

ORIGINAL ARTICLE

Police Detection of Deception: Beliefs About Behavioral Cues to Deception Are Strong Even Though Contextual Evidence Is More Useful

Jaume Masip & Carmen Herrero

School of Psychology, University of Salamanca, 37005 Salamanca, Spain

Research questions the validity of behavioral deception cues; however, people believe behavioral cues are reliable deception indicators. Police officers and community members indicated both how lies can be detected (beliefs), and how they discovered a lie in the past (revealing information). Officers did the latter twice, prompted with a professional versus a personal context. For both groups, beliefs were primarily behavioral (e.g., demeanor) and revealing information contextual (evidence, third-party information, etc.). Officers responded similarly regardless of context. Relative to community members, officers provided more cues and referred more often to verbal contradictions and active detection strategies when asked about beliefs. Practitioners should be made aware of the discrepancies between their beliefs about deception cues and useful information to detect lies.

Keywords: Deception, Deception Cues, Beliefs, Police, Context, TDT.

doi:10.1111/jcom.12135

Deception is a frequent phenomenon in many areas of life, and forensic contexts are a kind of situation where it may be most relevant (see Granhag & Strömwall, 2004). As the consequences of being caught after having broken the law are often negative, both guilty and innocent suspects might deny their involvement in a crime. In order to be able to perform their role, the police should be able to discriminate between truths and lies correctly. In an attempt to do so, officers (and nonofficers too) may pay attention to the senders' behavior.

Why should people focus on behavior when assessing veracity? Having accurate depictions of the world is adaptive. Evolution has provided species with the senses to gather information from the world. This makes it possible for organisms to safely interact with their environments, avoiding potential dangers while fulfilling their

Corresponding author: Jaume Masip; e-mail: jmasip@usal.es

basic needs. Because of this, organisms do not ordinarily question the reality of the sensations (information about the world) they get through their senses. As an extension of this, ordinarily human beings do not question the reality of the messages they get from others (Gilbert, 1991; Levine, 2014; Levine, Park, & McCornack, 1999). This may be adaptive, as most people tell the truth most of the time, but it also makes individuals vulnerable to occasional deceit (Levine, 2014). This poses a risk.

As a result, most cultures condemn deceit (Inglehart, Basáñez, Díez-Medrano, Halman, & Luijckx, 2004) and their members are socialized into truthfulness. In writing the report of the Global Deception Research Team's (2006) study, Charlie Bond reasoned that the worldwide-shared beliefs about indicators of deceit identified by the Team may not be descriptive but prescriptive. That is, rather than reflecting *actual* differences between truth tellers and liars, they indicate how liars *should* behave. The ultimate goal is to discourage lies: "Children should be ashamed when they lie to their parents, and liars should feel bad. Lying should not pay, and liars should be caught. ... The hope is that lying will be deterred or (at least) that the caregiver's prophesies of shame will be self-fulfilling" (Global Deception Research Team, 2006, pp. 69–70).

As a likely outcome of these socialization practices, people strongly believe that lies can be detected from nonverbal (mainly) and verbal (secondarily) behavior. In the Global Deception Research Team's (2006) study, when asked "How can you tell when people are lying?" citizens from 58 different countries listed a number of nonverbal and verbal behaviors (that, in reality, are poorly related to actual veracity). Professional detectors and lay people have similar beliefs about deception cues (Strömwall, Granhag, & Hartwig, 2004).

Not only do people mention behavioral cues when asked how lies can be detected, they also focus on behavioral information when it is available, overlooking contextual revealing information that would allow perfect accuracy (Bond, Howard, Hutchison, & Masip, 2013).

Just like lay people, deception researchers have embraced the premise that deception can be revealed through behavior. All major theoretical perspectives on deception, such as Ekman and Friesen's (1969) leakage hypothesis, Zuckerman, DePaulo, and Rosenthal's (1981) four-factor theory, Buller and Burgoon's (1996) Interpersonal Deception Theory, and DePaulo et al.'s (2003) self-presentational perspective, are predicated on the premise that observable behavior can provide cues to deceit (or may leak the concealed affect or information).

However, lay people's and scientists' focus on behavior as indicative of truth or deception may be misguided. If behavioral cues were clean pathways to veracity, then message recipients (at least those more experienced or properly trained) would be able to accurately assess honesty by observing the senders' behavior. This prediction is at odds with research findings accumulated over several decades. Meta-analyses show that humans' ability to separate truths from lies on the basis of the sender's behavior hardly exceeds chance probability (Aamodt & Custer, 2006; Bond & DePaulo, 2006), that practitioners whose jobs require skill at detecting deception are not any more accurate than lay persons (e.g., Aamodt & Custer, 2006; Bond & DePaulo, 2006; Vrij,

2008), and that cue training to detect deception has only limited effects on accuracy (Frank & Feeley, 2003), particularly if it focuses on nonverbal cues (Hauch, Sporer, Michael, & Meissner, 2014). Further, training effects can be explained at least in part by an increase in the trainees' motivation rather than by cue validity (Levine, Feeley, McCornack, Hughes, & Harms, 2005), and cue training may bias judgments toward deception rather than increasing accuracy (see Masip, Alonso, Garrido, & Herrero, 2009). Levine et al. (2011) identified a set of behaviors that prompted either truth or deception judgments but were poorly correlated with actual veracity. All of these findings question the notion that deceit is revealed through nonverbal cues (see Levine, 2010, 2014; Levine & Kim, 2010).

Even more conclusive evidence of the poor usefulness of behavioral deception indicators comes from meta-analyses focusing precisely on the behavioral differences between liars and truth tellers. These meta-analyses reveal only few and small differences, and the utility of the scant diagnostic cues is under the influence of a host of moderator variables (DePaulo et al., 2003; Sporer & Schwandt, 2006, 2007). Even the most reliable cues are poorly related with deception (Hartwig & Bond, 2011). As an aside, it is fair to say here that verbal cues appear to be more revealing than nonverbal cues (DePaulo et al., 2003), that training focusing on verbal cues is more effective than that focused on nonverbal cues (Hauch et al., 2014), and that receivers attain greater accuracy rates when they focus on verbal (rather than nonverbal) information (Bond & DePaulo, 2006). Still, the overall picture of the usefulness of behavioral deception cues as a whole is one of general skepticism.

In summary, laboratory studies show that human observers are poor truth and lie detectors. The reason appears to be that behavioral cues do not reveal deception. Yet, people believe that behavioral cues are powerful deception markers. They even take fallible nonverbal information into account even when almost perfectly diagnostic contextual information is available (Bond et al., 2013).

Contextual deception indicators

How could observers' detection accuracy be improved? Two promising approaches have been proposed recently. First, as behavioral differences between liars and truth tellers are minute, interviewing approaches to increase these differences (i.e., to increase the *transparency* of senders; see Levine, 2010) ought to be designed (Hartwig & Bond, 2011). Indeed, a wealth of exciting research has been conducted in recent times along this line (see, e.g., Hartwig, Granhag, & Luke, 2014; Levine, Blair, & Clare, 2014; Levine, Shaw, & Shulman, 2010; Vrij & Granhag, 2012). Second, as behavioral differences between liars and truth tellers are minute, nonbehavioral situational or contextual indicators have been suggested as a promising alternative (Blair, Levine, Reimer, & McCluskey, 2012). Contextual indicators involve situational information that either reveals deception or can be used to compare the sender's communication to assess its veracity. Examples include tangible evidence, third-party information, or confessions (Park, Levine, McCornack, Morrison, & Ferrara, 2002), as well as knowledge about the sender's normal behaviors, his or her incentives, or what is

normal or possible in a situation (see Blair, Levine, & Shaw's, 2010, notion of "content in context"). This differs from decontextualized verbal, nonverbal, or physiological leakage, or deception cues involuntarily displayed by the sender at the time of lying.

According to Blair et al. (2010), potentially revealing contextual information may be of three general kinds: contradictions between the message delivered and what the detector already knew, normative information (about the liars' routine activities, the laws of physics, or the normal ways of reacting to certain situations), and idiosyncratic information (e.g., shortages in a bank always stop when one of the employees goes on vacation and begin again when he/she returns; Blair et al., 2010). Blair et al. (2010) showed the potential of this approach. Across eight separate experiments, overall accuracy was 57% (63% for truths and 52% for lies) when only behavioral information was provided; when contextual information was given too, accuracy increased to 75% (74% for truths and 78% for lies).

The usefulness of situational information (relative to behavioral cues) to detect deception was first stressed in a provocative paper by Park et al. (2002). The authors contended that the low accuracy rates in judging veracity found in many studies might be an artifact. More specifically, in laboratory experiments potentially revealing contextual deception cues are absent by design; instead, only behavioral cues are available. Park et al. argued that in real life accuracy might be greater than in laboratory experiments because outside the laboratory people ordinarily do not have to make *immediate* veracity judgments of *strangers* based solely on *behavioral information*. The authors asked 202 respondents to recall a lie they had detected in the past and to indicate how they had detected it. In support for their contentions, they found that in real life lies are only rarely detected from behavioral cues; instead, third-party information, confessions, and physical evidence were the most useful sources of information. Further, they showed that normally lies are not detected immediately after being told, and that in virtually all cases the detector personally knew the liar.

Reinhard, Sporer, Scharmach, and Marksteiner's (2011) findings that receivers who are familiar with the situation are more accurate in making veracity judgments than receivers who are unfamiliar with the situation may appear consistent with the notion that contextual cues are useful to detect deception. As anticipated by Stiff et al. (1989), receivers in familiar situations are able to assess the verisimilitude of the message verbal content—for instance, by comparing verbal content with their knowledge about the situation. However, Reinhard, Scharmach, and Sporer (2012) showed that perceived (and not necessarily real) familiarity is enough for this effect to happen. This questions the notion that the situational familiarity effect is caused by accurate contextual information, although it does not question the more general notion that contextual cues may reveal deception.

The present study: Rationale and hypotheses

This study is an extension of Park et al.'s (2002) research with a number of crucial changes. Also, it aligns well with Levine's (2014) *Truth Default Theory* (TDT). TDT is a collection of separate but logically consistent notions, all supported by previous

Table 1 Propositions of Levine's Truth Deception Theory (See Levine, 2014, for More Details)

Proposition 1	Most people tell the truth most of the time.
Proposition 2	Most lies are told by a few prolific liars.
Proposition 3	Most people believe what others say most of the time.
Proposition 4	This is adaptive and enables efficient communication, yet it makes people vulnerable to occasional deception.
Proposition 5	Both truthful and deceptive messages are means to attain certain goals; most people do not resort to lies if these goals can be attained with the truth.
Proposition 6	When the truth is inconsistent with a desired goal of the sender, people may doubt veracity.
Proposition 7	Other factors raising suspicion ("triggers") are a lack of coherence (internal logical consistency) in message verbal content, a lack of correspondence between the message and the known reality, and third-party information revealing deception.
Proposition 8	If these triggers are strong enough, the person will scrutinize the message to assess its veracity.
Proposition 9	The person may eventually judge the message as deceptive on the basis of communication context and motive, sender demeanor, third-party information, and degree of coherence and correspondence.
Proposition 10	Deception triggers may not occur at the time of the deception.
Proposition 11	Except for a few transparent individuals, deception is not accurately detected by passively observing the senders' behavior at the time the lie is told because honest and deceptive demeanor are poorly related to actual veracity.
Proposition 12	Instead, if deception is detected this happens later in time via a confession by the liar, external evidence, or correspondence.
Proposition 13	Context-sensitive questioning of the potential liar may produce diagnostic information, whereas the wrong questioning may hinder detection accuracy.
Proposition 14	Deception detection expertise does not involve skill in passively detecting and interpreting behaviors but skill in generating diagnostic information from potential liars.

research conducted by Levine and his collaborators. It contains 13 modules and the 14 propositions summarized in Table 1.

In this study, police officers and a control group of lay people of identical gender and age as the officers were first asked *how lies can be detected (beliefs)*, and then had to describe *how they had discovered that someone had lied to them in the past (revealing information)*. Police officers had to answer the second question twice, once prompted with a professional context and once with a personal context. Below we are explaining how the present research differs from Park et al.'s (2002) and how it relates to TDT.

First, *previous research focused on either beliefs about deception cues or on revealing information. These two foci converge here.* Research into beliefs shows that people think behavioral cues are telltale deception indicators, whereas Park et al. (2002; see also Blair et al., 2010) showed that revealing information is contextual rather than behavioral. However, the participants in both lines of research were not the same individuals. *We sought to show that the same individuals would mention nonbehavioral information when asked about a lie successfully detected in the past, yet they would still believe that behavioral information is the most revealing kind of information.* This hypothesis is grounded on the following considerations.

1. In line with Bond (Global Deception Research Team, 2006), we argued above that worldwide socialization practices convey the notion that lies are transparent and can be detected through behavior. The effects of these practices are arguably powerful. Indeed, previous empirical research shows that people (including practitioners) report behavioral cues when asked how deception can be detected (Global Deception Research Team, 2006; Strömwall et al., 2004). Also, veracity judgments are influenced by behavioral information when this information is available (Bond et al., 2013).
2. These notions may be ubiquitous and persistent. People have difficulty in explaining the behavior of others on the basis of situational forces, and search instead for causes within the person (Gilbert & Malone, 1995; Ross, 1977; Ross & Nisbett, 1991). Similarly, indicators of deception may be searched for in the person rather than the situation.
3. However, as shown by Bond et al. (2013), in reality situations may prompt people to lie and hence access to situational information may increase accuracy. For this reason, respondents will mention situational cues when describing how they detected a lie in the past. This is consistent with the *Deception Motives* (Levine, Kim, & Hamel, 2010) and the *Projected Motive Model* (Levine, Kim, & Blair, 2010) modules of TDT, as well as with Propositions 5, 6, and 8 of TDT (see Table 1).
4. We expected the nonbehavioral information mentioned by participants when asked about a successfully detected lie to be of the same kinds reported by Park et al. (2002), including confessions, external evidence (correspondence), and third-party information. Thus, this prediction also aligns well with the *How People Really Detect Lies* module and Propositions 7 through 9, and particularly 12 of TDT.

Second, *participants in the present experiment were not only lay respondents but also police officers.* Exploring Park et al.'s (2002) contentions with law enforcement personnel was the main goal of this study. Much deception research has been conducted to date examining officers' lie detection *outcomes* (response bias, accuracy, confidence, and beliefs about deception cues; see, e.g., Garrido & Masip, 1999; Vrij, 2008). However, almost no study has focused on officers' *processes* in detecting deception, that is, the mechanisms and strategies they use to assess veracity. The findings of

the few extant studies (Masip, Garrido, Herrero, Antón, & Alonso, 2006; Nahari, 2012; Reinhard, Dahm, & Scharmach, 2012) are limited because of being laboratory-based, thus providing only behavioral—and not contextual—information to observers. In this study, officers were questioned about real lies they detected in real life. We explored whether, because of their frequent experience with deception, officers learned to rely more strongly on revealing nonbehavioral cues normally absent in laboratory studies. We also wondered whether the predicted mismatch between beliefs and revealing information would be smaller among officers than among lay respondents.

Third, *effective sources of information in police versus nonpolice contexts were compared for the police sample*. If officers acquired certain special knowledge about useful information to detect deception, do they use this knowledge only at work or also in their personal life? Are lies in family or relational contexts (where mutual trust is the norm) detected the same way as lies in police–suspect interactions?

Fourth, the present research also differs from Park et al.'s (2002) on *methodological grounds*. While Park et al.'s study was only descriptive, we used inferential statistics comparing behavioral versus nonbehavioral information, as well as sample and type of context, and sought to replicate Park et al.'s findings with older participants more representative of the entire population than undergraduates. Successful replication would strengthen Park et al.'s conclusions and the corresponding components of TDT. The replication attempt involved not only Park et al.'s findings concerning the types of information useful to detect lies, but also those concerning the time at which the lie is detected; this is relevant for Propositions 10–12 of TDT (see Table 1).

Method

Participants

Participants were 22 local officers and 22 community members. Officers' data were collected during a workshop given by J.M. at a police department in Spain. Officers' average length of job experience was $M = 13.5$ years, $SD = 8.6$, $Mdn = 11.0$. Community members' data were collected in public areas in the same town. They were selected so that each officer had his/her own control, matched in terms of gender (23% females in each sample) and age (in both samples, $M_{age} = 38.6$ years, $SD = 6.6$; $Mdn = 37.5$ range: 30–56 years).

Procedures

The participants first filled in a questionnaire (Q1) with the open prompt "Please indicate how you believe lies can be detected." They were given as much time as they needed to answer this question in writing.

Q1 was collected and the participants were given a second questionnaire analogous to Park et al.'s (2002). Specifically, the questionnaire asked participants to first think of a lie they had detected in the past and then, with this in mind, to (a) describe in detail the circumstances of the lie and what the lie was about (Question 1), (b)

Table 2 Interrater Reliability of Codings (Kalphas)

Kind of Information	Dichotomic Data	Frequency Data
Behavioral cues	.81	.82
Visible	.92	.92
Verbal	.79	.77
Paralinguistic	.84	.84
Physiological	.96	.95
Unspecified	.76	.76
Nonbehavioral information	.86	.85
Third-person information	.95	.94
Evidence	.79	.79
Confession	.89	.88
Inconsistency with knowledge	.59	.59
Other information	.27	.27
Total	1.00	.98
Contradictions	.84	.84
Active strategies	.87	.87

report how long ago they were told that lie (Question 2), (c) indicate the kind of relationship they had with the liar (Question 3), (d) describe in detail how they had found out that the person had lied to them (Question 4), and (e) indicate the amount of time that passed between the moment they were told the lie and the moment they found out that was a lie (Question 5). Officers completed this second questionnaire twice, first with professional (i.e., police-related) contexts in mind (Q2) and then with personal contexts in mind (Q3), whereas nonofficers did so for personal contexts only (Q3).

Coder 1 examined the respondents' answers and:

1. Created a list of categories for the responses to both Q1 (beliefs) and Question 4 of Q2 and Q3 (revealing information), and tallied the number of times each participant gave an answer corresponding to each category. A successful attempt was made to use Park et al.'s (2002) categories to code nonbehavioral information. The categories are listed in Table 2 under "kind of information."
2. Verbal contradictions within the statement (considered as a verbal cue) happened to be cited somewhat frequently, and sometimes respondents reported using "active" strategies to catch the liar (such as asking specific kinds of questions, hiding some information, giving false information to see whether the liar corrected them, and so forth). Because of this, Coder 1 also coded *contradictions* and *active strategies* separately (in addition to also coding them within the specific corresponding categories).

3. Created categories for Questions 2 (how long ago the lie was told), 3 (kind of relationship with the liar), and 5 (time until discovery) of Q2 and Q3 and coded the respondents' answers.
4. Wrote a booklet with category definitions.

Coder 2 read the booklet and tallied the responses independently. Reliability was calculated separately for dichotomic data and frequency counts. Dichotomic data captured whether each respondent *mentioned/did not mention* any information belonging to a specific category (e.g., *third-person information*). Frequency counts refers to the *number* of cues pertaining to that category mentioned by each participant (e.g., any individual respondent may have mentioned several different items of information that fit within *third-person information*). Because of the nature of the data (dichotomic in addition to frequency counts, some categories having small frequencies, etc.), intercoder reliabilities were calculated with Krippendorff's (1970; see also Hayes & Krippendorff, 2007) Alpha (Kalpha). Results are displayed in Table 2. It is clear that Kalphas were satisfactory for all cue categories but *inconsistency with knowledge* and the residual *other information* category (which was not included in subsequent analyses). Kalphas were also good for the time the lie was told (*Kalpha* = 1.00), the time until discovery (*Kalpha* = .93) and the relationship with the liar (*Kalpha* = .94). Intercoder discrepancies were resolved by discussion.

Results

In all, 311 cues were reported: 183 in response to Q1 (106 mentioned by officers and 77 by community members), 57 in response to Question 4 in Q2 (only officers), and 71 in response to Question 4 in Q3 (39 by officers and 32 by community members).

Table 3 shows the percentage of respondents who mentioned each type of information. The most frequent kind of behavioral information mentioned when asked about beliefs (Q1) was *visible behavior*, though officers also mentioned *verbal cues*. However, the most revealing behavioral cues (Q2 and Q3) appeared to be *unspecified behavioral information* (behaviors that could not be easily coded as visible, verbal, paralinguistic or physiological, or that could be coded in several of these subcategories; e.g., "her behavior," or "nervousness"). Concerning nonbehavioral information, *evidence* was the kind of information cited most often overall. Park et al. (2002) found that *third-person information* was mentioned most often. We found this for community members responding to Q3, which was the condition most similar to Park et al.'s (see Table 3).

Influence of profession and measure

As the frequency distributions were skewed, the data were square-root transformed before running the statistical tests. However, for comprehensibility, descriptive (*Ms* and *SDs*) statistics below are based on the untransformed data (*Fs*, *ps*, and partial η^2 were calculated using the transformed data).

Table 3 Percentage of Respondents Mentioning at Least One Cue of Each Category in Responding to Q1 (Beliefs), Q2 (Revealing Information—Police Contexts), or Q3 (Revealing Information—Personal Contexts)

Kind of Information	Police Officers			Community Members		Overall	
	Q1	Q2	Q3	Q1	Q3	Q1	Q3
Behavioral cues	95.45	38.10	33.33	86.36	40.91	90.91	37.21
Visible	81.82	9.52	9.52	72.73	4.55	77.27	6.98
Verbal	81.82	23.81	9.52	22.73	13.64	52.27	11.63
Paralinguistic	27.27	9.52	0.00	40.91	0.00	34.09	0.00
Physiological	27.27	4.76	4.76	22.73	4.55	25.00	4.65
Unspecified	50.00	23.81	23.81	31.82	31.82	40.91	27.91
Nonbehavioral information	40.91	95.24	85.71	31.82	72.73	36.36	79.07
Third-person information	9.09	33.33	23.81	9.09	36.36	9.09	30.23
Evidence	40.91	71.43	52.38	27.27	27.27	34.09	39.53
Confession	0.00	4.76	14.29	0.00	4.55	0.00	9.30
Inconsistency with knowledge	0.00	14.29	9.52	4.55	4.55	2.27	6.98
Other information	27.27	23.81	9.52	9.09	4.55	18.18	6.98
Total (all kinds of cues together)	100.00	100.00	100.00	100.00	100.00	100.00	100.00

To examine the influence of profession (officer vs. community member) and measure (beliefs vs. revealing information) on type of information (behavioral vs. nonbehavioral), we first run a Profession $\times$ Questionnaire (Q1 vs. Q3) $\times$ Type of Information mixed Analysis of Variance (ANOVA) with repeated measures in the latter two factors. All three main effects were significant, revealing that more cues were reported by police officers ($M = 1.54$, $SD = 0.57$) than by community members ($M = 1.19$, $SD = 0.52$), $F(1, 42) = 5.68$, $p = .022$, partial $\eta^2 = .12$; in response to Q1 ($M = 1.94$, $SD = 0.93$) than in response to Q3 ($M = 0.79$, $SD = 0.59$), $F(1, 42) = 43.43$, $p < .001$, partial $\eta^2 = .51$; and if they were behavioral ($M = 1.98$, $SD = 1.15$) rather than nonbehavioral ($M = 0.75$, $SD = 0.66$), $F(1, 42) = 12.05$, $p = .001$, partial $\eta^2 = .22$. More interestingly, the Questionnaire $\times$ Information Type interaction was significant too, $F(1, 42) = 61.12$, $p < .001$, partial $\eta^2 = .59$. In support of our prediction, Bonferroni-adjusted post hoc tests revealed that whereas in responding to Q1 (beliefs) the participants mentioned significantly more behavioral ($M = 3.32$, $SD = 2.05$) than nonbehavioral ($M = 0.57$, $SD = 0.93$) cues, $F(1, 42) = 50.11$, $p < .001$, partial $\eta^2 = .54$, in responding to Q3 (revealing information) the participants mentioned significantly fewer behavioral ($M = 0.65$, $SD = 1.12$) than nonbehavioral ($M = 0.94$, $SD = 0.66$) cues, $F(1, 42) = 6.50$, $p = .015$, partial $\eta^2 = .13$. This was so both for officers and nonofficers (no interaction involving profession was significant).

A glance at Table 3 reveals that percentage data also point to the same conclusion. The percentage of both officers and community members who mentioned behavioral cues was greater when they were asked about beliefs (Q1) than when they were asked about useful information in the past (Q3, and also Q2). The reverse pattern emerged for nonbehavioral information. These conclusions generally stand (with a few exceptions) even when looking at each individual category of behavioral and nonbehavioral information (see Table 3).

Influence of context

To examine whether police officers mentioned different sources of information when prompted with professional versus personal contexts, a Questionnaire (Q1 vs. Q2 vs. Q3) × Type of Cues (Behavioral vs. NonBehavioral) within-participants ANOVA was run for the police sample. A significant main effect of questionnaire, $F(2, 42) = 14.28, p < .001$, partial $\eta^2 = .41$, indicated that more cues were mentioned in response to Q1 ($M = 2.20, SD = 0.85$) than in response to Q2 ($M = 1.24, SD = 0.65; p = .020$) or Q3 ($M = 0.88, SD = 0.69; p < .001$); the difference between Q2 and Q3 was not significant. The interaction was significant, $F(2, 42) = 28.53, p < .001$, partial $\eta^2 = .58$. Again, whereas more behavioral ($M = 3.68, SD = 1.96$) than nonbehavioral ($M = 0.73, SD = 1.12$) information was mentioned in response to Q1, $F(1, 21) = 28.25, p < .001$, partial $\eta^2 = .57$, the means differed significantly in the opposite direction for both Q2 (behavioral, $M = 0.77, SD = 1.15$; nonbehavioral, $M = 1.59, SD = 1.22; F(1, 21) = 9.72, p = .005$, partial $\eta^2 = .32$) and Q3 ($M = 0.59, SD = 1.10; M = 1.09, SD = 0.81$, respectively; $F(1, 21) = 7.48, p = .012$, partial $\eta^2 = .26$.) It is remarkable that no significant differences between Q2 and Q3 were found for either behavioral or nonbehavioral cues.

Table 4 Mean Number of Times Contradictions and Active Strategies Were Mentioned, and Percentage of Respondents Who Mentioned These Cues in Responding to Q1 (Beliefs), Q2 (Revealing Information — Police Contexts) or Q3 (Revealing Information — Personal Contexts)

Kind of Information	Police Officers			Community Members	
	Q1	Q2	Q3	Q1	Q3
Mean frequency of cues					
Contradictions	0.86	0.19	0.05	0.00	0.05
Active strategies	0.86	0.05	0.05	0.09	0.00
Percentage (dichotomic data)					
Contradictions	72.73	19.05	4.76	0.00	4.55
Active strategies	54.55	4.76	4.76	4.55	0.00

Profession differences: Contradictions and active strategies

In coding the data, we noticed that many officers mentioned *verbal contradictions* or using *active strategies* when asked about their beliefs, so we decided to code these two categories separately to see whether they revealed any difference between officers and community members. They did; almost all instances of these cues were mentioned by officers in response to Q1 (Table 4). With some empty cells ($M = 0.00$), we could not run ANOVAs on mean values. However, we did run Fisher's exact tests with the dichotomic data. This supported our impressions that in responding to Q1 more officers than community members mentioned contradictions ($p < .001$) and active strategies ($p = .001$). The difference was not significant for Q3.

Also, more officers (81.82%) than nonofficers (22.73%) mentioned verbal cues when asked about beliefs (Q1) (see Table 3; for Fisher's exact test, $p < .001$; no other officer-community member comparison was significant for the data in Table 3).

Times and relationship

Like Park et al. (2002), in Q2 and Q3 we asked how recent the lie was, how long it took to detect the lie, and what relationship the detector had with the liar. Observation of Table 5 reveals some interesting trends. Park et al. (2002) reported that 42.8% of the lies they studied had been told within the last month, and 81.0% within the last year. Our figures for officers in responding to Q2 (last month: 52.63%; last year: 73.68%) and for community members responding to Q3 (42.86% and 71.43%, respectively) are similar to Park et al.'s, but when responding to Q3 officers appeared to report older lies (no lie told the same day and almost half of the lies told more than 1 year before data collection; see Table 5).

We expected to replicate Park et al.'s (2002) findings that lies were *not* detected immediately and, in general, were *not* told by strangers. These predictions were supported for personal contexts (Q3) but not for police-related contexts (Q2).

Police contexts

Of the lies reported by police officers in response to Q2, all were told by strangers; 80% of them were detected in <10 minutes, and all of them within the same day (Table 5). This finding will be explained later in the Discussion section.

Personal contexts

In line with our predictions, few lies were detected immediately; around one-third were detected after one day or longer, and 14% only after one month or longer. Further, only between 5% (officers) and 9% (community members) of the lies were told by an unknown person; most were told by friends, neighbors, or acquaintances (Table 5).

Observation of the Q3 columns in Table 5 also reveals that the lies reported by officers were, in addition to older, primarily told by friends, neighbors, or acquaintances (61.90%) rather than by coworkers (4.76%) or by clients/patients/shop assistants (4.76%). This contrasts with the corresponding figures for community members (27.27%, 22.73%, and 22.73%, respectively; see Table 5).

Table 5 How Long Ago the Lie Was Told, How Much Time Elapsed Until Discovery, and Relationship With the Liar

	Police Officers		Community Members
	Q2	Q3	Q3
Time told			
<1 day	15.79	0.00	9.52
2 days to 1 week	21.05	14.29	19.05
>1 week to 1 month	15.79	19.05	14.29
>1 month to 1 year	21.05	19.05	28.57
>1 year	26.32	47.62	28.57
Time until discovery			
<10 minutes	80.00	33.33	31.82
10 minutes to 1 hour	10.00	9.52	27.27
>1 hour to 1 day	10.00	28.57	4.55
>1 day to 1 month	0.00	14.29	22.73
>1 month	0.00	14.29	13.64
Relationship			
Unacquainted person	100.00	4.76	9.09
Professional relationship (client, patient, shop assistant ...)	0.00	4.76	22.73
Workmate, classmate, boss ...	0.00	4.76	22.73
Friend, neighbor, acquaintance ...	0.00	61.90	27.27
Relative, partner ...	0.00	23.81	18.18

Note: Entries are percentages. A few respondents did not provide this information.

Discussion

This study supported our main hypothesis, replicated Park et al.'s (2002) findings, and provided data consistent with TDT. The same participants were asked how lies can be detected and then to recall how they detected a lie in the past. Despite mentioning primarily nonbehavioral, contextual information in response to the latter question, most participants mentioned primarily behavioral cues in response to the former one. In other words, participants stuck to their beliefs about deception cues despite their experience showing that other kinds of information are far more revealing. This is further evidence supporting the allure of observable behavior as indicative of deceit (see Bond et al., 2013), which might be based on widespread socialization practices (Global Deception Research Team, 2006) and people's difficulty in acknowledging that certain situational forces may impact behavior (Gilbert & Malone, 1995; Ross, 1977; Ross & Nisbett, 1991), for example pushing individuals to lie, in which case situational information could become a telltale deception cue (see Bond et al., 2013).

The same pattern of results stood for both police officers (with $Mdn = 11$ years of job experience) and nonstudent community members. Police experience does

not seem to correct the officers' wrong beliefs. However, some interesting differences did emerge between officers and nonofficers. First, officers provided more cues in general; this suggests that they are more observant or mindful concerning deception detection. Second, when asked about their beliefs, officers mentioned *verbal cues*, *contradictions*, and *active strategies* much more often than community members. In comparison with nonverbal behaviors, *verbal cues* have been shown to be more revealing of deception (DePaulo et al., 2003), yield greater detection rates (Bond & DePaulo, 2006), and are more useful to improve accuracy through training (Hauch et al., 2014). Officers' greater motivation to, and concern about detecting lies may lead them to focus more often on verbal content cues (see Reinhard & Sporer, 2008, 2010) and not to downplay the usefulness of verbal cues when asked about beliefs. It is remarkable that in professional contexts (Q2) the percentage of officers mentioning verbal cues was as high as the percentage of officers mentioning unspecified cues (see Table 3). This suggests sometimes officers successfully pay attention to verbal information to detect deception in professional contexts.

Concerning *contradictions*, even though at times they have been unrelated to veracity (e.g., Granhag & Strömwall, 2002), a meta-analysis shows that *discrepancies/ambivalence* (which is a construct that includes contradictions but also channel discrepancies and speaker's ambivalence) is a valid deception indicator (DePaulo et al., 2003). Also, according to Blair et al. (2012) and to Levine's (2014) TDT, a lack of coherence in message verbal content may at times signal deception (TDT's *Correspondence and Coherence* module). Finally, recent research stresses the value of the interviewer taking an *active approach* in order to increase sender's transparency (Hartwig et al., 2014; Levine, Shaw, & Shulman, 2010; Levine et al., 2014; Vrij & Granhag, 2012). In short, officers seem to be right in mentioning these three cues when asked about beliefs. Their mentioning an active approach is in consonance with Proposition 14 of TDT (Table 1). It is however surprising that contradictions and active strategies were hardly mentioned at all in response to Q2 and Q3. It is apparent that though these cues may be useful, in real life other kinds of information are still more revealing.

The previous paragraphs should not blind us to the fact that, just like community members, officers mentioned more behavioral than nonbehavioral information when asked how lies can be detected. Thus, there is a mismatch between officers' beliefs and useful information to detect real-life lies. This is worrisome. If officers believe that questionable behavioral cues are useful indicators they might consciously focus on these cues (and overlook more revealing information) when they suspect deception. This might decrease their accuracy rate in assessing suspects' credibility. It should be noted here that the fact that few behavioral cues were mentioned in response to Q2 and Q3 does not indicate that these cues are not used when trying to detect deceit outside the laboratory; it only shows that they are of little utility. Indeed, in a rare field study examining the nonverbal behavior of innocent and guilty suspects during their interactions with the police, nonverbal behaviors stereotypically associated

with deception proved to be faulty indicators (Johnson, 2007). It seems imperative to make officers aware of the discrepancies between their beliefs and useful information to detect lies, so that in real-life situations where veracity is under question, they consciously focus on reliable contextual information instead of focusing on fallible behavioral cues. This would arguably increase their detection accuracy. It may also be useful to show practitioners how inaccurate judgments based on behavioral (particularly on nonverbal) cues might be, and to discourage them from making judgments unless contextual information is available.

No differences were found within the officers' group between revealing information in professional (Q2) versus personal (Q3) contexts. Making a specific prediction was hard. It could be argued that behavioral cues would be mentioned *less* often in police-related than in personal contexts because—unlike occasional liars—well-practiced offenders may have learned to inhibit the few valid behavioral deception cues. Alternatively, it could be argued that behavioral cues would be mentioned *more* often in police-related than in personal contexts. Indeed, unlike lies told in law-enforcement contexts, generally everyday lies are trivial and inconsequential (DePaulo, Kashy, Kirkendol, Wyer, & Epstein, 1996); therefore, they probably produce little arousal, worry, or feelings of guilt that could ultimately be revealed through observable behavior. The absence of significant differences between the kinds of cues mentioned by the officers in responding to Q2 and Q3 provides no clue as to which of these alternatives is correct; it is possible that they cancel out.

We replicated Park et al.'s (2002) main finding that real-life lies are primarily detected through contextual rather than behavioral information, and we did so with older participants, hence more representative of the entire population. We found *evidence* to be the nonbehavioral category cited most often. This is consistent with Proposition 12 (Table 1) of TDT but differs from Park et al.'s finding that *third-person information* was the category mentioned by more participants. However, among community members responding to Q3—that is, the condition most similar to Park et al.'s—*third-person information* was the cue mentioned most often (Table 3). Thus, in reality our findings are in line with Park et al.'s.

It is noteworthy that among the behavioral cues reported when respondents were asked how they detected a lie in the past *unspecified* behavioral information stood out. This is in line with both DePaulo et al.'s (2003) meta-analytical findings showing that subjective general impressions may be more valid deception cues than objectively measurable isolated cues, and some authors' contention that observers trying to detect deception should focus on multiple cues from different channels instead of single cues (see Ekman, O'Sullivan, Friesen, & Scherer, 1991; Vrij, Akehurst, Soukara, & Bull, 2004).

Our data also coincide with Park et al.'s (2002) conclusions that in personal contexts most lies were not detected on the spot (see TDT Propositions 10 and 11 in Table 1) and were told by familiar persons. The only exception was for officers reporting about professional contexts. The peculiarities of these officers' professional duties

may explain this exception. These were local police officers dealing primarily with traffic offenses, gender violence, and minor street crime. Most offenders were not major criminals or recidivists, so they were unknown to the police, and most issues did not involve a lengthy police investigation but were resolved on the spot. Future research should examine police forces ordinarily involved in lengthy investigations of offenses committed by major criminals.

The present findings have theoretical implications for Levine's (2014) TDT. Suggestions have been made for additional components to be incorporated into TDT (Van Swol, 2014). Here we would like to suggest a new module and proposition to be incorporated. There is compelling evidence that even though the diagnostic utility of behavioral cues is poor (DePaulo et al., 2003; Levine, 2010; Sporer & Schwandt, 2006, 2007) and outside the laboratory lies are primarily discovered from contextual rather than behavioral information (Park et al., 2002; this study), people strongly believe behavioral cues to be indicative of deception (e.g., Global Deception Research Team, 2006; Strömwall et al., 2004) and are influenced by useless nonverbal information even when highly diagnostic contextual information is available (Bond et al., 2013). The strength of people's beliefs and the discrepancy between beliefs and useful information are nicely shown in this study, where the same individuals reported on both kinds of information. These findings are coherent with TDT and should be incorporated into the theory.

Limitations and caveats

Because the police data were collected during a workshop with all police participants doing the task at the same time, counter balancing questionnaire order was not possible. For experimental control, the same order was used later with community members. It may be argued that in responding to Q2 and Q3 (revealing information) respondents reported those possible cues they had not reported in responding to Q1 (beliefs). This is unlikely because (a) our findings for Q2 and Q3 nicely replicate those of Park et al.'s (2002) study where respondents had not been previously asked about beliefs, (b) we did *not* ask respondents not to repeat the same cues in different questionnaires and, in fact, there was some overlap, and (c) respondents were free to mention as many cues as they wished in responding to Q1, and to spend as much time as they needed to fill in the questionnaire. Future research may nevertheless examine this issue by reversing the order of questionnaires. However, asking respondents about revealing cues first may taint their reporting of believed cues, as participants may mention the just recalled revealing cues in addition to their general beliefs. Among officers, asking about professional lies first and about personal lies later made sense, as the data were collected during a professional seminar covering several eyewitness psychology topics.

Another methodological limitation (shared with Park et al.'s study) is that because in Q2 and Q3 we asked respondents only about lies *successfully* detected, it is uncertain (a) whether using the kinds of information reported in Q2 or Q3 *always* led to successful detection, and (b) how useful these kinds of information are to confirm

truths instead of detecting lies. It is also possible that a few of the reported "lies" were actually truths misjudged as lies by the respondents. However, the probability of this having occurred is low. Even though third-person information may not be a very reliable deception indicator—the third person may provide purposefully or unwittingly false information—evidence (which was the revealing cue mentioned by more respondents), the liars' admissions (with no pressure), and consistency with (accurate) knowledge look like fairly reliable indicators. In any case, more real-life deception research is needed to shed light on these issues.

We argued above that officers' belief that behavioral cues are revealing may lead them to focus on these cues instead of contextual information. However, Hartwig and Bond (2011) found in a meta-analysis that the cues people actually use when judging veracity differ from the cues often reported in research about beliefs. Yet, it should be noted that the cues people use are still behavioral (see Hartwig & Bond, 2011). However, this may be a result of the studies examined by Hartwig and Bond being laboratory-based and not providing receivers with information other than behavioral cues. Again, future real-life deception research should explore this issue.

Finally, it may be argued that the difference between beliefs and revealing cues was caused by people in real life being strongly motivated to find the truth and making the necessary cognitive effort to process more complex (and revealing) cues than the sender's behavior. Processing verbal cues requires more cognitive effort than processing nonverbal cues, and research shows that people with high task involvement, cognitive capacity, or need for cognition use verbal cues more often than people low in these features (Reinhard, 2010; Reinhard & Sporer, 2008, 2010). Just like processing verbal cues, processing revealing contextual information may require substantial cognitive effort. This is unlikely. First, some kinds of revealing information, like third-person information and confessions, require no extraordinary processing effort from the detector. Second, in accordance with DePaulo et al.'s (1996) findings, many of the lies reported by our participants were not very serious; thus, presumably the receivers' motivation or involvement was not extreme. Third, if this explanation was true one would expect many participants to mention verbal behavior in Q2 and Q3; this happened only among officers when reporting a "professional lie." It is therefore unlikely that the present findings can be explained in terms of the receivers' motivation or task involvement.

Conclusions

Recent meta-analytical findings showing that the validity of behavioral cues to deceit is very limited have yielded a shift in deception research. This shift is characterized by the interviewer adopting an active approach when questioning a potential liar and by taking contextual information into account. The present findings support the potential of contextual information to detect deception. This shift in deception research is promising, and it is our hope that it will move the field forward and will provide useful guidelines to practitioners in law enforcement charged with the daunting responsibility of identifying deceptive offenders.

Acknowledgments

We are grateful to Arantxa Vegas and Abelardo Periañez for their kind invitation to J.M. to give a workshop on eyewitness psychology to the local police force of Salamanca (Spain), and to Luisa Velasco for her assistance; the police data were collected during that workshop. We are also grateful to research assistants Susana Cano, Irene Ramírez, María Castaño, Natalia Cerletti, and Elena Montes for their help. Thanks are also due to Dr. Iris Blandón-Gitlin (California State University, Fullerton) for her insightful comments on this research, as well as to three anonymous reviewers and Dr. Malcom Parks, the Editor in Chief.

References

- Aamodt, M., & Custer, H. (2006). Who can best catch a liar? A meta-analysis of individual differences in detecting deception. *Forensic Examiner*, **16**, 6–11.
- Blair, J. P., Levine, T. R., Reimer, T. O., & McCluskey, J. D. (2012). The gap between reality and research: Another look at detecting deception in field settings. *Policing: An International Journal of Police Strategies & Management*, **35**, 723–740. doi:10.1108/13639511211275553.
- Blair, J. P., Levine, T. R., & Shaw, A. S. (2010). Content in context improves deception detection accuracy. *Human Communication Research*, **36**, 423–442. doi:10.1111/j.1468-2958.2010.01382.x.
- Bond, C. F., Jr., & DePaulo, B. M. (2006). Accuracy of deception judgments. *Personality and Social Psychology Review*, **10**, 214–234. doi:10.1207/s15327957pspr1003_2.
- Bond, C. F., Jr., Howard, A. R., Hutchison, J. L., & Masip, J. (2013). Overlooking the obvious: Incentives to lie. *Basic and Applied Social Psychology*, **35**, 212–221. doi:10.1080/01973533.2013.764302.
- Buller, D. B., & Burgoon, J. K. (1996). Interpersonal deception theory. *Communication Theory*, **6**, 203–242. doi:10.1111/j.1468-2885.1996.tb00127.x.
- DePaulo, B. M., Kashy, D., Kirkendol, S. E., Wyer, M. W., & Epstein, J. A. (1996). Lying in everyday life. *Journal of Personality and Social Psychology*, **70**, 979–995. doi:10.1037/0022-3514.70.5.979.
- DePaulo, B. M., Lindsay, J. J., Malone, B. E., Muhlenbruck, L., Charlton, K., & Cooper, H. (2003). Cues to deception. *Psychological Bulletin*, **129**, 74–118. doi:10.1037/0033-2909.129.1.74.
- Ekman, P., & Friesen, W. V. (1969). Nonverbal leakage and clues to deception. *Psychiatry*, **32**, 88–106.
- Ekman, P., O'Sullivan, M., Friesen, W. V., & Scherer, K. R. (1991). Face, voice, and body in detecting deceit. *Journal of Nonverbal Behavior*, **15**, 125–135. doi:10.1007/BF00998267.
- Frank, M. G., & Feeley, T. H. (2003). To catch a liar: Challenges for research in lie detection training. *Journal of Applied Communication Research*, **31**, 58–75. doi:10.1080/00909880305377.
- Garrido, E., & Masip, J. (1999). How good are police officers at spotting lies? *Forensic Update*, **58**, 14–21.
- Gilbert, D. T. (1991). How mental systems believe. *American Psychologist*, **46**, 107–119. doi:10.1037/0003-066X.46.2.107.

- Gilbert, D. T., & Malone, P. S. (1995). The correspondence bias. *Psychological Bulletin*, **117**, 21–38. doi:10.1037/0033-2909.117.1.21.
- Global Deception Research Team (2006). A world of lies. *Journal of Cross-Cultural Psychology*, **37**, 60–74. doi:10.1177/0022022105282295.
- Granhag, P.-A., & Strömwall, L. A. (2002). Repeated interrogations: Verbal and non-verbal cues to deception. *Applied Cognitive Psychology*, **16**, 243–257. doi:10.1002/acp.784.
- Granhag, P.-A., & Strömwall, L. A. (Eds.) (2004). *Deception detection in forensic contexts*. Cambridge, England: Cambridge University Press.
- Hartwig, M., & Bond, C. F. (2011). Why do lie-catchers fail? A lens model meta-analysis of human lie judgments. *Psychological Bulletin*, **137**, 643–659. doi:10.1037/a0023589.
- Hartwig, M., Granhag, P.-A., & Luke, T. (2014). Strategic use of evidence during investigative interviews: The state of the science. In D. C. Raskin, C. R. Honts, & J. C. Kircher (Eds.), *Credibility assessment: Scientific research and applications* (pp. 1–36). San Diego, CA: Academic Press.
- Hauch, V., Sporer, S. L., Michael, S. W., & Meissner, C. A. (2014). Does training improve detection of deception? A meta-analysis. *Communication Research*. doi:10.1177/0093650214534974.
- Hayes, A. F., & Krippendorff, K. (2007). Answering the call for a standard reliability measure for coding data. *Communication Methods and Measures*, **1**, 77–89. doi:10.1080/19312450709336664.
- Inglehart, R., Basáñez, M., Díez-Medrano, J., Halman, L., & Luijckx, R. (2004). *Human beliefs and values: A cross-cultural sourcebook based on the 1999-2002 values surveys*. México, DF: Siglo XXI Editores.
- Johnson, R. R. (2007). Race and police reliance on suspicious non-verbal cues. *Policing: An International Journal of Police Strategies & Management*, **30**, 277–290. doi:10.1108/13639510710753261.
- Krippendorff, K. (1970). Estimating the reliability, systematic error and random error of interval data. *Educational and Psychological Measurement*, **30**, 61–70. doi:10.1177/001316447003000105.
- Levine, T. R. (2010). A few transparent liars. *Communication Yearbook*, **34**, 41–61.
- Levine, T. R. (2014).). Truth-default theory (TDT). A theory of human deception and deception detection. *Journal of Language and Social Psychology*, **33**, 378–392. doi:10.1177/0261927X14535916.
- Levine, T. R., Blair, J. P., & Clare, D. D. (2014). Diagnostic utility: Experimental demonstrations and replications of powerful question effects in high-stakes deception detection. *Human Communication Research*, **40**, 262–289. doi:10.1111/hcre.12021.
- Levine, T. R., Feeley, T. H., McCornack, S. A., Hughes, M., & Harms, C. M. (2005). Testing the effects of nonverbal behavior training on accuracy in deception detection with the inclusion of a bogus training control group. *Western Journal of Communication*, **69**, 203–217. doi:10.1080/10570310500202355.
- Levine, T. R., & Kim, R. K. (2010). Some considerations for a new theory of deceptive communication. In M. S. Glone & M. L. Knapp (Eds.), *The interplay of truth and deception: New agendas in communication* (pp. 16–34). New York, NY: Routledge.
- Levine, T. R., Kim, R. K., & Blair, J. P. (2010). (In)accuracy at detecting true and false confessions and denials: An initial test of a projected motive model of veracity judgments. *Human Communication Research*, **36**, 82–102. doi:10.1111/j.1468-2958.2009.01369.x.

- Levine, T. R., Kim, R. K., & Hamel, L. M. (2010). People lie for a reason: Three experiments documenting the principle of veracity. *Communication Research Reports*, **27**, 271–285. doi:10.1080/08824096.2010.496334.
- Levine, T. R., Park, H. S., & McCornack, S. A. (1999). Accuracy in detecting truths and lies: Documenting the “veracity effect.” *Communication Monographs*, **66**, 125–144. doi:10.1080/03637759909376468.
- Levine, T. R., Serota, K. B., Shulman, H., Clare, D., Park, H. S., Shaw, A., ... Lee, J. H. (2011). Sender demeanor: Individual differences in sender believability have a powerful impact on deception detection judgments. *Human Communication Research*, **37**, 377–403. doi:10.1111/j.1468-2958.2011.01407.x.
- Levine, T. R., Shaw, A., & Shulman, H. C. (2010). Increasing deception detection accuracy with strategic questioning. *Human Communication Research*, **36**, 216–231. doi:10.1111/j.1468-2958.2010.01374.x.
- Masip, J., Alonso, H., Garrido, E., & Herrero, C. (2009). Training to detect what? The biasing effects of training on veracity judgments. *Applied Cognitive Psychology*, **23**, 1282–1296. doi:10.1002/acp.1535.
- Masip, J., Garrido, E., Herrero, C., Antón, C., & Alonso, H. (2006). Officers as lie detectors: Guilty before charged. In D. Chadee & J. Young (Eds.), *Current themes in social psychology* (pp. 187–205). Mona, Jamaica: The University of the West Indies Press.
- Nahari, G. (2012). Elaborations on credibility judgments by professional lie detectors and laypersons: Strategies of judgment and justification. *Psychology, Crime & Law*, **18**, 567–577. doi:10.1080/1068316X.2010.511222.
- Park, H. S., Levine, T. R., McCornack, S. A., Morrison, K., & Ferrara, S. (2002). How people really detect lies. *Communication Monographs*, **69**, 144–157. doi:10.1080/714041710.
- Reinhard, M.-A. (2010). Need for cognition and the process of lie detection. *Journal of Experimental Social Psychology*, **46**, 961–971. doi:10.1016/j.jesp.2010.06.002.
- Reinhard, M.-A., Dahm, J., & Scharmach, M. (2012). Perceived experience and police officers' ability to detect deception. *Policing: An International Journal of Police Strategies & Management*, **35**, 822–834. doi:10.1108/13639511211275805.
- Reinhard, M.-A., Scharmach, M., & Sporer, S. L. (2012). Situational familiarity, efficacy expectations, and the process of credibility attribution. *Basic and Applied Social Psychology*, **34**, 107–127. doi:10.1080/01973533.2012.655992.
- Reinhard, M.-A., & Sporer, S. L. (2008). Verbal and nonverbal behavior as a basis for credibility attribution: The impact of task involvement and cognitive capacity. *Journal of Experimental Social Psychology*, **44**, 477–488. doi:10.1016/j.jesp.2007.07.012.
- Reinhard, M.-A., & Sporer, S. L. (2010). Content versus source cue information as a basis for credibility judgments. The impact of task involvement. *Social Psychology*, **41**, 93–104. doi:10.1027/1864-9335/a000014.
- Reinhard, M.-A., Sporer, S. L., Scharmach, M., & Marksteiner, T. (2011). Listening, not watching: Situational familiarity and the ability to detect deception. *Journal of Personality and Social Psychology*, **101**, 467–484. doi:10.1037/a0023726.
- Ross, L. (1977). The intuitive psychologist and his shortcomings: Distortions in the attribution process. *Advances in Experimental Social Psychology*, **10**, 173–220.
- Ross, L., & Nisbett, R. E. (1991). *The person and the situation: Perspectives of social psychology*. New York, NY: McGraw-Hill.
- Sporer, S. L., & Schwandt, B. (2006). Paraverbal indicators of deception: A meta-analytic synthesis. *Applied Cognitive Psychology*, **20**, 421–446. doi:10.1002/acp.1190.

- Sporer, S. L., & Schwandt, B. (2007). Moderators of nonverbal indicators of deception: A meta-analytic synthesis. *Psychology, Public Policy, and Law*, **13**, 1–34. doi:10.1037/1076-8971.13.1.1.
- Stiff, J. B., Miller, G. R., Sleight, C., Mongeau, P., Garlick, R., & Rogan, R. (1989). Explanations for visual cue primacy in judgments of honesty and deceit. *Journal of Personality and Social Psychology*, **56**, 555–564. doi:10.1037/0022-3514.56.4.555.
- Strömwall, L., Granhag, P. A., & Hartwig, M. (2004). Practitioners' beliefs about deception. In P. -A. Granhag & L. A. Strömwall (Eds.), *Deception detection in forensic contexts* (pp. 229–250). Cambridge, England: Cambridge University Press.
- Van Swol, L. M. (2014). Questioning the assumptions of deception research. *Journal of Language and Social Psychology*, **33**, 411–416. doi:10.1177/0261927X14528914.
- Vrij, A. (2008). *Detecting lies and deceit: Pitfalls and opportunities*. Chichester, England: Wiley.
- Vrij, A., Akehurst, L., Soukara, S., & Bull, R. (2004). Detecting deceit via analyses of verbal and nonverbal behavior in children and adults. *Human Communication Research*, **30**, 8–41. doi:10.1111/j.1468-2958.2004.tb00723.x.
- Vrij, A., & Granhag, P.-A. (2012). Eliciting cues to deception and truth: What matters are the questions asked. *Journal of Applied Research in Memory and Cognition*, **1**, 110–117. doi:10.1016/j.jarmac.2012.02.004.
- Zuckerman, M., DePaulo, B. M., & Rosenthal, R. (1981). Verbal and nonverbal communication of deception. *Advances in Experimental Social Psychology*, **14**, 1–59.