



Texting and the Language of Everyday Deception

Thomas Holtgraves & Elizabeth Jenkins

To cite this article: Thomas Holtgraves & Elizabeth Jenkins (2020) Texting and the Language of Everyday Deception, Discourse Processes, 57:7, 535-550, DOI: [10.1080/0163853X.2019.1711347](https://doi.org/10.1080/0163853X.2019.1711347)

To link to this article: <https://doi.org/10.1080/0163853X.2019.1711347>



Published online: 29 Jan 2020.



Submit your article to this journal [↗](#)



Article views: 141



View related articles [↗](#)



View Crossmark data [↗](#)



Texting and the Language of Everyday Deception

Thomas Holtgraves  and Elizabeth Jenkins

Department of Psychological Science; Ball State University; Department of Psychological Science, Ball State University; School of Communication Studies, Ohio University

ABSTRACT

Two experiments were conducted to examine the production and detection of common, everyday deception. Experiment 1 was a naturalistic study in which participants provided their most recent truthful and deceptive (both sent and received) text messages. Participants in Experiment 2 were asked to generate text messages that were either deceptive or truthful. Messages in both experiments were analyzed with the Linguistic Inquiry and Word Count (LIWC) program and presented to other participants for their judgments of truthfulness. LIWC analyses yielded both similarities (e.g., more negations in deceptive texts) and differences (e.g., more first-person pronouns in deceptive texts) with past deception research. In contrast to prior deception research, participants in both experiments were able to significantly differentiate between deceptive and nondeceptive messages, and some of the LIWC variables that differentiated deceptive from nondeceptive texts were significantly related to judgments of truthfulness.

Introduction

Deception, typically defined as the intentional, conscious, and purposeful misleading of another person, appears to be both relatively common (George & Robb, 2008; Hancock, Thom-Santelli, & Ritchie, 2004; Levine, 2014) and multifaceted (Burgoon, Buller, Guerrero, Afifi, & Feldman, 1996). Because deception is multifaceted, it is important to establish whether findings regarding deception generalize across different types of deception and modalities. In this research we examined everyday deception, or what some researchers refer to as small and easy lies (Harwood, 2014) or white lies (DePaulo & Kashy, 1998; DePaulo, Kashy, Kirkendol, Wyer, & Epstein, 1996), and examined them as both naturally occurring deceptive texts and as text messages created and judged in laboratory contexts. Our primary focus was on the linguistic markers of deception and the ability of people to recognize these types of deception.

Detecting deception

There is a long history of research examining the ability of people to detect others' lies. For example, several meta-analyses have been conducted in recent years that have examined the influence of individual differences (Aamodt & Custer, 2006; Bond & DePaulo, 2008) and training (Hauch, Sporer, Michael, & Meissner, 2016) on deception detection abilities. Overwhelmingly, these meta-analyses demonstrate that deception detection is often no better than chance, regardless of intense training (Hauch et al., 2016), working in a career field that requires deception detection ability (e.g., a police officer), confidence levels, age, experience, or sex (Aamodt & Custer, 2006).

Although deception detection rates are quite low, various moderators have been examined in this regard, such as motivation (DePaulo, Kirkendol, Tang, & O'Brian, 1988), deception medium (Bond & DePaulo, 2006), cues (e.g., transgressions, planning, duration, etc.; DePaulo et al., 2003), level of

expertise (Burgoon, Buller, & Guerrero, 1995), amount of prior interpersonal interaction (e.g., strangers versus friends; Anderson, DePaulo & Ansfield, 2002; Morris et al., 2016), confidence in judgment (DePaulo, Charlton, Cooper, Lindsay, & Muhlenbruck, 1997), and interactant as well as observer effects (Burgoon & Buller, 1994), to name a few. These and other moderators have demonstrated mixed results concerning the overall impact on deception detection.

Language of deception

A relatively recent development has been the systematic examination of some of the linguistic features of deception, particularly when people communicate using computer-mediated communication (CMC) modalities (e.g., texting, e-mail, etc.; Hancock et al., 2004). Several linguistic features have emerged in this research. Some of these studies have demonstrated that people engaging in deceptive CMC tend to use fewer self-references (Toma & Hancock, 2012; Zhou, Burgoon, Zhang, & Nunmaker, 2004) and fewer first-person singular pronouns (Hancock, Curry, Goorha, & Woodworth, 2008; Newman, Pennebaker, Berry, & Richards, 2003) than those who are not deceptive (see Van Swol, Braun, & Malhotra, 2012 for an exception). Some authors posit that people use fewer self-references and fewer first-person singular pronouns to distance themselves psychologically from the deception (Toma & Hancock, 2012).

In addition, deceptive messages in CMC are more likely to contain more words than nondeceptive CMC messages (Hancock et al., 2008; Zhou et al., 2004). Also, liars communicating via CMC tend to use more verbs overall (Zhou et al., 2004), specifically more motion (Newman et al., 2003) and modal (i.e., could, should, might, etc.; Zhou et al., 2004) verbs than truthful people. A negative slant to deceptive CMC also seems apparent because liars tend to use more negations (Toma & Hancock, 2012) and more negative emotion words (Newman et al., 2003) than nonliars.

Note, however, that the language of deception will likely vary considerably as a function of the context. For example, although some studies have demonstrated higher rates of negative emotion words for deceptive (vs. nondeceptive) messages (Newman et al., 2003), such effects do not always occur and in fact are reversed in some contexts such as messages on dating sites (Toma & Hancock, 2012). Similarly, although deceivers have been demonstrated to use more words overall than nondeceivers (Hancock et al., 2008; Zhou et al., 2004), this pattern is reversed in an on-line dating context (Toma & Hancock, 2012).

White lies, politeness, and face-work

DePaulo and colleagues (DePaulo & Kashy, 1998; DePaulo et al., 1996; Kashy & DePaulo, 1996) popularized the study of everyday lying. In two diary studies involving both college students and community members, these researchers documented the frequency with which people lie, the characteristics of these lies, and the motivations behind these lies. College participants averaged two lies per day (one of three social interactions on average); community members averaged one lie per day (one of five social interactions on average). Note, however, that a reanalysis of these data suggests that the high frequency of everyday lies may be due to a relatively small percentage of people who are very prolific liars (Serota, Levine, & Boster, 2010; but see Smith, Hancock, Reynolds, & Birnholtz, 2014). The reported motives for lying were most frequently self-serving and occurred primarily as a means of impression management. That is, people reported lying not for material gain but for psychological gain as a means of creating and presenting a more positive impression than is warranted. At the same time, not all lies were self-serving. Another common motive for lying was to protect the feelings of others, and close to 25% of all lies were of this type (termed “other-oriented lies”). With these lies people portray their feelings (and opinions, evaluations, preferences) as more positive than they really are. Although people told fewer lies to others with whom they were close, when they did lie to close others, the lies were more often altruistic rather than self-serving.

Other researchers have referred to these other-oriented lies as white lies or prosocial lies and have documented their general usefulness at lessening conflict and facilitating smooth interactions (Camden, Motley, & Wilson, 1984), although there is some evidence for cultural variability in the norms regarding prosocial lies (Lee, Cameron, Xu, Fu, & Board, 1997). A more specific type of white lie, termed a “butler lie,” has been documented as a strategy of using deception to manage the entry and exit of online social interactions (e.g., instant messaging). This form of lie is defined as “a particular linguistic strategy for social inattention ... which people use to avoid social interaction or account for a failure to communicate” (Birnholtz, Reynolds, Smith, & Hancock, 2013, p. 2231). Specifically, these lies are told in mediated conversations when impression management is important, such as when one communicator has a desire to leave the conversation or a lack of interest in communicating at all (Hancock, Woodworth, & Goorha, 2010). Participants in one study perceived butler lies to be an ordinary part of CMC and that both senders and receivers have expectations that these lies will happen via mediated communication (Birnholtz et al., 2013).

Other-oriented white lies also overlap with certain types of politeness. Politeness, as defined by Brown and Levinson (1987), refers to saying things in such a way so as to avoid threatening another person’s face or identity (Goffman, 1967). In other words, politeness is a linguistic means of managing face, that of both the speaker and the recipient. For example, when asked to respond to a question requesting potentially face-threatening information, respondents may sometimes dissemble and not provide the requested (truthful but negative) information. One way to do so is to avoid directly responding to the question, providing a subtle topic shift (Holtgraves, 1998) or equivocation (Bavelas, 1983; Bavelas, Black, Chovil, & Mullet, 1990). According to Bavelas, equivocation occurs when a person is faced with a conflict such that neither truth nor falsification is desirable. As result, a speaker may be intentionally vague, indirect, or evasive.

Although Bavelas does not consider equivocation to be deception (she instead regards it as manipulative discourse), it clearly falls within the typical definition of deception (intentionally, knowingly, and/or purposely misleading another person; Levine, 2014). In fact, many researchers do consider equivocation (and related discourse moves) to be deception. For example, in Burgoon et al.’s (1996) Interpersonal Deception Theory, equivocation is considered to be one of several types of deception, along with concealment and falsification. Moreover, their research demonstrates that people do consider equivocation to be a type of deception; observers in their studies perceived equivocations to be significantly less truthful than nonequivocation messages.

Present research

Although everyday deception (i.e., indirect replies, equivocation, butler lies, white lies, etc.) has been examined empirically in terms of its interpersonal antecedents and consequences (Birnholtz et al., 2013; Buller, Burgoon, Buslig, & Roiger, 1994; Buller, Burgoon, White, & Ebesu, 1994; Holtgraves, 1986, 1998), it has not been subject to an empirical examination of its qualities as deception. We pursued two general issues in this regard. First, do people recognize these types of utterances as deception, and if so, what linguistic features of messages drive this judgment? Although substantial research demonstrates that people typically are not able to detect deception successfully, there are reasons to believe these types of deception will be recognized. This is because these small types of lies generally serve face-management functions. That is, they allow speakers to convey potentially negative information in a way that helps manage the face of the interactants. However, this is accomplished only if the recipient recognizes the face-work that is being undertaken and that it is negative information that is being (nicely) conveyed (Holtgraves, 1998). Hence, individuals should be able to recognize these messages as deceptive. The second issue we pursued was an examination of the linguistic markers of this type of deception. To do this we both examined language variables that have previously been identified as being related to deception in digital contexts and conducted exploratory analyses of other possible linguistic markers of everyday deception. A corollary issue we

pursued was whether the language variables differentiating deceptive from nondeceptive texts would be the same variables that drove observers' judgments of deception.

Experiment 1

The purpose of this experiment was to examine the language of naturally occurring deception in text messages and the extent to which these deceptions would be recognized by others. This was a two-part study. In the first part participants provided deceptive text messages they had sent and received and text messages not involving deception that served as a control. In Part 2 the deceptive and nondeceptive messages were shown to a separate group of participants who were asked to judge the deceptiveness of each message. We expected these naturalistic deceptions to be recognized as more deceptive than the truthful control messages. Second, we used the 2015 version of the Linguistic Inquiry and Word Count (LIWC-15) program (Pennebaker, Booth, Boyd, & Francis, 2015) to identify language variables that differentiated deceptive from nondeceptive texts. One set of hypotheses was based on prior research examining deception in CMC; hence, we expected deceptive texts relative to nondeceptive texts to have more words, verbs, negations and negative emotion words, and fewer first- and third-person pronouns. In addition, we conducted exploratory analyses of other linguistic variables that differentiated deceptive from nondeceptive texts.

Methods

Participants

Participants were recruited from introductory psychology classes at a midwestern university and received partial course credit for their participation. Sixty-five students (14 men) participated in Part 1 and a separate group of 42 students (13 men) participated in Part 2.

Procedure

Participants in Part 1 were instructed to bring their cell phones with them to the lab so their text messages could be examined. During the session, participants were asked to access their phone and to provide the last 10 text messages they had sent, the last 10 deceptive text messages they had sent, and the last 10 deceptive text messages they had received. Participants typed each of their text messages into a computer. Deceptive sent text messages were defined as "texts in which you said something that may have been slightly different from how you actually felt or acted." Deceptive received text messages were defined as "texts in which the sender said something that may have been slightly different from how that person actually felt or acted."

For each text message, participants indicated how well they liked the recipient (1 = *dislike strongly* to 7 = *like strongly*), how close they were to the recipient (1 = *extremely distant* to 7 = *extremely close*), how long they have known the recipient (1 = *less than 24 hours* to 9 = *all my life*), the recipient's gender, relative age (1 = *much younger than me* to 5 = *much older than me*), and an open-ended question regarding the nature of their relationship with the recipient (e.g., friend, family, spouse or boy/girlfriend, roommate, other). For deceptive texts, participants rated (1 = *absolutely did not believe* to 7 = *absolutely did believe*) the extent to which they initially believed the deceptive message (for deceptive received) and the extent to which they believed the recipient believed their message (for deceptive sent). In addition, participants provided demographic information and completed several personality measures (after providing their text messages) including the self-monitoring scale (Snyder, 1974) and a brief 10-item measure of the five-factor model of personality (Ten-Item Personality Inventory; Muck, Hell, & Gosling, 2007). Analyses indicated no relevant effects with these variables, and they are not discussed further.

Part 2 was conducted online using the Qualtrics platform. Participants received one set of 20 text messages (10 nondeceptive texts and 10 deceptive texts sent) created by a participant in Part 1. Only texts from participants who provided a complete set of texts (20) were included ($n = 42$). Also, only

deceptive messages sent were included due to the possibility that the deceptive messages received may not have been deceptive. The text messages (absent any context) were presented in a random order, and for each text message participants rated the message in terms of their perceptions of the truthfulness of the text (1 = *not at all truthful* to 10 = *completely truthful*; Hancock et al., 2010), their confidence in their judgment of the truthfulness of the text (1 = *not at all confident* to 7 = *extremely confident*), and their perceptions of how believable the text would be to the recipient (1 = *not at all believable* to 7 = *completely believable*).

Results

Part 1

Examples of text messages that were rated (by the sender) as highly truthful and highly deceptive are presented in Table 1. Each text message was analyzed with LIWC-15 (Pennebaker et al., 2015), and the resulting linguistic categories were analyzed with a linear mixed-effects model that included text type (nondeceptive vs. deceptive sent vs. deceptive received) as a fixed effect and the intercepts for participants and text messages as random intercepts. When the Hessian matrix was not positive definitive, indicating the best estimate for the variance of the text message was 0, the text message variable was dropped from the analysis (indicated with † in all tables). In no cases did this change the results. We calculated and report semipartial effect sizes (R^2) for any variables that reached statistical significance (Edwards, Muller, Wolfinger, Qaqish, & Schabenberger, 2008).

Initial analyses indicated that liking for the recipient, $F(2, 1542.18) = 58.87$, $p < .001$, and closeness with the recipient, $F(2, 1541.59) = 33.76$, $p < .001$, were significantly lower for deceptive messages (both sent and received) than for truthful messages. Because our interest was in differences between deceptive and truthful texts, independent of the nature of the relationship between texter and receiver, we combined liking and closeness ($r = .77$) to use as a covariate in all subsequent analyses. We also examined whether recipients' age, gender, and relationship length varied as a function of type of text. There were no significant effects for any of these variables and so they were not included in any analyses.¹

We first conducted analyses for LIWC categories for which researchers have reported significant differences. These results are presented in Table 2. Consistent with past laboratory research, deceptive texts (both sent and received) contained significantly more negations (e.g., can't, didn't, don't) and negative emotions (e.g., afraid, fool, warn) than nondeceptive texts. In contrast to past laboratory studies, there were no differences between deceptive and nondeceptive texts for word count and third-person references (they and s/he). Finally, significant differences occurred for first-person pronouns. However, the direction was the opposite of that reported in prior research: Deceptive texts contained significantly more (rather than fewer) first-person pronouns than nondeceptive texts.

We then conducted exploratory analyses of the other LIWC linguistic categories, including the new composite variables included in LIWC-15 (clout, analytic, tone, authentic²). We used a false discovery rate criterion of 5%. The false discovery rate adaptively controls the false-positive rate for only those associations deemed significant rather than for all tests conducted (Benjamini &

Table 1. Sample Text Messages: Experiment 1

Truthful texts (sender's truthful rating = 10)
Hello!
How was class?
Deceptive texts sent (sender's truthful rating = 1)
Was I upset at all?
That's totally fine:) sounds good!
Deceptive texts received (receiver's truthful rating = 1)
Good talk.
Ok cool

Table 2. Means (SEs) for LIWC Categories as a Function of Text Type: Experiment 1

LIWC Category	Nondeceptive	Deceptive Sent	Deceptive Received	<i>F</i>	<i>R</i> ²
Word count	9.022 (.536)	9.106 (.567)	8.357 (.581)	1.313	
First-person pronoun	7.445 ^a (.580)	10.640 ^b (.635)	10.993 ^b (.657)	16.650**	.021
Third-person pronoun†	1.204 (.166)	0.993 (.187)	1.142 (.196)	<1	
Verbst	20.774 ^a (.723)	23.072 ^b (.807)	22.041 ^{ab} (.842)	2.929	
Negationst	2.465 ^a (.365)	4.432 ^b (.413)	3.670 ^b (.432)	7.289**	.009
Negative emotions	2.612 ^a (.542)	4.188 ^b (.594)	4.115 ^b (.616)	3.975*	.005

Row means without a superscript in common are significantly different at $p < .05$.

†Random intercept for text message not included in the model due to Hessian matrix not being positive definite. ** $p < .001$, * $p < .01$

Table 3. Means (SEs) for LIWC Categories that Differed Significantly as a Function of Text Type: Experiment 1.

LIWC Category	Nondeceptive	Deceptive Sent	Deceptive Received	<i>F</i>	<i>R</i> ²
Clout	53.295 ^a (1.828)	42.026 ^b (2.102)	43.829 ^b (2.088)	15.058	.019
Analytic	38.694 ^a (1.587)	26.182 ^b (1.781)	27.775 ^b (1.861)	20.246	.025
Pronouns	19.371 ^a (.726)	22.914 ^b (.803)	22.490 ^b (.834)	9.369	.012
Personal pronouns	14.065 ^a (.648)	16.739 ^b (.719)	17.551 ^b (.748)	9.553	.012
Articles†	2.547 ^a (.218)	1.811 ^b (.244)	1.820 ^b (.255)	4.187	.005
Auxiliary verb	11.287 ^a (.557)	13.109 ^b (.624)	13.246 ^b (.651)	4.293	.005
Function	46.767 ^a (1.127)	51.264 ^b (1.245)	50.349 ^b (1.294)	5.962	.008
Adjectives	4.087 ^a (.556)	6.263 ^b (.622)	7.635 ^b (.650)	11.223	.041
Interrogatives†	2.292 ^a (.272)	1.301 ^b (.307)	1.428 ^b (.321)	5.056	.005
Numbers†	2.257 ^a (.349)	1.155 ^b (.369)	0.937 ^b (.381)	7.334	.009

Row means without a superscript in common are significantly different at $p < .05$. The overall false discovery rate for reported effects is .05.

†Random intercept for text message not included in the model due to Hessian matrix not being positive definite.

Hochberg, 1995; Benjamini & Yekutieli, 2001). All effects reported in Table 3 survived the false discovery rate correction and had, on average, only a 5% probability of being false positives.

As can be seen in Table 3, nondeceptive texts were more analytic (formal, logical, hierarchical) and had higher levels of clout (i.e., confidence and low tentativeness) than deceptive texts. Also, there were significantly fewer articles, interrogatives, and numbers in deceptive texts (both sent and received) than nondeceptive texts, with the reverse occurring for pronouns, personal pronouns, function words, adjectives, and auxiliary verbs (i.e., more of each in deceptive texts relative to nondeceptive texts).

Part 2

Participants were clearly able to identify the truthfulness of the texts. Nondeceptive texts were perceived as significantly more truthful (7.065 vs. 6.606), $F(1,338.901) = 8.287$, $p < .01$, and believable (5.221 vs. 4.950), $F(1,338.813) = 7.127$, $p < .01$, than the deceptive texts. Confidence judgments were not significantly different ($p > .05$).

Table 4. Significant ($p < .05$) Predictors of Judged Truthfulness/Believability of Texts: Experiment 1

LIWC Category	<i>B</i>	<i>t</i>
	.041	2.309
Authentic	.001	2.019
Article	.075	2.279
Auxiliary verbs	−.042	−2.496
Interrogatives	.082	3.578

We then conducted an analysis to determine the linguistic variables that predicted observers' judgments of the truthfulness of texts. We combined observers' judgments of text truthfulness and believability into a single composite ($\alpha = .84$) and used a linear mixed-effects model to predict truthfulness/believability from the set of LIWC linguistic variables; LIWC variables were treated as fixed effects, and the intercept for the text messages by participants interaction was included as a random effect in the model.

There were five significant predictors (Table 4). Three of these predictors—articles, auxiliary verbs, and interrogatives—paralleled the earlier results. That is, more articles and interrogatives and fewer auxiliary verbs were associated with perceptions of greater truthfulness. Hence, these categories significantly differentiated deceptive from nondeceptive texts, and the Part 2 participants' judgments were significantly influenced, in the same direction, by these predictors. Although the authentic and word count categories did not differentiate deceptive from nondeceptive texts, both were significantly and positively related to observer's judgments of truthfulness.

Discussion

The purpose of this experiment was to examine natural, everyday deceptive text messages. Some of the findings paralleled past deception research. Specifically, deceptive texts contained significantly more negations, negative emotion words, and verbs (for deceptive sent) than nondeceptive texts. In contrast, some of the other linguistic differences (e.g. number of words, third-person pronouns) documented in previous deception research did not emerge in this study. Also, one difference was reversed in this study. Specifically, deceptive texts contained significantly more (rather than fewer) first-person pronouns than did the nondeceptive texts. There are at least two possible reasons for this reversal. First, naturalistic deception, unlike deception generated in a laboratory, often involves something about the speaker, something he or she did (or didn't do) or felt (or didn't feel) and so on, hence the heightened use of first-person pronouns. Note that this difference occurred for all pronouns and especially personal pronouns, not just first-person pronouns. Deceptive texts are about people, and as a result there are higher rates of personal pronouns. The second possibility concerns the nature of the nondeceptive texts in this study. Unlike laboratory studies in which participants are asked to either lie or be truthful, naturally occurring texts often do not have clear truth values. People are performing various speech acts, or actions, with their texts (e.g., greetings, questions, etc.). According to Speech Act Theory (Searle, 1969) these messages can be defective in various ways (e.g., the sender may not have the requisite beliefs or attitudes), but truthfulness is simply not relevant. It is possible, then, that the first-person pronoun reversal observed here is due to the nature of the nondeceptive comparison messages, an issue we address further in the General Discussion.

Although not predicted, there were significant differences for several of the new LIWC-15 composite variables.² Large and strong effects were observed for the clout and analytic categories; nondeceptive texts were significantly higher than deceptive texts on both measures. High scores on clout reflect confidence and expertise, whereas low scores indicate a more tentative and anxious style. Hence, the clout category may be capturing some of the tentativeness and uncertainty that DePaulo et al. (2003) argue are characteristic of deception (but see Hauch, Blandón-Gitlin, Masip, & Sporer, 2015). Scores on the analytic category reference variations in logical and hierarchical thinking. The low analytic scores for the deceptive texts likely reflect

a more narrative and less logical style. This is consistent with our finding of significantly fewer articles and more pronouns and negations in deceptive texts relative to nondeceptive texts, a pattern that Pennebaker, Chung, Frazee, Lavergne, and Beaver (2014) refer to as a dynamic or narrative language style. Note that as with the reversal for first-person pronouns, it is possible that these differences are a function of the type of speech acts performed with the nondeceptive texts.

Finally, the deceptive texts sent by the Part 1 participants were readily identified as being deceptive by the Part 2 participants. This is a striking finding because the judgments made by the Part 2 participants were based solely on the text itself; there was no information about the context. Were these judgments based on the same linguistic categories that differentiated deceptive from nondeceptive texts? Yes and no. Some of the categories (articles, auxiliary verbs, interrogatives) that differentiated deceptive from nondeceptive texts did play a significant role in judgments of truthfulness. For these categories, judges were relatively accurate and attuned to linguistic differences that appear to be markers of deception in these contexts. However, some categories that differentiated deceptive from nondeceptive texts were not related to judgments of truthfulness, and hence were missed by observers. Also two categories, word count and authenticity, played a role in truthfulness judgments, although neither significantly differentiated deceptive from nondeceptive texts.

Several potential limitations of this study should be noted. First, the deceptive texts received were potentially ambiguous because they were based on participants' judgments about whether a text that they had received was deceptive. This is why we used only the deceptive texts sent when examining judgments of deception. Still, the results for the deceptive sent and deceptive received texts were almost identical, providing support for the validity of the deceptiveness of the deceptive texts received. Second, we were not able to identify the specific type of deception that occurred in the deceptive texts, in part because the context was not known. Moreover, because we examined naturally occurring deception rather than manipulating deception, it is possible that there were differences in the type of speech acts being performed in the deceptive and nondeceptive conditions. So, to pursue these issues further we conducted an experimental study in which we manipulated deception and had participants produce deceptive and nondeceptive texts in response to specific contextual prompts for face-threatening information; other participants then judged the truthfulness of the texts.

Experiment 2

The purpose of this experiment was to examine deceptive text messages under controlled conditions. As in Experiment 1, this was a two-part study. We asked participants in Part 1 to respond to a hypothetical text message from another person with either a deceptive or a nondeceptive response. We expected to replicate the linguistic differences between the deceptive and nondeceptive texts found in Experiment 1. In addition, because prior research has demonstrated that power can influence concerns with face management, and hence the motivation to be deceptive (Koning, Steinel, van Beest, & van Dijk, 2011; Olekalns & Smith, 2009; Raven, 1992), we also manipulated the power of the recipient in this study. Participants responded to someone of higher power than themselves for half of the scenarios and to someone of equal power for the other half of the scenarios. We expected the predicted linguistic differences between the deceptive and nondeceptive messages to be greater when responding to someone higher in power. In Part 2, a new set of participants read the text messages generated by a participant in Part 1 and provided judgments regarding the truthfulness of each message. Based on the results of Experiment 1, we expected that the Part 2 participants would be able to successfully discriminate between deceptive and nondeceptive texts.

Methods

Participants

Participants were recruited from introductory psychology classes at a midwestern university and received partial course credit for their participation. Forty-six students (17 men) participated in Part 1 and a separate group of 46 students (12 men) participated in Part 2.

Materials

Eight scenarios (see [Appendix 1](#)) were constructed. Each scenario was three to six sentences and consisted of background information and a description of a main character who “sends” the participant a text message. Part 1 participants were asked to imagine that they received this text and to create a text message in response. For four scenarios they were told to tell the truth and for four scenarios to tell a white lie. In addition, power was manipulated; the scenario character was either relatively high in power (e.g., a boss) or not (e.g., coworker). It was thus a 2 (high-power vs. equal-power) $\times$ 2 (deceptive vs. nondeceptive) within-subjects factorial design. Four booklets were created with each having an equal number of each power–truth combination. Across the experiment, an equal number of participants saw each version of each scenario.

Procedure

Participants were run in small groups of up to five participants and completed the study at a computer terminal. Participants were told they would be reading scenarios and responding (via text) to a question from the character in the scenarios and that for some scenarios they would be asked to tell the truth and for others they would be asked to tell a white lie. Participants were told that “White lies are common in everyday interactions. These small lies are often not meant to hurt the receiver, rather oftentimes people tell these lies to actually help the other person.” Participants completed two practice trials and then responded to the eight experimental scenarios. The order of the scenarios was randomized for each participant. Following Hancock et al. (2010), participants were instructed to review and rate the truthfulness of each of their text messages on a 10-point scale (0 = *not at all truthful* to 10 = *completely truthful*). After completing all tasks participants provided basic demographic information (age, gender, and ethnicity).

Participants in Part 2 were told that the purpose of the experiment was to detect deceptive text messages. Participants were told they were going to read scenarios that had been given to another participant who had provided a text message in response to the character in the scenario. They were also told that some of these responses were truthful and others were white lies. Participants were provided with the same definition of white lies as used in Part 1. After completing two practice trials, participants read the eight scenarios and corresponding text messages. The order was randomized for each participant. For each text message, participants provided ratings regarding their perceptions of the truthfulness of the text message (0 = *not at all truthful* to 10 = *completely truthful*; Hancock et al., 2010), their confidence in their judgment of the truthfulness of the message (1 = *not at all confident* to 7 = *extremely confident*), and their perceptions of how believable the message would be to the recipient (1 = *not at all believable* to 7 = *completely believable*). In addition, participants indicated the extent to which each text message was hurtful to the receiver (1 = *not at all hurtful* to 7 = *completely hurtful*) and helpful to the receiver (1 = *not at all helpful* to 7 = *completely helpful*). After completing all tasks participants provided basic demographic information (age, gender, and ethnicity).

Results

Part 1

We first conducted analyses for LIWC categories for which researchers have reported significant differences. These LIWC categories were analyzed with a linear mixed-effects model that included

Table 5. Means (SEs) for LIWC Categories as a Function of Text Type: Experiment 2

LIWC Category	Nondeceptive	Deceptive	<i>F</i>	<i>R</i> ²
Word count	18.68 (.674)	15.00 (.547)	33.362**	.097
First-person pronoun	8.84 (.469)	9.11 (.461)	<1	
Third-person pronoun	6.30 (.456)	6.39 (.532)	<1	
Verbs	25.83 (.683)	25.85 (.711)	<1	
Negations	4.16 (.416)	5.66 (.528)	6.842*	.021
Negative Emotions	3.16 (.356)	2.89 (.434)	<1	

p* < .01 *p* < .001

deception and power as fixed effects and the intercepts for participants and text messages as random variables. The results are presented in Table 5.

Consistent with Experiment 1, there was a significant effect for negations; deceptive texts again contained more negations than truthful responses. The effects for first -and third-person pronouns were in the same direction as Experiment 1 but were not significant. In contrast to Experiment 1, there was a significant effect for word count, with truthful responses containing more words than deceptive responses. The effect of power and its interaction with deception was not significant for any dependent measures.

As in Experiment 1, we conducted exploratory analyses (setting the false discovery rate to 5%) of the LIWC-15 linguistic variables, including the four composite measures (clout, analytic, authentic, tone). In this analysis the only significant effect occurred for tone, $F(1,312) = 8.622$, $p < .01$, $R^2 = .027$, with high values of tone (and hence more positive valence) for the deceptive texts (74.607) than for the nondeceptive texts (64.122).³

Part 2

The results for the raters' perceptions of the texts paralleled those in Experiment 1 (Table 6). Truthful responses were rated as more truthful and believable than deceptive responses, and again, there was no difference in confidence. Hence, as in Experiment 1, participants were able to detect deception. In addition, deceptive messages were rated as less hurtful and less helpful than the nondeceptive messages (Table 6). The effect of power and its interaction with deception was not significant for any dependent measures. We then conducted a linear mixed-effects analysis to determine the LIWC linguistic variables that predicted observers' judgments of the truthfulness of texts; LIWC variables were treated as fixed effects and the intercepts for participants and scenarios were included as a random effect in the model. The results are presented in Table 7.

There were four significant predictors (Table 7). Three of these predictors, word count, negations, and tone, significantly differentiated deceptive from nondeceptive texts, and the judgments of the Part 2 participants were significantly influenced, in the same direction, by these predictors. The other predictor, prepositions, did not differentiate deceptive from nondeceptive texts but was significantly and negatively related to observer's judgments of truthfulness.

Table 6. Means (SEs) for Perceptions of Deceptive and Nondeceptive Texts: Experiment 2

Judgment	Nondeceptive	Deceptive	<i>F</i>	<i>R</i> ²
Truthfulness	7.84 (.256)	4.92 (.255)	82.265**	.209
Believability	5.58 (.120)	5.20 (.116)	6.530*	.021
Hurtful	3.50 (.143)	2.79 (.124)	18.131**	.057
Helpful	4.09 (.150)	3.27 (.151)	18.976**	.055
Confidence	6.15 (.088)	6.00 (.097)	1.907	

p* < .05 *p* < .001

Table 7. Significant ($p < .05$) Predictors of Judged Truthfulness of Texts: Experiment 2

LIWC Category	<i>B</i>	<i>t</i>
Word Count	.084	2.926
Tone	-.016	-3.010
Negations	-.093	-2.073
Prepositions	-.087	-2.054

Discussion

Participants in this experiment were asked to produce text messages that were either truthful or deceptive. Hence, in this study we examined experimentally controlled deception rather than naturally occurring deception. And again, as in Experiment 1, participants were able to successfully differentiate between truthful and deceptive messages.

Consistent with Experiment 1, deceptive texts contained more negations than did the nondeceptive texts. In contrast to Experiment 1, there was a significant difference in word length, with deceptive messages containing significantly fewer words than truthful messages, a pattern that is the opposite of what has been reported previously (Hancock et al., 2008; Zhou et al., 2004). Truthful responses in this experiment would be hurtful as they conveyed negative information that would be threatening to the recipient; as a result, participants were motivated to keep their messages shorter. They may have attempted to follow an “If you can’t say something nice about someone don’t say anything” strategy. Importantly, word count was also a significant predictor of truthfulness judgments (as it was in Experiment 1); apparently, lay observers are sensitive to the veracity implications of shorter texts. Note also that deceptive texts in this study had a more positive tone than nondeceptive texts, a finding that contrasts with the first study where deceptive texts contained more negative emotion words (but not fewer positive emotion words) than nondeceptive texts. In this study, it is likely that participants went overboard in terms of praise when being deceptive. Importantly, observers picked up on this as well; judgments of truthfulness were negatively related to message tone.

General discussion

According to one of the most well-known theories of conversational meaning (Grice, 1975), truthfulness is a working assumption that guides conversations and the interpretation of conversation remarks. In other words, speaker truthfulness is the default assumption (see also Levine, 2014). However, people deceive for a variety of reasons and do so in a variety of ways. In this research, we examined everyday, digital deception, a type of deception often motivated by prosocial reasons. In two experiments we observed both similarities and differences with the findings of past deception studies. In both experiments, deceptive texts contained significantly more negations (e.g., no, isn’t, doesn’t, etc.) than nondeceptive texts, a finding that has been reported previously (Toma & Hancock, 2012). Deception, especially everyday deception, is about denial, saying that one doesn’t feel this way (e.g., that’s not true) or did not do something untoward (e.g., I didn’t do it) that would be threatening to the recipient. Note this reflects the interactional nature of these deceptions, in contrast to one-off deceptive messages created in the laboratory. Also consistent with prior research, deceptive naturalistic texts (i.e., in Experiment 1) contained more verbs and negative emotion words, both findings that parallel prior research (Newman et al., 2003; Zhou et al., 2004).

Some of our other results are unique and extend our understanding of what everyday deception looks like. Most notable in this regard was the cluster of language variables in the naturalistic study that were both positively (pronouns, especially personal pronouns, function words) and negatively (clout, analytic, articles) associated with deceptive texts. One of these patterns suggest a tentativeness and uncertainty (low clout) for deceptive texts, a pattern that has been noted previously by DePaulo

et al. (2003). Tentativeness may reflect the ambiguity that is often used when conveying negative, face-threatening information, as with indirect replies (Holtgraves, 1998). The other pattern is a tendency for deceptive texts to reflect a more narrative and less logical style (low analytic score, fewer articles, more pronouns and negations) (Pennebaker et al., 2014). That is, naturally occurring deceptive texts are relatively more immediate and involved than nondeceptive texts and are personal stories referencing specific people and things.

The elevated rates of first-person pronouns in deceptive messages in Experiment 1 is at odds with previous research (e.g., Hancock et al., 2008; Newman et al., 2003). Although this may reflect a less analytic and more narrative style in everyday deception, it is also possible that this difference is a function of the nature of the nondeceptive texts in this study. In laboratory studies of deception, the critical comparison is typically between intentionally deceptive and intentionally truthful messages; hence, all messages have clear truth values. In contrast, in everyday conversation truthfulness is not always relevant. This is because people are performing speech acts (e.g., greetings, requests, etc.), and truthfulness is simply not a relevant dimension (Searle, 1969). The differences we observe in Experiment 1 may be due in part to this different comparison.

Note also that because we examined naturally occurring text messages in Experiment 1, there was no control over the deception context. As has been noted by others (e.g., Smith et al., 2014), the occurrence of deception, as well as its detection, is heavily context dependent. Note in this regard that a larger number of variables distinguished deceptive from nondeceptive texts in Experiment 1 relative to 2. This is likely due in part to the more constrained and unnatural setting for Experiment 2 relative to the natural texts that likely occurred in a variety of contexts in Experiment 1 and to the nature of the nondeceptive texts in Experiment 1.

A final major finding, and one that occurred in both studies, was participants' ability to correctly identify deceptive and nondeceptive texts, a finding at odds with prior demonstrations of deception detection being only slightly better than chance. So, why were participants able to detect the lies in this research? In Experiment 2 the context was available to participants, and hence they could discern a likely motive for deception. Past research has demonstrated that one variable that improves deception detection is context, that is, when outside raters have background knowledge of the lie's intent and meaning. In these instances, these human raters are able to detect deception (Blair, Levine, & Shaw, 2010). In short, when there is an easily discernable reason for the lie, people are more likely to suspect deception (Levine, 2014). Note, however, that this does not explain the successful deception detection observed in Experiment 1 when there was no context. Participants in that study were able to accurately differentiate between deceptive and nondeceptive texts based solely on the text message itself.

To explore this issue, we examined linguistic variables that were associated with participants' judgments of message truthfulness. In effect, we examined the linguistic variables driving lay deception detection. Interestingly, we found judges were sensitive to some of the variables that distinguished deceptive from nondeceptive texts. Specifically, in Experiment 1 more interrogatives and articles and fewer auxiliary verbs were all significant predictors of judged truthfulness. Judges were also sensitive to variables, authentic and word count, that were not significantly associated with deception. And in Experiment 2, word count, tone, and negations were significant predictors of truthfulness, all variables that significantly distinguished deceptive from nondeceptive texts. Hence, participants in both studies were not only able to distinguish between deceptive and nondeceptive texts but did so based on linguistic variables that actually did differentiate deceptive from nondeceptive texts. In a Brunsickian lens model, then, these linguistic cues were both valid and utilized (Brunswick, 1956).

The ability to successfully identify deceptive texts may be one of the most noteworthy findings in this research. In our view it is the recognition that the deceptive sender is engaging in face-work that makes these lies recognizable. This is clearly the case for the deceptive texts produced in Experiment 2 as well as a subset of the texts generated in Experiment 1. These types of lies are designed to protect the face of the recipient (or sender) but simultaneously designed to allow for some recognition that

the information is not entirely positive. When a speaker with a negative opinion equivocates in his or her reply, the recipient might appreciate the lack of an overt negative opinion but also realize that the opinion is indeed negative. Consistent with this, participants in Experiment 2 rated the deceptive messages as significantly less hurtful than the truthful messages. Surprisingly, however, they also rated them as less helpful than the truthful messages, reflecting, perhaps, a view that ultimately it is more helpful in the long run to be truthful. The person who is sending the everyday lie, however, might be more likely to view the message as helpful (e.g., see DePaulo et al., 1996). A useful avenue for future research would be to examine differences between senders and receivers in their perceived motivation for these types of lies.

It is also possible that the ability to detect these lies within CMC may be based on the affordances of text-based communication modalities. That is, specific features of CMC may allow deception to be more easily detectable. For example, Walther (2011) articulated that various features of CMC could impact why people choose to communicate through specific mediums. Future research should study which specific features of CMC may add to the sender's perception of deception success or deception failure (e.g., asynchronicity, editability, missing cues).

In addition to the limitations noted above, we should note also how the use of a word-counting program like LIWC has certain limitations when studying digital deception. The counting of single words in isolation results in a fair amount of ambiguity. For example, the negative text message "It's not good" would be categorized by LIWC as positive due to the coding of "good" in isolation. Also, digital communication includes much more than simply plain text. Many messages include emotion indicators such as emoji (Panger, 2016), indicators that may be critical for determining the meaning of a message. Clearly, examining the role of emoji and emoticons in deceptive texting is warranted.

Overall, our general interpretation of the deceptive texts in both studies is that they were primarily other-oriented and, as such, were designed with a degree of opaqueness. Recipients could take them at face value (i.e., truthfully) and/or simultaneously recognize them for what they were, minor deceptions with the recipient, and sometimes with the speaker, in mind. Of course, this raises the issue of whether these texts should even be considered deception. At the conceptual and theoretical level that is an open issue. At the empirical level, however, they clearly constitute deception, according to participants.

There are many ways to be deceptive and a variety of motives for doing so. It is likely that different types of deception have different linguistic signatures and vary in their detectability. In this research we examined the production and comprehension of everyday, digital deception. Our results demonstrate some of the unique characteristics of this type of deception.

Notes

1. Although we were able to control for relationship differences, it is possible that there were other differences between deceptive and nondeceptive texts that we did not assess (e.g., being part of a conversational thread).
2. These LIWC composite variables are proprietary, and hence it is not possible to provide examples of the words used in their computation.
3. Significant effects for the power variable occurred for the clout and analytic categories. There were lower levels of clout when addressing a higher power recipient ($M = 45.52$) than a lower power recipient ($M = 52.08$), $F(1, 312.05) = 4.24$, $p = .04$, and analytic scores were higher when responding to the higher power recipient ($M = 14.61$) than when responding to the lower power recipient ($M = 10.16$), $F(1, 312.533) = 5.18$, $p = .023$.

Disclosure statement

No potential conflict of interest was reported by the authors.

Funding

The research was supported by grants from the National Science Foundation [BCS-1224553 and BCS-1917631] awarded to the first author. The second experiment was conducted as part of the second author's master's thesis under the direction of the first author.

ORCID

Thomas Holtgraves  <http://orcid.org/0000-0003-0764-6923>

References

- Aamodt, M. G., & Custer, H. (2006). Who can best catch a liar? A meta-analysis of individual differences in detecting deception. *The Forensic Examiner*, 15, 6–11.
- Anderson, D. E., DePaulo, B. M., & Ansfield, M. E. (2002). The development of deception detection skill: A longitudinal study of same-sex friends. *Personality and Social Psychology Bulletin*, 28, 536–545. doi:10.1177/0146167202287010
- Bavelas, J. B. (1983). Situations that lead to disqualifications. *Human Communication Research*, 9, 130–143. doi:10.1111/j.1468-2958.1983.tb00688.x
- Bavelas, J. B., Black, A., Chovil, N., & Mullet, J. (1990). *Equivocal communication*. Newbury Park, CA: Sage.
- Benjamini, Y., & Hochberg, Y. (1995). Controlling the false discovery rate: A practical and powerful approach to multiple testing. *Journal of the Royal Statistical Society*, 57, 289–300.
- Benjamini, Y., & Yekutieli, D. (2001). The control of the false discovery rate in multiple testing under dependency. *The Annals of Statistics*, 29, 1165–1188.
- Birnholtz, J., Reynolds, L., Smith, M. E., & Hancock, J. (2013). “Everyone has to do it.” A joint action approach to managing social inattention. *Computers in Human Behavior*, 29, 2230–2238. doi:10.1016/j.chb.2013.05.004
- 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(3), 214–234. doi:10.1207/s15327957pspr1003_2
- Bond, C. F., Jr., & DePaulo, B. M. (2008). Individual differences in judging deception: Accuracy and bias. *Psychological Bulletin*, 134, 477–492. doi:10.1037/0033-2909.134.4.477
- Brown, P., & Levinson, S. C. (1987). *Politeness: Some universals in language usage*. Cambridge, UK: Cambridge University Press.
- Brunswick, E. (1956). *Perception and the representative design of psychological experiments*. Berkeley: University of California Press.
- Buller, D. B., Burgoon, J. K., Buslig, A. L., & Roiger, J. F. (1994). Interpersonal deception VIII: Further analysis of nonverbal and verbal correlates of equivocation from the Bavelas et al. (1990) research. *Journal of Language and Social Psychology*, 13, 396–417. doi:10.1177/0261927x94134003
- Buller, D. B., Burgoon, J. K., White, C. H., & Ebesu, A. S. (1994). Interpersonal deception VII: Behavioral profiles of falsification, equivocation, and concealment. *Journal of Language and Social Psychology*, 13, 366–395. doi:10.1177/0261927x94134002
- Burgoon, J. K., & Buller, D. B. (1994). Interpersonal deception III: Effects of deceit on perceived communication and nonverbal behavior dynamics. *Journal of Nonverbal Behavior*, 18, 155–184. doi:10.1007/bf02170076
- Burgoon, J. K., Buller, D. B., & Guerrero, L. K. (1995). Interpersonal deception IX: Effects of social skill and nonverbal communication on deception success and detection accuracy. *Journal of Language and Social Psychology*, 14, 289–311. doi:10.1177/0261927x95143003
- Burgoon, J. K., Buller, D. B., Guerrero, L. K., Afifi, W. A., & Feldman, C. M. (1996). Interpersonal deception XXI: Information management dimensions underlying deceptive and truthful messages. *Communication Monographs*, 63, 50–69. doi:10.1080/03637759609376374
- Camden, C., Motley, M. T., & Wilson, A. (1984). White lies in interpersonal communication: A taxonomy and preliminary investigation of social motivations. *Western Journal of Speech Communication*, 48, 309–325. doi:10.1080/10570318409374167
- DePaulo, B. M., Charlton, K., Cooper, H., Lindsay, J. J., & Muhlenbruck, L. (1997). The accuracy-confidence correlation in the detection of deception. *Personality and Social Psychology Review*, 1, 346–357. doi:10.1207/s15327957pspr0104_5
- DePaulo, B. M., & Kashy, D. A. (1998). Everyday lies in close relationships. *Journal of Personality and Social Psychology*, 74, 63–79. doi:10.1037/0022-3514.74.1.63

- DePaulo, B. M., Kashy, D. A., Kirkendol, S. E., Wyer, M. M., & 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](https://doi.org/10.1037/0022-3514.70.5.979)
- DePaulo, B. M., Kirkendol, S. E., Tang, J., & O'Brian, T. P. (1988). The motivational impairment effect in the communication of deception: Replications and extensions. *Journal of Nonverbal Behavior*, 12, 177–202. doi:[10.1007/bf00987487](https://doi.org/10.1007/bf00987487)
- 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](https://doi.org/10.1037/0033-2909.129.1.74)
- Edwards, L. J., Muller, K. E., Wolfinger, R. D., Qaqish, B. F., & Schabenberger, O. (2008). An R2 statistic for fixed effects in the linear mixed model. *Statistics in Medicine*, 27(29), 6137–6157. doi:[10.1002/sim.v27:29](https://doi.org/10.1002/sim.v27:29)
- George, J. F., & Robb, A. (2008). Deception and computer-mediated communication in daily life. *Communication Reports*, 21, 92–103. doi:[10.1080/08934210802298108](https://doi.org/10.1080/08934210802298108)
- Goffman, E. (1967). *Interaction ritual: Essays in face to face behavior*. New Brunswick, NJ: Transaction Publishers.
- Grice, H. P. (1975). Logic and conversation. In P. Cole & J. Morgan (Eds.), *Syntax and semantics* (Vol. 3, pp. 41–58). New York, NY: Academic Press.
- Hancock, J. T., Curry, L. E., Goorha, S., & Woodworth, M. (2008). On lying and being lied to: A linguistic analysis of deception in computer-mediated communication. *Discourse Processes*, 45, 1–23. doi:[10.1080/01638530701739181](https://doi.org/10.1080/01638530701739181)
- Hancock, J. T., Thom-Santelli, J., & Ritchie, T. (2004, April). *Deception and design: The impact of communication technology on lying behavior*. Paper presented at the meeting of CHI, Vienna, Austria.
- Hancock, J. T., Woodworth, M. T., & Goorha, S. (2010). See no evil: The effect of communication medium and motivation on deception detection. *Group Decision and Negotiation*, 19, 327–343. doi:[10.1007/s10726-009-9169-7](https://doi.org/10.1007/s10726-009-9169-7)
- Harwood, J. (2014). Easy lies. *Journal of Language and Social Psychology*, 33, 405–410. doi:[10.1177/0261927X14534657](https://doi.org/10.1177/0261927X14534657)
- Hauch, V., Blandón-Gitlin, I., Masip, J., & Sporer, S. L. (2015). Are computers effective lie detectors? A meta-analysis of linguistic cues to deception. *Personality and Social Psychology Review*, 19(4), 307–342. doi:[10.1177/1088868314556539](https://doi.org/10.1177/1088868314556539)
- Hauch, V., Sporer, S. L., Michael, S. W., & Meissner, C. A. (2016). Does training improve the detection of deception? A meta-analysis. *Communication Research*, 43, 283–343. doi:[10.1177/0093650214534974](https://doi.org/10.1177/0093650214534974)
- Holtgraves, T. (1986). Language structure in social interaction: Perceptions of direct and indirect speech acts and interactants who use them. *Journal of Personality and Social Psychology*, 51, 305–314. doi:[10.1037/0022-3514.51.2.305](https://doi.org/10.1037/0022-3514.51.2.305)
- Holtgraves, T. (1998). Interpreting indirect replies. *Cognitive Psychology*, 37, 1–27. doi:[10.1006/cogp.1998.0689](https://doi.org/10.1006/cogp.1998.0689)
- Kashy, D. A., & DePaulo, B. M. (1996). Who lies? *Journal of Personality and Social Psychology*, 70, 1037–1051. doi:[10.1037/0022-3514.70.5.1037](https://doi.org/10.1037/0022-3514.70.5.1037)
- Koning, L., Steinel, W., van Beest, I., & van Dijk, E. (2011). Power and deception in ultimatum bargaining. *Organizational Behavior and Human Decision Processes*, 115, 35–42. doi:[10.1016/j.obhdp.2011.01.007](https://doi.org/10.1016/j.obhdp.2011.01.007)
- Lee, K., Cameron, C. A., Xu, F., Fu, G., & Board, J. (1997). Chinese and Canadian children's evaluations of lying and truth telling: Similarities and differences in the context of pro- and anti-social behaviors. *Child Development*, 68, 924–934. doi:[10.2307/1132042](https://doi.org/10.2307/1132042)
- Levine, T. R. (2014). Truth-Default Theory (TDT) A theory of human deception and deception detection. *Journal of Language and Social Psychology*, 33(4), 378–392. doi:[10.1177/0261927X14535916](https://doi.org/10.1177/0261927X14535916)
- Morris, W. L., Sternglanz, R. W., Ansfield, M. E., Anderson, D. E., Synder, J. L. H., & DePaulo, B. M. (2016). A longitudinal study of the development of emotional deception detection within new same-sex friendships. *Personality and Social Psychology Bulletin*, 42, 204–218. doi:[10.1177/0146167215619876](https://doi.org/10.1177/0146167215619876)
- Muck, P. M., Hell, B., & Gosling, S. D. (2007). Construct validation of a short five-factor model instrument: A self-peer study of the German adaptation of the Ten-Item Personality Inventory (TIPI-G). *European Journal of Psychological Assessment*, 23, 166–175. doi:[10.1027/1015-5759.23.3.166](https://doi.org/10.1027/1015-5759.23.3.166)
- Newman, M. L., Pennebaker, J. W., Berry, D. S., & Richards, J. M. (2003). Lying words: Predicting deception from linguistic styles. *Personality and Social Psychology Bulletin*, 29, 665–675. doi:[10.1177/0146167203029005010](https://doi.org/10.1177/0146167203029005010)
- Olekalns, M., & Smith, P. L. (2009). Mutually dependent: Affect and the use of deception in negotiation. *Journal of Business Ethics*, 85, 347–365. doi:[10.1007/s10551-008-9774-4](https://doi.org/10.1007/s10551-008-9774-4)
- Panger, G. (2016). Reassessing the Facebook experiment: Critical thinking about the validity of Big Data research. *Information, Communication & Society*, 19(8), 1108–1126. doi:[10.1080/1369118X.2015.1093525](https://doi.org/10.1080/1369118X.2015.1093525)
- Pennebaker, J. W., Booth, R. J., Boyd, R. L., & Francis, M. E. (2015). *Linguistic inquiry and word count: LIWC 2015*. Austin, TX: Pennebaker Conglomerates. Retrieved from www.LIWC.net
- Pennebaker, J. W., Chung, C. K., Frazee, J., Lavergne, G. M., & Beaver, D. I. (2014). When small words foretell academic success: The college admissions essays. *PLoS One*, 9, e115844. doi:[10.1371/journal.pone.0115844](https://doi.org/10.1371/journal.pone.0115844)
- Raven, B. H. (1992). A power/interaction model of interpersonal influence: French and Raven thirty years later. *Journal of Social Behavior and Personality*, 7, 217–244.
- Searle, J. (1969). *Speech acts: An essay in the philosophy of language*. Cambridge, UK: Cambridge University Press.
- Serota, K. B., Levine, T. R., & Boster, F. J. (2010). The prevalence of lying in America: Three studies of self-reported lies. *Human Communication Research*, 36, 2–25. doi:[10.1111/hcre.2010.36.issue-1](https://doi.org/10.1111/hcre.2010.36.issue-1)

- Smith, M. E., Hancock, J. T., Reynolds, L., & Birnholtz, J. (2014). Everyday deception or a few prolific liars? The prevalence of lies in text messaging. *Computers in Human Behavior*, 41, 220–227. doi:[10.1016/j.chb.2014.05.032](https://doi.org/10.1016/j.chb.2014.05.032)
- Snyder, M. (1974). Self-monitoring of expressive behavior. *Journal of Personality and Social Psychology*, 30, 526–537. doi:[10.1037/h0037039](https://doi.org/10.1037/h0037039)
- Toma, C. L., & Hancock, J. T. (2012). What lies beneath: The linguistic traces of deception in online dating profiles. *Journal of Communication*, 62, 78–97. doi:[10.1111/j.1460-2466.2011.01619.x](https://doi.org/10.1111/j.1460-2466.2011.01619.x)
- Van Swol, L. M., Braun, M. T., & Malhotra, D. (2012). Evidence for the Pinocchio effect: Linguistic differences between lies, deception by omissions, and truths. *Discourse Processes*, 49, 79–106. doi:[10.1080/0163853X.2011.633331](https://doi.org/10.1080/0163853X.2011.633331)
- Walther, J. B. (2011). Theories of computer-mediated communication and interpersonal relations. In M. L. Knapp & J. A. Daly (Eds.), *The handbook of interpersonal communication* (pp. 443–479). Thousand Oaks, CA: SAGE.
- Zhou, L., Burgoon, J. K., Zhang, D., & Nunmaker, J. F. (2004). Language dominance in interpersonal deception in computer-mediated communication. *Computers in Human Behavior*, 20, 381–402. doi:[10.1016/s0747-5632\(03\)00051-7](https://doi.org/10.1016/s0747-5632(03)00051-7)

Appendix 1

Experiment 2 Scenarios

Scenario 1

Your sibling (parent) looked terrible at a family event and seemed exhausted through the entire event. One hour after the event, your sibling (parent) texts you the following message: “Did I look terrible today? I hope I didn’t make a bad impression!”

Scenario 2

Your friend (parent) made you dinner one night. However, the meal was not very good and rather bland. One hour after you get home, your friend (parent) texts you the following message: “What did you think of my new recipe? I worked really hard to make you a great dinner!”

Scenario 3

Your acquaintance (supervisor) invites you to lunch to celebrate your birthday. He or she has told you multiple times that he or she is excited to give you your gift. Upon opening the gift, you realize that the gift is nothing to be excited about. One hour after getting home, your acquaintance (superior) texts you the following message: “What did you think about the awesome gift? I am so excited for you to have it!”

Scenario 4

Your coworker (boss) asks you to spend time together over the upcoming weekend, but you have other plans with someone else. You feel bad because you haven’t wanted to see your coworker (boss) in quite some time. One hour after your alternative plans take place, your coworker (boss) texts you the following message: “How did your weekend work out? I hope we can get together soon!”

Scenario 5

Your boss (sibling) told you to do something that you do not think is appropriate. So, you decide not to do it. One hour later, your boss (sibling) texts you the following message: “Were you able to get the task done? I really hope so because it was important!”

Scenario 6

Your parent (friend) met you for coffee and he or she was incredibly mean and rude. This was out of character for him or her. One hour after you get home, your parent (friend) texts you the following message: “Did you think I was rude today? I am really sorry if I was!”

Scenario 7

Your supervisor (acquaintance) made plans with you to spend time together one evening. At the last minute, you cancel because you hate spending one-on-one time with him or her, but you tell your supervisor (acquaintance) it is because you aren’t feeling well. One hour after you should have showed up, your supervisor (acquaintance) texts you the following message: “How are you feeling? I hope we get to spend time together soon!”

Scenario 8

Your boss (coworker) met with you for the fifth time this week. You are somewhat annoyed by how needy he or she has been lately and how much of your time he or she is taking up. One hour after meeting, your boss (coworker) texts you the following message: “Does it bother you that I take so much of your time? I promise to be less needy in the future!”