que también tenga la capacidad de generalizar a datos futuros.

Modelos ARIMA

Como ya se ha explicado, la metodología de Box-Jenkins comienza con la fase de identificación, y más concretamente con el análisis de la estacionariedad. Observando el gráfico de series temporales, se aprecia una clara falta de

estacionariedad tanto en media, ya que los ciclos no oscilan alrededor de un valor constante y muestran tendencias, como en varianza, dado que la amplitud de los ciclos no es uniforme. Por tanto, se aplica una transformación logarítmica para abordar la no estacionariedad en la varianza y una diferenciación para hacer la serie estacionaria en media.

En la figura 5.1 se presentan las series logarítmicas diferenciadas divididas en conjuntos de entrenamiento y prueba, donde se puede observar una evidente estacionariedad. Además, mediante la prueba de raíz unitaria de Dickey-Fuller Aumentada (ADF), se confirma formalmente que las series transformadas son estacionarias y que, por tanto, se trata de modelos $ARIMA(p, 1, q)$.

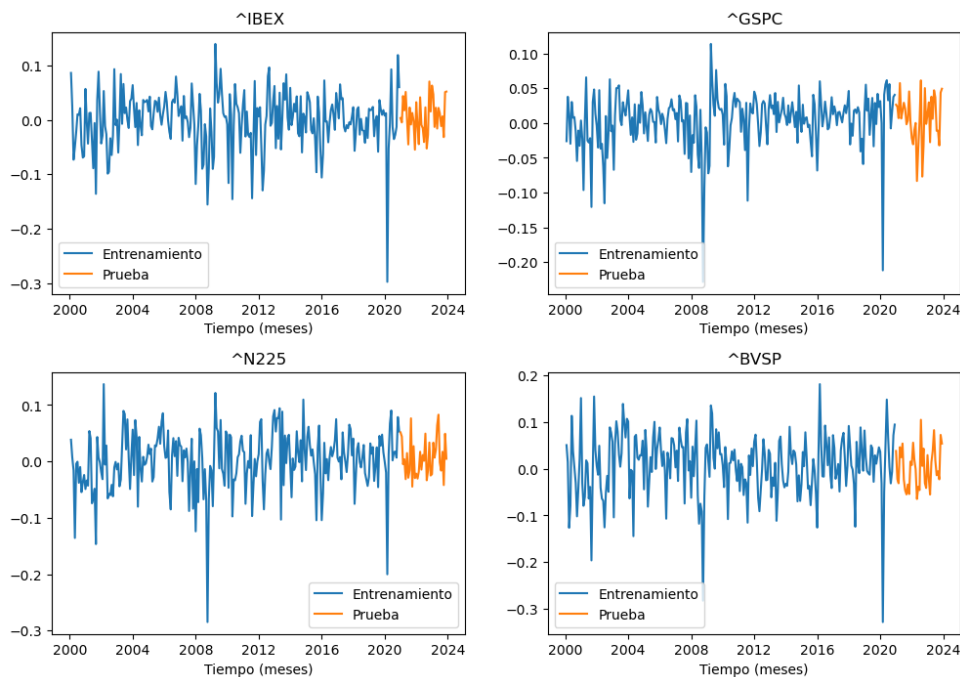


Figura 5.1: Series logarítmicas diferenciadas divididas en conjunto de entrenamiento y prueba

Por otro lado, se ha investigado la componente estacional tanto visualmente, a través del periodograma, como matemáticamente, mediante el test de Canova-Hansen. Ambos métodos han indicado la ausencia de estacionalidad, lo que ha permitido descartar la necesidad de diferenciaciones estacionales y, por consiguiente, los modelos SARIMA.

A continuación, se han analizado las funciones de autocorrelación simple y parcial, como se muestra en la figura 5.2 para seleccionar los órdenes autorregresivos y de medias móviles (véase el anexo II para los otros índices

bursátiles).

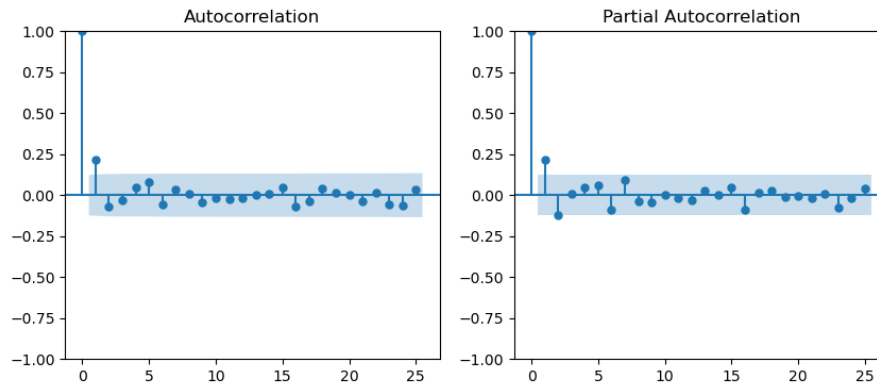


Figura 5.2: Correlograma de la serie transformada del índice IBEX 35

Se han identificado varios modelos para compararlos posteriormente y seleccionar el más adecuado. Para cada índice bursátil, se ha fijado un máximo de p y q , evaluando todas las combinaciones posibles de órdenes inferiores. De esta manera, para los índices IBEX 35 y IBOVESPA, se han considerado combinaciones de órdenes inferiores a $p = 2$ y $q = 1$. Para los índices S&P 500 y Nikkei 225, se han probado las combinaciones de órdenes inferiores a $p = 1$ y $q = 1$. Finalmente, para concluir la fase de identificación, se ha determinado que no es necesario incluir el término constante δ en el modelo.

Una vez identificados todos los posibles modelos para cada índice, se ha procedido a ajustarlos utilizando la librería `statsmodels` y a validarlos posteriormente. Esta etapa de validación ha permitido descartar aquellos modelos que no cumplen con las condiciones necesarias respecto a los coeficientes y los residuos. Por último, se ha estudiado cuál es el modelo óptimo entre los propuestos, seleccionando aquel con el menor Criterio de Información AIC.

Para comparar los modelos obtenidos manualmente con el desempeño de su correspondiente función automática, se ha utilizado la función `auto_arima` de la librería `pmdarima`. Esta función ha sugerido los mismos modelos identificados de forma manual: $ARIMA(0, 1, 1)$ para los índices IBEX 35, S&P 500 e IBOVESPA, y $ARIMA(1, 1, 0)$ para el índice Nikkei 225.

En último lugar, se han realizado las predicciones del conjunto de prueba utilizando la estrategia de predicciones multietapa recursivas, y se han calculado las métricas de precisión. Al examinar la tabla [5.1](#), que detalla estos resultados, se observa que los modelos ARIMA muestran una variabilidad considerable en su rendimiento entre diferentes mercados. En general, los modelos ARIMA

demuestran una precisión predictiva insuficiente en todos los mercados evaluados. Si bien el S&P 500 exhibe mejores resultados en términos de errores absolutos, el desempeño es menos eficiente para el Nikkei 225 y el IBOVESPA.

	MAE	RMSE	MAPE
IBEX 35	647.08	795.70	7.11
S&P 500	493.34	563.14	11.33
Nikkei 225	1893.44	2673.49	6.15
IBOVESPA	7369.98	8939.05	6.72

Tabla 5.1: Métricas de precisión de los modelos ARIMA

La figura 5.3 muestra la comparación entre los valores observados en el conjunto de prueba (en azul) y las predicciones realizadas por modelos (en naranja) para los cuatro índices bursátiles. Las líneas de predicción, consistentemente horizontales para todos los índices, sugieren que los modelos utilizados no son adecuados para capturar las fluctuaciones, tendencias y cambios dinámicos que caracterizan a los datos reales de los mercados bursátiles.

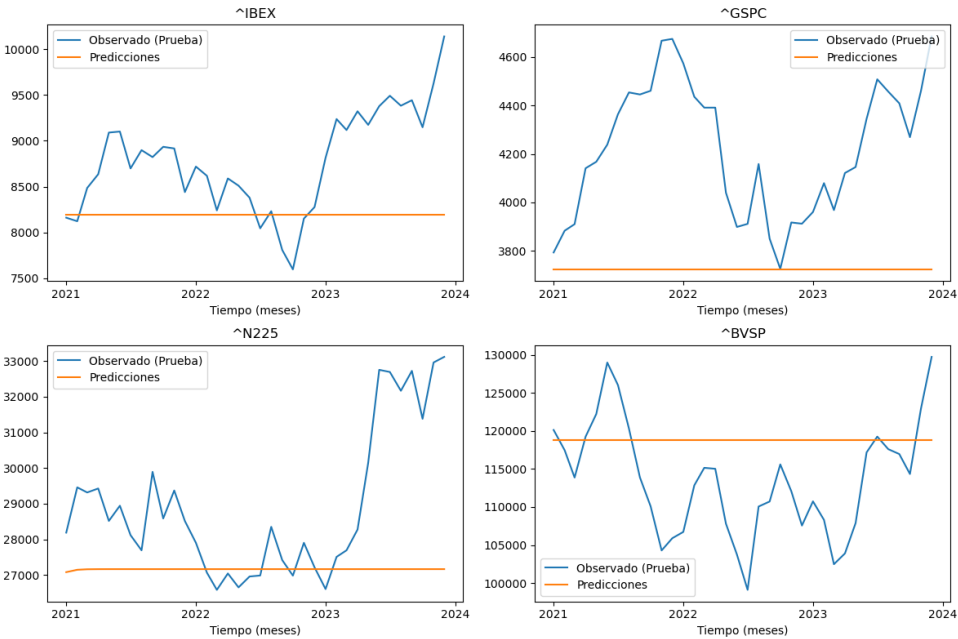


Figura 5.3: Predicciones sobre el conjunto de prueba de los modelos ARIMA

Gradient Boosting

En la configuración de *Gradient Boosting*, los hiperparámetros desempeñan un papel fundamental en la configuración y optimización del modelo. Cada conjunto

diferente de valores que los hiperparámetros pueden tomar influirá en la complejidad del algoritmo utilizado. En otras palabras, la forma en que se configuren los hiperparámetros afectará directamente la forma en que funciona el modelo o algoritmo, haciendo que sea más o menos complejo en términos de rendimiento, ajuste a los datos, capacidad de generalización, entre otros aspectos. La tabla 5.2 presenta los distintos valores que se han evaluado para cada hiperparámetro.

Hiperparámetro	Valores
subsample	[0.5, 0.75, 1.0]
n_estimators	[100, 200, 300]
learning_rate	[0.01, 0.05, 0.1, 0.2]
random_state	[32]
loss	['squared_error', 'absolute_error', 'huber']
max_depth	[3, 5, 7]
look_back	[6, 12, 18, 24, 30, 36]

Tabla 5.2: Posibles valores de hiperparámetros en *Gradient Boosting*

Una vez establecido el rango de posibles valores para la optimización de hiperparámetros, se ha buscado la combinación óptima de estos, primero de manera manual y luego utilizando la función automática GridSearchCV de la biblioteca scikit-learn.

El proceso manual ha sido el siguiente: tras definir una cuadrícula con los diferentes valores posibles para cada hiperparámetro del modelo, se ha iterado sobre todas las combinaciones posibles, entrenando y validando el modelo con diferentes tamaños de ventana deslizante para predecir los datos. Así, se ha seleccionado la combinación de hiperparámetros que ha resultado en el menor error cuadrático medio (véase el anexo III).

De nuevo, habiendo encontrado la combinación óptima de hiperparámetros en cada caso utilizando el conjunto de entrenamiento, se han realizado las predicciones del conjunto de prueba mediante la estrategia multietapa recursiva, obteniendo las métricas de precisión que se muestran en las tablas 5.3 y 5.4.

	MAE	RMSE	MAPE
IBEX 35	251.25	310.93	2.91
S&P 500	434.19	508.94	9.93
Nikkei 225	2500.06	3212.24	8.22
IBOVESPA	6051.89	7601.60	5.25

Tabla 5.3: Métricas de precisión de los modelos *Gradient Boosting* manuales

	MAE	RMSE	MAPE
IBEX 35	310.34	388.20	3.60
S&P 500	482.50	552.15	11.07
Nikkei 225	5800.36	6239.52	19.72
IBOVESPA	8774.14	10672.77	7.46

Tabla 5.4: Métricas de precisión de los modelos *Gradient Boosting* automáticos

En todas las métricas y para todos los índices bursátiles evaluados, los modelos de *Gradient Boosting* ajustados manualmente exhiben un desempeño superior en comparación con aquellos ajustados automáticamente. Los modelos automáticos registran errores más elevados en todas las medidas evaluadas - errores absolutos medios, errores cuadráticos medios y errores porcentuales medios absolutos - siendo particularmente notable en el IBOVESPA, donde el RMSE es significativamente alto.

La figura 5.4, que muestra las predicciones sobre el conjunto de prueba de los modelos de Gradient Boosting para los cuatro índices bursátiles, facilita la comparación de las predicciones hechas por los modelos encontrados de forma manual y automática frente a los valores observados. Se evidencia que los modelos presentan más dificultades a la hora de capturar los patrones y hacer predicciones precisas para el S&P 500 y el Nikkei 225. Por otro lado, ambos modelos del IBEX 35 muestran un buen desempeño en sus predicciones.



Figura 5.4: Predicciones sobre el conjunto de prueba de los modelos *Gradient Boosting*

Support Vector Regression (SVR)

Al igual que en los modelos *Gradient Boosting*, la determinación de los hiperparámetros es esencial para ajustar y optimizar el modelo de *Support Vector Regression* (SVR). En particular, la selección de parámetros como el tipo de kernel y el parámetro de regularización C tienen un impacto significativo en el rendimiento del modelo y su habilidad para generalizar.

El proceso general es el mismo seguido para los algoritmos *Gradient Boosting*: se ha definido el rango de los posibles diferentes valores para los hiperparámetros (véase la tabla 5.5) y se ha entrenado y validado de forma iterativa cada modelo hasta encontrar la combinación de hiperparámetros que produce el menor error cuadrático medio. Esto ha permitido una evaluación sistemática y eficiente de todas las combinaciones posibles, garantizando la selección del modelo más preciso. En cuanto a la función automática de búsqueda de hiperparámetros, se ha utilizado también la función `GridSearchCv` que ofrece `scikit-learn` (véase el anexo IV).

Hiperparámetro	Valores
kernel	['linear', 'poly', 'rbf', 'sigmoid']
C	[1, 5, 10]

Tabla 5.5: Posibles valores de hiperparámetros en SVR

Por último, se han realizado las predicciones del conjunto de prueba utilizando la estrategia multietapa recursiva, para cada modelo. Las métricas de precisión obtenidas se presentan en las tablas 5.6 y 5.7.

	MAE	RMSE	MAPE
IBEX 35	241.77	288.29	2.79
S&P 500	108.68	139.62	2.60
Nikkei 225	949.08	1145.09	3.24
IBOVESPA	5122.57	6354.79	4.55

Tabla 5.6: Métricas de precisión de los modelos SVR manuales

	MAE	RMSE	MAPE
IBEX 35	851.20	1098.24	9.72
S&P 500	121.10	149.54	2.89
Nikkei 225	9730.60	10593.93	33.27
IBOVESPA	20405.03	23480.73	18.60

Tabla 5.7: Métricas de precisión de los modelos SVR automáticos

Como se puede comprobar, los modelos SVR cuya combinación de hiperparámetros ha sido encontrada de forma manual muestran un MAE y RMSE relativamente bajos para el S&P 500 comparado con otros índices, lo que sugiere una mejor precisión en las predicciones de este mercado. El IBOVESPA, en cambio, muestra los valores más altos en todas las métricas, indicando una menor precisión en las predicciones para este índice. En general, el MAPE presenta valores satisfactorios en todos los casos.

Por otro lado, en los modelos SVR con hiperparámetros optimizados de forma automática, se observa una tendencia similar en términos de MAE y RMSE. Sin embargo, los valores son generalmente más altos en comparación con los modelos manuales, especialmente para el Nikkei 225 y el IBOVESPA, lo que podría sugerir una menor eficacia de la optimización automática en estos mercados. El MAPE es notablemente alto para el Nikkei 225, lo que indica un error porcentual considerable.

En la figura 5.5, de nuevo, se comparan las predicciones obtenidas mediante la configuración manual y automática de hiperparámetros contra los valores observados.

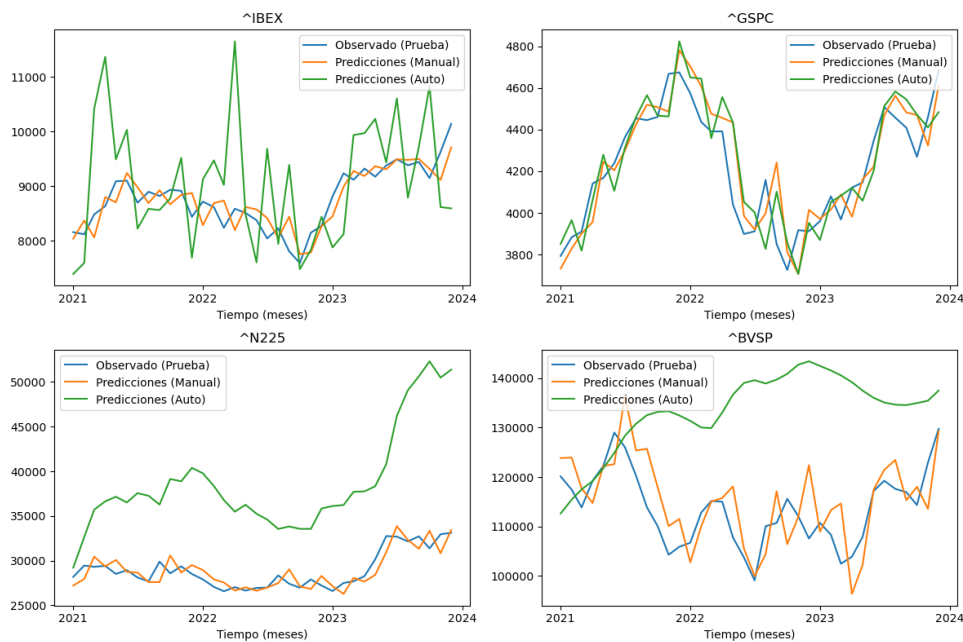


Figura 5.5: Predicciones sobre el conjunto de prueba de los modelos SVR

En cuanto al IBEX 35, las predicciones de los modelos con hiperparámetros obtenidos manualmente parecen seguir la tendencia general un poco mejor que los modelos optimizados automáticamente. El modelo manual sigue de manera aproximada los valores reales, mostrando una adecuada alineación con la

tendencia observada. Sin embargo, el modelo automático presenta una mayor variabilidad, con numerosos picos que no reflejan fielmente la trayectoria real del índice.

En el S&P 500, ambos modelos, manual y automático, muestran un desempeño similar al seguir de cerca la tendencia observada, sin diferencias significativas en su capacidad predictiva. Sin embargo, para los índices Nikkei 225 y IBOVESPA, los modelos manuales demuestran una mayor precisión al alinearse más estrechamente con los valores reales y capturar las tendencias del mercado con más exactitud. En cambio, los modelos automáticos tienden a sobreestimar los valores, posicionando sus líneas de predicción notablemente por encima de los datos observados. Esto refleja una limitación en la capacidad de estos modelos automáticos para adaptarse a las dinámicas particulares de estos mercados.

Comparación entre modelos

Para facilitar la comparación entre los modelos más efectivos para cada índice, se presentan las métricas de precisión de aquellos modelos que han demostrado ser los más precisos, cuyos hiperparámetros han sido optimizados de manera manual. Estos resultados se detallan en las tablas [5.8](#), [5.9](#), [5.10](#) y [5.11](#).

	MAE	RMSE	MAPE
ARIMA	647.08	795.70	7.11
Gradient Boosting	251.25	310.93	2.91
SVR)	241.77	288.29	2.79

Tabla 5.8: Métricas de precisión de los mejores modelos para el IBEX 35

	MAE	RMSE	MAPE
ARIMA	493.34	563.14	11.33
Gradient Boosting	434.19	508.94	9.93
SVR	108.68	139.62	2.60

Tabla 5.9: Métricas de precisión de los mejores modelos para el S&P 500

	MAE	RMSE	MAPE
ARIMA	1893.44	2673.49	6.15
Gradient Boosting	2500.06	3212.24	8.22
SVR	949.08	1145.09	3.24

Tabla 5.10: Métricas de precisión de los mejores modelos para el Nikkei 225

	MAE	RMSE	MAPE
ARIMA	7369.98	8939.05	6.72
Gradient Boosting	6051.89	7601.60	5.25
SVR	5122.57	6354.79	4.55

Tabla 5.11: Métricas de precisión de los mejores modelos para el IBOVESPA

A través de los índices, el modelo SVR consistentemente muestra un rendimiento superior, registrando los valores más bajos en todas las métricas de precisión, lo que indica su capacidad para capturar con eficacia las dinámicas del mercado. En contraste, el modelo ARIMA exhibe el peor rendimiento, con los valores más altos de MAE y RMSE, mientras que el Gradient Boosting se sitúa entre ambos, generalmente más cerca del rendimiento de SVR. Estos resultados subrayan la efectividad del SVR en la predicción precisa de movimientos de índices bursátiles bajo variadas condiciones de mercado, sugiriendo que es la opción más robusta para estos casos específicos.

6. Discusión

La predicción de series temporales en el contexto de índices bursátiles es de gran importancia para inversores, analistas y gestores de fondos, ya que estos índices son indicadores clave de la salud del mercado (Romero-Meza y cols., 2021). Sin embargo, predecir estos índices es un desafío debido a la naturaleza volátil y compleja de los mercados financieros, influenciados por factores económicos, políticos y sociales (Hernández, 2000).

Cada vez es de mayor interés evaluar y mejorar la precisión de las predicciones de índices bursátiles, debido a su impacto significativo. Por este motivo, se busca comparar el rendimiento de distintos modelos de predicción, incluyendo ARIMA, *Gradient Boosting* y *Support Vector Regression* (SVR).

Aunque los modelos ARIMA han demostrado ser útiles en otros campos, su efectividad parece limitada cuando se aplican a índices bursátiles. Si bien pueden realizar predicciones efectivas sobre los precios de las acciones en el corto plazo (Ariyo, Adewumi, y Ayo (2014), prever el mercado de valores en períodos más largos representa un desafío debido a la falta de estacionalidad o tendencias claras en los datos (Kulkarni, Jadha, y Dhingra, 2020). Además, según (Petrică, Stancu, y Tindeche (2016) estas series suelen mostrar asimetrías, cambios abruptos en intervalos no regulares y fluctuaciones en la volatilidad, aspectos que los modelos ARIMA no pueden capturar de manera adecuada.

Por ello, surge la propuesta de considerar modelos más complejos como los algoritmos de aprendizaje automático, entre ellos el *Gradient Boosting* y el SVR. De acuerdo con (Nabipour y cols. (2020), los algoritmos basados en árboles tienen un notable potencial en problemas de regresión para prever los valores futuros de los índices bursátiles. Además, (Ferrouhi y Bouabdallaoui (2024) resaltan la precisión y solidez de los algoritmos *boosting* en la predicción de los precios de las acciones. Sin embargo, a pesar de ser un algoritmo atractivo para proyectos de modelización predictiva en clasificación y regresión debido a su alto rendimiento, *Gradient Boosting* suele requerir mucho tiempo para converger a la solución (Sonkavde y cols., 2023).

Por su parte, aunque no son tan frecuentemente utilizados en la literatura, los algoritmos SVM son considerados entre los más adecuados para predecir el precio de mercado de una acción basándose en sus datos históricos, como han indicado

autores como [Jyothirmayee, Kumar, Rao, y Reddy \(2019\)](#) y [Kurani, Doshi, Vakharia, y Shah \(2023\)](#). Estos algoritmos son reconocidos por su capacidad para manejar relaciones complejas entre variables y capturar patrones no lineales en los datos ([Kapinus, Liashenko, Ljepava, Liashenko, y Danylov, 2024](#)). Según [Yang \(2023\)](#), la rápida convergencia y alta precisión de los modelos SVR les permite predecir eficazmente los datos bursátiles, generando resultados de predicción que se aproximen considerablemente a los valores reales.

7. Conclusiones

Los hallazgos derivados del estudio permiten inferir las siguientes conclusiones:

- Aunque funciones automáticas como `auto.arima` y `GridSearchCV` facilitan el proceso de optimización de hiperparámetros, no necesariamente garantizan una mayor precisión en las predicciones. Esto enfatiza la importancia de un ajuste cuidadoso y personalizado de los modelos de predicción en finanzas. La variabilidad y la naturaleza dinámica de los datos pueden hacer que los enfoques más generalizados sean menos efectivos.
- Mientras que los modelos ARIMA pueden ser útiles para series temporales estables y con relaciones lineales, no son necesariamente la herramienta más adecuada para prever mercados financieros. Las características inherentes a los datos financieros, tales como la no linealidad, la volatilidad y el ruido, plantean desafíos significativos para la modelización y predicciones precisas de los precios de las acciones mediante modelos ARIMA.
- La naturaleza de estos datos implica la necesidad de aplicar técnicas avanzadas de *machine learning*. Esto subraya la importancia de emplear métodos que no solo se adapten a la estructura de los datos financieros, sino que también sean capaces de evolucionar y responder a las dinámicas cambiantes de estos mercados.
- El hecho de que el SVR supere al *Gradient Boosting* en la precisión de las predicciones resalta la importancia de seleccionar un modelo que no solo maneje efectivamente la complejidad y la dimensionalidad de los datos, sino que también ofrezca un margen de decisión robusto y una buena capacidad de generalización. En los mercados financieros, donde la exactitud en las predicciones es crítica y los costos de predicciones erróneas pueden ser elevados, modelos como el SVR, que optimizan estos aspectos, resultan particularmente valiosos.
- En general, el IBEX 35 parece ser el índice mejor predicho, lo que podría deberse a características inherentes al mercado español como mayor estabilidad o menor volatilidad en comparación con otros índices. Por el contrario, el IBOVESPA representa el mayor desafío para los modelos predictivos, probablemente debido a su alta volatilidad y a los impactos de factores económicos y políticos menos predecibles.

Referencias

- Acevedo Arango, H. D. (2012). *Revisión sistemática de literatura: Modelos de pronóstico índice Standard and Poor's 500 (S&P500)* (Trabajo Fin de Grado, Universidad Nacional de Colombia). Descargado de <https://repositorio.unal.edu.co/bitstream/handle/unal/11691/8359354.2012.pdf.pdf?sequence=1&isAllowed=y>
- Aguirre Ariza, L. E., y Castañeda Rueda, D. (2010). *Análisis de pronóstico de precios mediante series temporales aplicado al caso de la bolsa de valores colombiana* (Trabajo Fin de Grado, Universidad Tecnológica de Pereira). Descargado de <https://hdl.handle.net/11059/2195>
- Akaike, H. (1974). A new look at the statistical model identification. *IEEE Transactions on Automatic Control*, 19(6), 716-723. doi: 10.1109/TAC.1974.1100705
- Amat Rodrigo, J. (2023). *Gradient Boosting con Python*. Descargado de https://cienciadedatos.net/documentos/py09_gradient_boosting_python
- Andrade Flores, K. M., y Flores Carvajal, L. V. (2023). *Evaluación de producción de energía renovable utilizando tres modelos autorregresivos* (Trabajo Fin de Grado, Universidad Politécnica Salesiana). Descargado de <http://dspace.ups.edu.ec/handle/123456789/26242>
- Ariyo, A. A., Adewumi, A. O., y Ayo, C. K. (2014). Stock price prediction using the ARIMA model. En *Proceedings - UKSim-AMSS 16th International Conference on Computer Modelling and Simulation, UKSim 2014* (p. 106-112). Institute of Electrical and Electronics Engineers Inc. doi: 10.1109/UKSim.2014.67
- Atsalakis, G. S., y Valavanis, K. P. (2009). Surveying stock market forecasting techniques - Part II: Soft computing methods. *Expert Systems with Applications*, 36(3 PART 2), 5932-5941. doi: 10.1016/j.eswa.2008.07.006
- Ayala Castrejón, R. F., y Bucio Pacheco, C. (2020). Modelo ARIMA aplicado al tipo de cambio peso-dólar en el periodo 2016-2017 mediante ventanas temporales deslizantes. *Revista Mexicana de Economía y Finanzas Nueva Época*, 15(3), 331-354. doi: 10.21919/remef.v15i3.466
- Ben Taieb, S., Bontempi, G., Atiya, A. F., y Sorjamaa, A. (2012). A review and comparison of strategies for multi-step ahead time series forecasting based on the nn5 forecasting competition. *Expert Systems with Applications*, 39(8), 7067-7083. doi: 10.1016/j.eswa.2012.01.039
- Bentéjac, C., Csörgő, A., y Martínez-Muñoz, G. (2021). A comparative analysis of

- Gradient Boosting algorithms. *Artificial Intelligence Review*, 54(3), 1937-1967. doi: 10.1007/s10462-020-09896-5
- Bishop, C. M. (2006). *Pattern recognition and machine learning*. Springer New York.
- Bollerslev, T. (1986). Generalized autoregressive conditional heteroskedasticity. *Journal of Econometrics*, 31(3), 307-327. doi: 10.1016/0304-4076(86)90063-1
- Bonilla Ruiz, A. (2019). *Estudio estadístico de un caso real. Series temporales* (Trabajo Fin de Grado, Universidad de Almería). Descargado de <http://hdl.handle.net/10835/7878>
- Box, G. E. P., y Cox, D. R. (1964). An analysis of transformations. *Journal of the Royal Statistical Society. Series B (Methodological)*, 26(2), 211-252. Descargado de <https://www.jstor.org/stable/2984418>
- Box, G. E. P., y Jenkins, G. M. (1976). *Time series analysis: Forecasting and control*. Holden Day.
- Brugger, S., y Ortiz, E. (2012). Mercados accionarios y su relación con la economía real en américa latina. *Problemas del Desarrollo*, 43(168), 63-93. doi: 10.22201/iiec.20078951e.2012.168.28638
- Cáceres Hernández, J. J., Martín Rodríguez, G., y Martín Álvarez, F. J. (2008). *Introducción al análisis univariante de series temporales económicas*. Delta Publicaciones.
- Camacho, M., Ramallo, S., y Ruiz Marín, M. (2021). Árboles de decisión en economía: Una aplicación a la determinación del precio de la vivienda. En D. Peña, P. Poncela, y E. Ruiz (Eds.), *Nuevos métodos de predicción con datos masivos* (p. 61-92). Funcas. Descargado de <https://www.funcas.es/libro/nuevos-metodos-de-prediccion-economica-con-datos-masivos/>
- Canales, R. J. (2017). Estado actual de los Índices bursátil en el mundo. *Revista Electrónica de Investigación en Ciencias Económicas*, 5(9), 65-84. doi: 10.5377/reice.v5i9.4363
- Cao, L., y Tay, F. (2001). Financial forecasting using Support Vector Machines. *Neural Computing & Applications*, 10(2), 184-192. doi: 10.1007/s005210170010
- Cardona Ochoa, S. (2023). *Análisis de factores externos que influyen en el desempeño bursátil de Ecopetrol* (Trabajo Final de la asignatura Macroeconomía avanzada, Universidad Pontificia Bolivariana). Descargado de <https://repository.upb.edu.co/handle/20.500.11912/10908>
- Catani, P. S., y Ahlgren, N. J. (2017). Combined Lagrange multiplier test for ARCH in vector autoregressive models. *Econometrics and Statistics*, 1, 62-84. doi: 10.1016/j.ecosta.2016.10.006
- Comisión Nacional del Mercado de Valores. (2020). El mercado de valores y los productos de inversión: Manual para universitarios [Manual de

- software informático]. Descargado de <https://www.cnmv.es/DocPortal/Publicaciones/Guias/ManualUniversitarios.pdf>
- Dash, R., y Dash, P. K. (2016). A hybrid stock trading framework integrating technical analysis with machine learning techniques. *Journal of Finance and Data Science*, 2(1), 42-57. doi: 10.1016/j.jfds.2016.03.002
- de la Fuente Fernández, S. (2013). *Series temporales: modelo arima*. Portal Estadística Aplicada. Descargado de <https://www.estadistica.net/ECONOMETRIA/SERIES-TEMPORALES/modelo-arima.pdf>
- Díez García, J. (2020). *Machine learning aplicado al trading: Marco teórico y análisis de casos prácticos* (Trabajo Fin de Grado, Universidad Pontificia de Comillas). Descargado de <http://hdl.handle.net/11531/49353>
- Drucker, H., Burges, C. J. C., Kaufman, L., Smola, A. J., y Vapnik, V. (1996). Support Vector Regression Machines. *Advances in Neural Information Processing Systems*, 9, 155-161.
- Egea Rosas, C. (2020). *¿Cómo invertir en el mercado de acciones?* (Trabajo Fin de Grado, Universidad Politécnica de Cartagena). Descargado de <http://hdl.handle.net/10317/9165>
- El Amrani Joutey, E. H. (2020). *Análisis de series temporales: Uso del transporte público en Barcelona* (Trabajo Fin de Grado, Universitat Politècnica de Catalunya). Descargado de <http://hdl.handle.net/2117/340448>
- Engle, R. F. (1982). Autoregressive conditional heteroscedasticity with estimates of the variance of United Kingdom inflation. *Econometrica*, 50(4), 987-1007. doi: 10.2307/1912773
- Ferrari, F. (2023). *Análisis del comportamiento del precio de una acción en el mercado de valores argentino utilizando un modelo ARIMA* (Trabajo Fin de Grado, Universidad Nacional de Río Negro). Descargado de <http://rid.unrn.edu.ar/handle/20.500.12049/10116>
- Ferrouhi, E. M., y Bouabdallaoui, I. (2024). A comparative study of ensemble learning algorithms for high-frequency trading. *Scientific African*, 24(e02161), e02161. doi: 10.1016/j.sciaf.2024.e02161
- Fierro, A. A. (2020). *Predicción de series temporales con redes neuronales* (Trabajo Fin de Grado). Descargado de <http://sedici.unlp.edu.ar/handle/10915/114857>
- Fisher, I. (1922). *The making of index numbers: A study of their varieties, tests, and reliability*. Houghton Mifflin Company.
- Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *The Annals of Statistics*, 29(5), 1189-1232. doi: 10.1214/aos/1013203451
- Fu, Z., y Wang, Z. (2021). Prediction of financial economic time series based on

- group intelligence algorithm based on machine learning. *Revista Argentina de Clínica Psicológica*(2), 938-947. doi: 10.24205/03276716.2020.4101
- García Soler, G. (2018). *Modelos econométricos para la predicción a corto plazo de cotizaciones del IBEX 35* (Trabajo Fin de Grado, Universidad Miguel Hernández). Descargado de <https://hdl.handle.net/11000/7432>
- Gitman, L. J., y Joehnk, M. D. (2009). *Fundamentos de inversiones* (10.^a ed.). Pearson Educación.
- Glosten, L. R., Jagannathan, R., y Runkle, D. E. (1993). On the relation between the expected value and the volatility of the nominal excess return on stocks. *The Journal of Finance*, 48(5), 1779-1801. doi: 10.1111/j.1540-6261.1993.tb05128.x
- Goldfeld, S., y Quandt, R. (1965). Some tests for heteroscedasticity. *Journal of the American Statistical Association*, 60(310), 539-547. doi: 10.1080/01621459.1965.10480811
- Gómez Rueda, D. M. (2016). *Modelo predictivo del índice bursátil colombiano de deuda pública tasa fija COLTES LP entre enero 2016 y mayo del 2021* (Trabajo Fin de Grado, Fundación Universitaria Los Libertadores). Descargado de <http://hdl.handle.net/11371/4148>
- González Casimiro, M. P. (2009). *Técnicas de predicción económica* (Trabajo Fin de Grado, Universidad del País Vasco). Descargado de <http://hdl.handle.net/10810/12493>
- Gujarati, D. N., y Porter, D. C. (2009). *Econometría* (5.^a ed.). McGraw-Hill.
- Gutiérrez, E. D., Puche, R., y Hernández, F. (2020). Estimación de casos de COVID-19 en países de Suramérica empleando modelos ARIMA (Autorregresivo Integrado de Promedio Móvil). *Observador del Conocimiento*, 5(3), 11-25. Descargado de <https://revistaoc.oncti.gob.ve/index.php/odc/article/view/155>
- Gutiérrez Ramos, D. (2022). *Análisis de series temporales: Estudio de un caso práctico* (Trabajo Fin de Grado, Universidad Politécnica de Madrid). Descargado de <https://oa.upm.es/70926/>
- Hernández, B. (2000). *Bolsa y estadística bursátil*. Ediciones Díaz de Santos.
- Hernández Rodríguez, C. (2020). *Predicción y clasificación de series temporales bursátiles mediante redes neuronales recurrentes* (Trabajo Fin de Grado, Universidad de Valladolid). Descargado de <http://uvadoc.uva.es/handle/10324/44399>
- Hom Serra, N. (2021). *Análisis bursátil de 6 empresas del IBEX 35 durante el período 2019-20* (Trabajo Fin de Grado, Universitat de Barcelona). Descargado de <http://hdl.handle.net/2445/180953>
- Hyndman, R., y Athanasopoulos, G. (2021). *Forecasting: Principles and practice* (3.^a ed.). OTexts. Descargado de <https://otexts.com/fpp3/>

- Jain, A. (2022). *Complete machine learning guide to parameter tuning in Gradient Boosting (GBM) in Python*. Descargado de <https://www.analyticsvidhya.com/blog/2016/02/complete-guide-parameter-tuning-gradient-boosting-gbm-python/>
- James, G., Witten, D., Hastie, T., y Tibshirani, R. (2013). *An introduction to statistical learning with applications in R*. Springer Texts in Statistics. doi: 10.1007/978-1-4614-7138-7
- Janiesch, C., Zschech, P., y Heinrich, K. (2021). Machine learning and deep learning. *Electronic Markets*, 31, 685-695. doi: 10.1007/s12525-021-00475-2
- Jaramillo, J. A., Velásquez, J. D., y Franco, C. J. (2017). Research in financial time series forecasting with SVM: Contributions from literature. *IEEE Latin America Transactions*, 15(1), 145-153. doi: 10.1109/TLA.2017.7827918
- Jarque, C. M., y Bera, A. K. (1980). Efficient tests for normality, homoscedasticity and serial independence of regression residuals. *Economics Letters*, 6(3), 255-259. doi: 10.1016/0165-1765(80)90024-5
- Jiménez Almaraz, L. (2010). *LATIBEX: El mercado latinoamericano de valores* (Trabajo Fin de Máster, Universidad Internacional de Andalucía). Descargado de <http://hdl.handle.net/10334/396>
- Jyothirmayee, S., Kumar, V., Rao, C., y Reddy, S. (2019). Predicting stock exchange using supervised learning algorithms. *International Journal of Innovative Technology and Exploring Engineering*, 9(1), 4081-4090. doi: 10.35940/ijitee.A4144.119119
- Kapinus, M., Liashenko, K., Ljepava, N., Liashenko, L., y Danylov, V. (2024). Predicting stock market trends with Python. *Grail of Science*(40), 109-116. doi: 10.36074/grail-of-science.07.06.2024.012
- Kendory, E., Kareem, N. A., y Humadi, S. M. (2020). Using the nonlinear regression to predict the prices of stocks and its impact on the investor decisions in the financial market. *International Journal of Psychosocial Rehabilitation*, 24(8), 8825-8844. doi: 10.37200/IJPR/V2418/PR280880
- Kinjang Li Ye, K., Aguilar Game, E. M., y Muñoz Franco, M. G. (2023). Pronóstico de la producción de petróleo de las empresas públicas del Ecuador mediante la aplicación del modelo estadístico ARIMA. *Visión Empresarial*, 2(3). Descargado de <https://revistasdigitales.uniboyaca.edu.co/index.php/viem/article/view/1226>
- Kulkarni, M., Jadha, A., y Dhingra, D. (2020). Time series data analysis for stock market prediction. En *Proceedings of the International Conference on Innovative Computing & Communications (ICICC)*. Elsevier BV. doi: 10.2139/SSRN.3563111

- Kurani, A., Doshi, P., Vakharia, A., y Shah, M. (2023). A comprehensive comparative study of Artificial Neural Network (ANN) and Support Vector Machines (SVM) on stock forecasting. *Annals of Data Science*, 10, 183-208. doi: 10.1007/s40745-021-00344-x
- Lanteri, L. (2011). Desarrollo del mercado accionario y crecimiento económico. Alguna evidencia para la Argentina. *Ensayos de Economía*(38), 117-145. Descargado de <https://repositorio.unal.edu.co/handle/unal/39384>
- Laspeyres, (1871). Die berechnung einer mittleren waarenpreissteigerung. *Jahrbücher für Nationalökonomie und Statistik*, 16, 296-314.
- Lillo León, M. (2016). *Índices bursátiles: el IBEX 35* (Trabajo Fin de Grado, Universidad de Jaén). Descargado de <https://hdl.handle.net/10953.1/7081>
- Ljung, G. M., y Box, G. E. P. (1978). On a measure of lack of fit in time series models. *Biometrika*, 65(2), 297-303. doi: 10.2307/2335207
- López Rial, R. (2016). *Análisis histórico de la performance de la Bolsa de Madrid (1990-2015)* (Trabajo Fin de Grado, Universidad de Cádiz). Descargado de <http://hdl.handle.net/10498/18766>
- Marcos Rubio, M. A. (2023). *Estimación de modelos ARMA usando el modelo espacio de los estados* (Trabajo Fin de Grado, Universidad Politécnica de Madrid). Descargado de <https://oa.upm.es/74963/>
- Martín, J. L. (1994). *Mercados derivados sobre índices bursátiles*. Gesmovasa.
- Martín González, D. (2023). *Clasificación estrella-galaxia con métodos de aprendizaje automático* (Trabajo Fin de Máster, Universidad de Salamanca). Descargado de <http://hdl.handle.net/10366/157836>
- Martín Guareño, J. J. (2016). *Support Vector Regression: Propiedades y aplicaciones* (Trabajo Fin de Grado, Universidad de Sevilla). Descargado de <http://hdl.handle.net/11441/43808>
- Martí Sanahuja, P. (2021). *Métricas de evaluación de rendimiento para predicciones de series temporales*. Descargado de <https://polmartisanahuja.com/metricas-de-evaluacion-de-rendimiento-para-predicciones-de-series-temporales/>
- Mauricio, J. A. (2017). *Introducción al análisis de series temporales*. Descargado de <https://www.ucm.es/data/cont/docs/518-2013-11-11-JAM-IASST-Libro.pdf>
- Mendoza-Rivera, R. J., Lozano-Díez, J. A., y Venegas-Martínez, F. (2020). Impacto de la pandemia COVID-19 en variables financieras relevantes en las principales economías de Latinoamérica. *Economía: teoría y práctica*(spe5), 125-144. doi: 10.24275/etypuam/ne/e052020/mendoza
- Mesafint Belete, D., y Huchaiah, M. D. (2021). Grid search in hyperparameter optimization of machine learning models for prediction of HIV/AIDS test results. *International Journal of Computers and Applications*, 44(9), 875-886. doi:

10.1080/1206212X.2021.1974663

- Miranda Chinlli, C. (2021). *Modelización de series temporales. Modelos clásicos y SARIMA* (Trabajo Fin de Máster, Universidad de Granada). Descargado de https://masteres.ugr.es/estadistica-aplicada/sites/master/moea/public/inline-files/TFM_MIRANDA_CHINLLI_CARLOS.pdf
- Moreno Gómez, L. (2020). *La dimensión fractal de la Bolsa* (Trabajo Fin de Grado, Universidad Rey Juan Carlos). Descargado de <http://hdl.handle.net/10115/19650>
- Nabipour, M., Nayyeri, P., Jabani, H., Mosavi, A., Salwana, E., y Shahab, S. (2020). Deep learning for stock market prediction. *Entropy*, 22(8). doi: 10.3390/e22080840
- Navales Cardona, E. S. (2013). *Una revisión sistemática acerca de las metodologías para el pronóstico de índices de mercado: Su estado actual y tendencias futuras* (Trabajo Fin de Grado, Universidad Nacional de Colombia). Descargado de <https://repositorio.unal.edu.co/handle/unal/12079>
- Nicolás Martínez, M., y Devesa Santamaria, V. (2021). *Análisis financiero de un índice bursátil. NASDAQ-100* (Trabajo Fin de Grado, Universidad Católica de Murcia). Descargado de <http://hdl.handle.net/10952/5199>
- Nieto Benito, G. J. (2023). *Predicción de series temporales mediante técnicas de aprendizaje automático. Automatización del proceso* (Trabajo Fin de Grado, Universidad de Sevilla). Descargado de <https://biblus.us.es/bibing/proyectos/abreproy/94769/fichero/TFG-4769+Nieto+Benito.pdf>
- Nieves Moina, P. (2019). *Análisis del comportamiento del mercado bursátil mediante modelos ARIMA* (Trabajo Fin de Grado, Universidad del País Vasco). Descargado de <http://hdl.handle.net/10810/36817>
- Ortiz Martín, R. (2022). *Metodología ARIMAX aplicada al análisis de la evolución de datos de suicidios en España* (Trabajo Fin de Grado, Universitat de Barcelona). Descargado de <http://hdl.handle.net/2117/385123>
- Osorio, J. E., y Ángel, V. H. (2019). *Aplicación del modelo Box-Jenkins en la estimación del flujo de llamadas en servicios de atención al cliente en el sector contact center* (Trabajo Fin de Grado, Fundación Universitaria Los Libertadores). Descargado de <http://hdl.handle.net/11371/2798>
- Osorio-Barreto, D., Mejía-Rubio, P. P., y Mora-Mora, J. U. (2022). Inflation expectations: A systematic literature review and bibliometric analysis. *Revista de Economía Contemporánea*, 26, 1-24. doi: 10.1590/198055272609
- Paasche, H. (1874). Über die preise. *Journal des Economistes*, 20, 5-42.
- Pérez García, M. (2018). *Técnicas boosting* (Trabajo Fin de Grado, Universidad de Sevilla). Descargado de <https://hdl.handle.net/11441/81665>

- Petrică, A., Stancu, S., y Tindeche, A. (2016). Limitation of ARIMA models in financial and monetary economics. En *Theoretical and Applied Economics* (p. 19-42). Descargado de <https://api.semanticscholar.org/CorpusID:157904704>
- Pineda Pertuz, C. M. (2022). *Aprendizaje automatico y profundo en Python: Una mirada hacia la inteligencia artificial*. Ra-Ma.
- Pisner, D. A., y Schnyer, D. M. (2019). Support Vector Machine. En *Machine learning: Methods and applications to brain disorders* (p. 101-121). Elsevier. doi: 10.1016/B978-0-12-815739-8.00006-7
- Podadera Molero, A. A. (2022). *Series temporales aplicadas* (Trabajo Fin de Máster, Universidad de Granada). Descargado de <https://masteres.ugr.es/estadistica-aplicada/docencia/trabajo-fin-master/antiguo>
- Pring, M. J. (1991). *Technical analysis explained: The successful investor's guide to spotting investment trends and turning points* (5.ª ed.). McGraw-Hill.
- Ramón Escolano, N., y López Espín, J. J. (2016). *Econometría: Series temporales y modelos de ecuaciones simultáneas*. Universidad Miguel Hernández de Elche.
- Rey, C., y Ramil, M. (2007). *Introducción a la estadística descriptiva* (2.ª ed.). Netbiblio. Descargado de <http://hdl.handle.net/2183/11897>
- Rodríguez Aranday, F. (2020). *Análisis bursátil: Metodología y práctica* (1.ª ed.). Instituto Mexicano de Contadores Públicos.
- Rodríguez-Marín, M. (2024). Demanda de turistas internacionales hacia México: Construcción de un modelo predictivo. *Contaduría y Administración*, 69(4), 327-344. doi: 10.22201/fca.24488410e.2024.5092
- Romero-Meza, R., Coronado, S., y Ibañez-Veizaga, F. (2021). COVID-19 y causalidad en la volatilidad del mercado accionario chileno. *Estudios Gerenciales*, 37(159), 242-250. doi: 10.18046/j.estger.2021.159.4412
- Roque García, A. D. (2021). *Modelo de Box-Jenkins para el pronóstico de la recaudación de ingreso de la genérica 1.3 de una entidad pública, 2017 – 2021* (Trabajo Fin de Grado, Universidad Nacional Mayor de San Marcos). Descargado de <https://hdl.handle.net/20.500.12672/17379>
- Rosero, L. (2010). El desarrollo del mercado de valores en el Ecuador: Una aproximación. *Ecuador debate*, 80, 23-34. Descargado de <http://hdl.handle.net/10469/3490>
- Rossi, P., Arbocco, G., Vacas, Y., Maco, C., y Posadas, L. H. (2012). *Índices bursátiles* (Trabajo Fin de Grado, Universidad ESAN). doi: 10.2139/ssrn.2129499
- Salvador Maceira, M. (2019). *Machine learning aplicado al trading* (Trabajo Fin de Grado, Universidad Pontificia de Comillas). Descargado de <https://repositorio.comillas.edu/rest/bitstreams/295750/retrieve>
- Sánchez García, J. (2015). *Análisis bursátil. Análisis técnico y análisis fundamental*

- (Trabajo Fin de Máster, ICADE Business School). Descargado de <http://hdl.handle.net/11531/6433>
- Scherk, A. (2011). *Manual de análisis fundamental* (5.^a ed.). Inversor Ediciones. Descargado de [https://www.academia.edu/26801419/Manual Analisis Fundamental](https://www.academia.edu/26801419/Manual_Analisis_Fundamental)
- Shen, Z., Wan, Q., y Leatham, D. J. (2021). Bitcoin return volatility forecasting: A comparative study between GARCH and RNN. *Journal of Risk and Financial Management*, 14(7), 337-355. doi: 10.3390/jrfm14070337
- Sonkavde, G., Dharrao, D. S., Bongale, A. M., Deokate, S. T., Doreswamy, D., y Bhat, S. K. (2023). Forecasting stock market prices using machine learning and deep learning models: A systematic review, performance analysis and discussion of implications. *International Journal of Financial Studies*, 11(3), 94. doi: 10.3390/ijfs11030094
- Sotaquirá, M. (2023). *Parámetros e hiperparámetros en el machine learning*. Descargado de <https://www.codificandobits.com/blog/parametros-hiperparametros-machine-learning/>
- Tsay, R. S. (2005). *Analysis of financial time series*. John Wiley & Sons. doi: 10.1002/0471746193
- Valova, I., Gueorguieva, N., Aayushi, T., Nikitha, P., y Mohamed, H. (2023). Hybrid models for predicting stock market performance. *Journal of Machine Intelligence and Data Science (JMIDS)*, 4, 6-14. doi: 10.11159/jmids.2023.002
- Vapnik, V. N. (1998). *Statistical learning theory*. Wiley.
- Villanueva Gonzales, A. (2007). Mercados financieros: Una aproximación a la Bolsa de Valores de Lima. *Contabilidad y Negocios*, 2(3), 23-33. doi: 10.18800/contabilidad.200701.003
- Walker, E. (1998). Mercado accionario y crecimiento económico en Chile. *Cuadernos de Economía*, 35(104), 49-72.
- Yang, J. (2023). Support vector machine-based stock prediction analysis. *Business, Economics and Management (MSIE)*, 3, 12-18. doi: 10.54097/hbem.v3i.4625
- Yoon, J. (2021). Forecasting of real GDP growth using machine learning models: Gradient Boosting and Random Forest approach. *Computational Economics*, 57(1), 247-265. doi: 10.1007/s10614-020-10054-w
- Zhang, F., y O'Donnell, L. J. (2019). Support Vector Regression. En *Machine learning: Methods and applications to brain disorders* (p. 123-140). Elsevier. doi: 10.1016/B978-0-12-815739-8.00007-9