

**EVIDENTIALITY AS A DECEPTIVE FUNCTION: A
CROSS-LINGUISTIC STUDY IN ENGLISH, FRENCH, TURKISH,
AND JAPANESE**

by
SELMA BERFİN TANIŞ ŞAPCI

Submitted to the Graduate School of Social Sciences
in partial fulfilment of
the requirements for the degree of Master of Science

Sabanci University
July 2025

**EVIDENTIALITY AS A DECEPTIVE FUNCTION: A
CROSS-LINGUISTIC STUDY IN ENGLISH, FRENCH, TURKISH,
AND JAPANESE**

Approved by:

Assoc. Prof. ÇAĞLA AYDIN
(Thesis Supervisor)

Asst. Prof. JUNKO KANERO

Asst. Prof. SEÇKİN ARSLAN

Date of Approval: July 22, 2025

SELMA BERFİN TANIŞ ŞAPCI 2025 ©

All Rights Reserved

ABSTRACT

EVIDENTIALITY AS A DECEPTIVE FUNCTION: A CROSS-LINGUISTIC STUDY IN ENGLISH, FRENCH, TURKISH, AND JAPANESE

SELMA BERFİN TANIŞ ŞAPCI

Psychology, M.S. Thesis, July 2025

Thesis Supervisor: Assoc. Prof. ÇAĞLA AYDIN

Keywords: evidentiality, lying, deception detection, linguistics, crosslinguistic

This thesis study examined the linguistic structure of lying, focusing particularly on the use of evidentiality—a grammatical marker indexing the information source. In particular, this study investigated how the obligatory (evidential) and the optional (non-evidential) categories of grammatical evidentiality in language impact the production of certain linguistic structures as subtle cues of lying in deceitful versus truthful retellings. In a fully crossed, counterbalanced design, participants ($N = 217$) from language groups typologically diverse in evidential marking (i.e., non-evidential languages of English and French, evidential languages of Turkish and Japanese) provided written accounts of the events from brief stories under two manipulations: presentation modality (i.e., silent video clips or audio recordings) and veracity (i.e., either truthfully or deceitfully). Narratives were coded for the frequencies of the grammatical and lexical forms covering tenses, negations, and evidential markers, and then, cross-linguistically examined for their distribution across the conditions. The results revealed that lie tellers in evidential languages exploit direct evidentials in their statements—not the indirect ones, suggesting a pragmatic use of evidentials as firsthand accounts. The speakers of non-evidential languages used evidential markers without uniformity in deception conditions, such that perception verbs were fewer in English and affirmation adverb rates were higher in French, except both languages adopted a more negative tone in their lies, confirming the previous evidence. Overall, this study identifies an overlooked grammatical category, evidentiality, as a deceptive cue for the first time and provides a comprehensive approach to deception detection research cross-linguistically.

ÖZET

ALDATICI BİR İŞLEV OLARAK KANITSALLIK: İNGİLİZCE, FRANSIZCA, TÜRKÇE VE JAPONCADA DİLLER ARASI BİR ÇALIŞMA

SELMA BERFİN TANIŞ ŞAPCI

Psikoloji, Yüksek Lisans Tezi, Temmuz 2025

Tez Danışmanı: Doç. Dr. ÇAĞLA AYDIN

Anahtar Kelimeler: kanıtsallık, yalan, aldatma tespiti, dil bilimi, diller arası

Bu tez çalışması, yalan söylemenin dilsel yapısını incelemiş ve bilgi kaynağını dizinleyen bir dil bilgisi işareti olarak kanıtsallığa odaklanmıştır. Bu çalışma, özellikle kanıtsallığın zorunlu (kanıtsal) veya isteğe bağlı olduğu (kanıtsal olmayan) dillerde kullanımının, doğru ve yalan içerikli anlatımlar arasında, aldatmanın ipucu olarak belirli dilsel yapıların üretimindeki rolünü araştırmıştır. Tam çaprazlanmış, dengelenmiş bir tasarımda, tipolojik olarak farklı kanıtsallığa sahip dil gruplarından (kanıtsal olmayan İngilizce ve Fransızca, kanıtsal Türkçe ve Japonca) katılımcılar ($N = 217$), kısa öyküler biçiminde sunulan olayları iki ayrı manipülasyon ile (sunum formatı: sessiz video görüntüleri veya ses kayıtları; doğruluk: dürüst veya aldatıcı) dört farklı kombinasyonda yazılı olarak anlatmışlardır. Anlatılar, içerdikleri zaman kipleri, olumsuzluk ifadeleri ve kanıtsal belirteçleri kapsayan dil bilgisi ve sözcük biçimlerinin sıklıklarına göre kodlanmış ve koşullar arasındaki dağılımlar diller arasında incelenmiştir. Bulgular, kanıtsal dillerde yalan söyleyenlerin dolaylı kanıtlayıcıların aksine doğrudan kanıtlayıcıları daha sık kullandıklarını ve kanıtsal belirteçlerden birinci elden anlatımlar olarak faydalandıklarını ortaya koymuştur. Kanıtsal olmayan dilleri konuşan katılımcılar kanıtsal yapıları farklı biçimlerde kullanmıştır: aldatma koşullarında algısal fiiller İngilizcede daha az ve onaylama zarfları Fransızcada daha yüksek oranda gözlemlenmiştir. Bununla birlikte her iki dilde de aldatıcı anlatımlarda daha olumsuz bir ton benimsenmiş, önceki bulgular doğrulanmıştır. Bu çalışma, bugüne dek göz ardı edilen bir dil bilgisi kategorisi olan kanıtsallığı ilk kez aldatıcı bir ipucu olarak tanımlamış ve aldatma tespiti araştırmalarına diller arası kapsamlı bir yaklaşım sunmuştur.

ACKNOWLEDGEMENTS

Being a self-proclaimed, notoriously unfortunate human being, I have the notion that I must have made a deal early on in life to spend all my luck at cross paths with the *best* advisor, collabs, and family and partner of a lifetime that I could not have even dreamt of. I will always remain grateful to everyone who taught me anything, supported me in my growth through hell and high water, as well as the moments of shine.

The Best Advisor Oscar goes to my dearest and esteemed mentor, *Dr. Çağla Aydın*. Her ability to see kindness and curious questions in every little detail, come up with the most exciting ideas, and be an impeccable role model as a profound scientist led my path to being a strong, independent researcher for the last five years. I have learnt how to make science, how to teach, and how to be a co-worker from her brilliant mind, wisdom, and guidance. It has always been delightful to work beside her, think and talk with her. I deeply admire her for finding a balance in life, and am grateful for everything wholeheartedly, especially for bringing peace to our work, accommodating every challenge and adversity in calmness and presenting them in the most simplistic fashion possible to soothe our anxieties. I appreciate your decisions and stance as an advisor greater each passing day.

I would like to express my gratitude to my thesis jury committee members, *Dr. Seçkin Arslan* and *Dr. Junko Kanero*, for their valuable feedback and their contributions to and guidance as collaborators of this project. As a student eager to dig deep in and learn more about language research, I appreciate the opportunity of leading this project, which paved the way through this thesis study, allowing me to learn a thrilling research line and familiarize myself with other languages. Thereby, I am thankful to TÜBİTAK and the French Ministry of Foreign Affairs for providing a grant through a joint program (grant no: 122N747), and all the participants and coders from Sabancı University and Côte d’Azur University for their support in the realization of this thesis project.

I would also like to extend my warmest gratitude to the Clam Lab members, a.k.a. the *Clam girls*: My good time, bad time, and all-time friend *Yağmur Damla Şentürk*, for being the one that I run first to ask for help with the ease of knowing you never turn me back. You are one of the few people that I can trust with their perfection in science, blindfolded, who excel in everything, and achieve anything they do. I also

very much appreciate our times together with *Ege Ötenen* as being my first colleague and only graduate student supervisor, bringing joy to my life during the pandemic with 3-hour meetings on Friday nights. I have always cherished your felicitous tips and advice and am stunned by how on point they are. I look up to you for not losing your excitement as a scientist, even for a bit, during the hassles of all these years. You girls are the pure examples of how someone should be in academia. Thank you for everything we have done together as the *Dream Team*. With many *unpopular opinions* and common troubles in life, I owe special thanks to my loveliest friend, *Zeynep Tuna Bozkır*, for being one of the kindest, caring, and hardworking people. During the times of despair and fatigue, you have made some witchcraft and created all those moments of blissfulness with your great company and genuine care. The following lines are to *Seyma Kalender*, my comrade, fellow in this project, and my dear roommate. It has always been a pleasure to think about our work and be stuck together, not knowing what to do next. Your artsy, colorful soul inspires me and fuels me, whether in brief encounters or during our pillow talks on many nights. I would like to thank *Eda Melin Develioğlu* for taking responsibility all the time, and willingly helping me with my alternative thesis project, which will hopefully be the first of many future collabs. I am lucky to have you all girls around me, both academically and in my personal life.

I am also blessed with the most chaotic and fun family in human history, which I am deeply grateful for: my dear father, *ONT*, the man of principles, for teaching me how to be a moral person like himself and grow as an intellectual, which always sets the bar to excel in what I do and comply with ethics faithfully. And to my tender-hearted mother, *HHT*, for her never-ending love and endeavors for our family. My precious, vibrant sisters, *VBT*, *ZBT*, and *BBT*, for *literally* introducing me to the psychological science, and listening to me for endless hours. This appreciation extends to my second family, the *Şapcıs*: I am genuinely grateful to you for folding me in your arms as a daughter and a sister. Thank you for creating a safe haven that I can rest on during the most hectic times for the past few years.

My final words are to my one and only *AOBS*, my partner in excellence and harsh judgments as the soulmate in life. Thank you for being there as the one I trust the most, and reminding me of my true self when I feel bewildered and lost. Without you, I would not be able to rise above any bitterness nor finish this project with happiness. I am grateful for finding you at a young age, so as to have a long passage ahead and eternity together.

To my beloved parents and the memory of my dear grandparents

TABLE OF CONTENTS

ABSTRACT	iv
ÖZET	v
LIST OF TABLES	xi
LIST OF FIGURES	xii
LIST OF ABBREVIATIONS	xiii
1. INTRODUCTION.....	1
1.1. Lying and Deception Detection	2
1.2. Linguistics of Lying	4
1.2.1. Lexicon and Grammar of Lies.....	5
1.2.2. Cultural Implications	6
1.3. Epistemicity Markers: Evidentiality Across Languages.....	7
1.3.1. Epistemic Modality	9
1.3.2. Pragmatics of Evidential Use	10
1.3.2.1. Developmental trajectory.....	11
1.4. The Present Study: Evidentiality and Deception	13
2. MATERIALS AND METHODS	16
2.1. Participants	16
2.2. Materials.....	18
2.2.1. Stories.....	18
2.2.2. Auditory Stimuli	19
2.2.3. Visual Stimuli	19
2.3. Procedure	20
2.3.1. Experimental Procedure	20
2.3.2. Coding of the Retellings and Analysis	22

3. RESULTS	24
3.1. Data Analytics Strategy	24
3.2. Narrative Characteristics	25
3.3. Examination of the Deceptive Cues Across Conditions	26
3.3.1. Evidential Languages	26
3.3.1.1. Direct evidential markers	26
3.3.1.2. Indirect evidential markers	28
3.3.1.3. Grammar hoppings	29
3.3.2. Non-Evidential Languages	30
3.3.2.1. Tense hoppings	30
3.3.2.2. Perception verbs	31
3.3.2.3. Plausibility modals	31
3.3.2.4. Affirmation adverbs	32
3.3.2.5. Negation words	32
3.3.2.6. First-person pronoun	33
4. DISCUSSION	36
4.1. Limitations and Future Directions	39
4.2. Conclusion	41
BIBLIOGRAPHY	43
APPENDIX A	52
APPENDIX B	54
APPENDIX C	59
APPENDIX D	64
APPENDIX E	70
APPENDIX F	71
APPENDIX G	72

LIST OF TABLES

Table 2.1. General characteristics of participant groups by languages	17
Table 2.2. Mean (<i>SD</i>) frequencies of the norming study	19
Table 3.1. Mean (<i>SD</i>) frequencies of the narrative characteristics by con- ditions and languages	27
Table 3.2. Statistical outputs from mixed-effects models in evidential lan- guages.....	30
Table 3.3. Statistical outputs from mixed-effects models in non-evidential languages	35

LIST OF FIGURES

Figure 1.1. Associated reliability of an information source and evidential marking in Turkish	8
Figure 1.2. Associated reliability of an information source and evidential marking in Japanese	9
Figure 2.1. Experimental conditions.....	20
Figure 2.2. Example arrays from the visually animated video clips for story 1	21

LIST OF ABBREVIATONS

ACC	Accusative Case
CI	Confidence Interval
DIR	Direct Evidential Marker
ICC	Intraclass Correlation Coefficient
INDIR	Indirect Evidential Marker
INDIR/HR	Indirect Evidential/Heard Marker
INDIR/INFR	Indirect Evidential/Inferred Marker
IRR	Incidence Rate Ratio
LOC	Locative Case
M	Mean
SD	Standard Deviation
SE	Standard Error
TOP	Topic Marker

1. INTRODUCTION

Lying is inherently a linguistic occurrence as it is communicated through a verbal coding system (see Antomo 2025, for a counterargument on lying with gestures). While language can provide the ingredients and the tools required for telling a lie, the recipe and how the deceitful narratives were prepared determine their credibility; hence, the liar should manipulate and craft the language they use skillfully (e.g., Meibauer 2005; Meibauer 1992). In return, the receiver should analyze the message that the interlocutor conveys and assess its honesty without taking the information for granted (*epistemic vigilance*; see Sperber et al. 2010). Although the need for detecting lies and filtering misleading messages is widespread, the rising number of deceitful statements and news on social media has made this a critical area of application. As such, recent efforts to automate deception detection and cybercrime in social media have accelerated the scholarly work in psycholinguistics and natural language processing (see e.g., Elbatanouny et al. 2025; Mbaziira and Jones 2016, for a review).

In this context, language sets the cues of deception in certain linguistic forms. However, the literature primarily focuses on lexical, prosodic, and discourse-level indicators, and findings regarding how grammatical markers could potentially serve as cues are rather limited. One such marker is evidentiality, referring to how individuals grammatically mark the source of the information they deliver—whether the information was witnessed, heard from another source, or inferred based on indirect evidence Aikhenvald (2004). As evidentiality coding varies in different languages (e.g., optional and lexically expressed versus obligatory and grammatically encoded), this thesis asks the following question: How are evidential and other grammatical forms used in deceptive narratives across languages with varying forms of evidentiality? To this end, languages from two typologically distinct categories were selected and compared: Turkish and Japanese, which feature grammatical evidentiality, and English and French, which do not. In evidential languages, the grammar requires speakers to specify how they know what they are telling—such as through suffixes or particles—while in non-evidential languages, speakers can choose whether or not

to mark the information source, typically using optional lexical expressions (e.g., *I heard*, *apparently*). By examining how speakers use evidential markers when telling the truth versus lying, the aim is to uncover whether evidential strategies differ systematically across languages and whether evidential marking itself becomes a tool for deception. Moreover, by manipulating the modality of the information source, this study aimed to examine whether reliability attribution will differ, which in turn leads to the adoption of diverse strategies in evidential marking.

The following sections will start by presenting an overview of the linguistic and communicative characteristics of lying, and key findings on the lexical and syntactic cues of deceptive speech and their variation across languages will be detailed. Then, the concept of evidentiality will be examined in depth, considering its use in epistemic and pragmatic contexts and going through the target languages to demonstrate cross-linguistic variation. Finally, the literature on evidentiality and deception will be combined for a discussion.

1.1 Lying and Deception Detection

Lying has been broadly defined as the intentional act of passing deceitful information on to the recipients or asserting false beliefs (Meibauer 2005, 2014) to make them become common ground (Stokke 2013). Simpson (1992) describes the core features of lying as being intentional, aiming to put someone in error (except for *bald-faced lies*, see Stokke 2013, for a review) while hiding the intentions behind the deception, and differentiates it from falsehood by presenting insincere beliefs. What counts to be a lie (or a type of *deception*) has been a debated construct; however, it can be classified under diverse categories, such as paltering/misleading or deceptive implications (García-Carpintero 2023; Rogers et al. 2017), concealment (i.e., omission of key details; Van Swol, Braun, and Malhotra 2012), bald-faced lying (Meibauer 2016), and bullshitting (Meibauer 2018). Motivation types for telling a lie, on the other hand, can be examined under the categories of self-focused, other-focused, and altruistic (Ketterer and Hoerger 2014, 239) with two levels: high-stakes (i.e., personal gains are towering, such as denial of a crime) and low-stakes (i.e., enforcement is low, such as making a compliment as a form of white lie). The type of lie and the motivation behind it can significantly alter the composition of lies (e.g., Hancock et al. 2007; Le 2016); therefore, examination of lies requires a catch-all approach from diverse disciplines, or a well-defined framework is a must to understand the scope of a study. This paper specifically focuses on altruistic lies in denial scenarios

with low-stakes motivation. Please see the Methods section for an overview.

Methods to study lying behavior are diverse: Many studies adopt rather objective approaches, such as neuroimaging techniques (e.g., Merzagora et al. 2006; Lakshan et al. 2019), heart rate, and skin conductance measurements (e.g., Gödert, Rill, and Vossel 2001; Wang et al. 2022). However, linguistic examinations are potentially more convenient and accessible as they do not require real-time assessments and allow researchers to conduct analyses directly on big data sets of written statements. Existing examples from the literature:

Previous research on deception detection has demonstrated that a number of linguistic indicators can predict untruthful content. A common technique referred to as *Statement Analysis* (SA; e.g., Hwang, Matsumoto, and Sandoval 2016; Matsumoto, Hwang, and Sandoval 2015a; 2015b), which analyzes linguistic features to differentiate between true and false narratives, has been employed in many studies. Forensics, for instance, benefits from linguistic analysis of the eyewitness testimonies (e.g., Solà-Sales et al. 2023; Villar, Arciuli, and Paterson 2013).

Similarly, techniques utilizing network analysis (Taskin, Kucuksille, and Topal 2022), content analysis (Xu et al. 2015), speech signal (Ekici et al. 2017), manual and automated tagging of deception-marking linguistic patterns (e.g., unassertive language, negative expressions, inconsistent uses of verbs and nouns; Bachenko, Fitzpatrick, and Schonwetter 2008, or using large corpora were found thriving (e.g., Almela 2021; Levitan, Maredia, and Hirschberg 2018; Ott et al. 2011; Vogler and Pearl 2020). However, natural language processing technologies focusing solely on the language of the input still fall short in successfully operating for detecting deceitful content online (e.g., fake news, trolls, spammers) and require contextual information and additional approaches (e.g., stylometry: *content comparison*, *similarity detection*, and *metadata*; Gröndahl and Asokan 2019).

Therefore, this study particularly focused on the linguistic examination of deceptive narratives with the detailing of the diverse conditions (i.e., the modality of the information source) and context they applied. Below, the linguistic characteristics of deception will be defined first, then the contextual factors that influence them will be discussed.

1.2 Linguistics of Lying

While the acts of deception can be observed in masterfully designed sentences, simple mistakes can make the lie apparent to the listener, creating skepticism about the content. Accordingly, being the prosodic indicators of increased cognitive load, hesitations, pauses, errors in speech (Vrij 2008), decreased volume, rising intonation, higher pitch, and slower speech are perceived as uncertainty as well as the dishonesty of the speaker, regardless of the language spoken (French, English, or Spanish; Goupil et al. 2021). However, unlike uncertainty judgments, dishonesty judgments require additional contextual information to become consistent and hold stable across time (Goupil et al. 2021).

In addition to non-verbal cues of deception, previous literature focused on the linguistic and grammatical structure of deceitful statements. As opposed to the *Cognitive Strain Model* (Adams-Quackenbush 2015; Gombos 2006), which argues that lying constrains higher cognitive demands resulting in less complex narratives, the *Strategic Model* (Conway III et al. 2008) suggests that liars adjust their narratives based on the audience and deceptive needs (Repke, Conway III, and Houck 2018). Nevertheless, studies analyzing speech from TV shows and laboratory experiments found that the use of the filler word “um” was more frequent in true statements than deceptive ones, for instance (Arciuli, Mallard, and Villar 2010; Villar and Castillo 2017; Villar, Arciuli, and Mallard 2012), indicating that cognitive load eliciting cues such as filler words are not reliable sources to detect lies.

Previous findings point out that although lying increases cognitive load, people can use more complex linguistic structures in lie conditions, specifically for the perspective they have been lying from (*elaborative complexity*; e.g., Conway III et al. 2008; Repke, Conway III, and Houck 2018) and adapt their lies depending on the audience (e.g., Anolli, Balconi, and Ciceri 2003), supporting the Strategic Model. In a similar vein, Dor (2017) asserts that, by encouraging a complex use of grammatical structure, lying has served to evolve and enrich languages over time. Thus, a lie cannot be conceived and conceptualized apart from its linguistic form and complexity. As such, given this centrality, this thesis focuses on linguistic cues, in particular, the relationship between certain grammatical structures and lying. However, findings regarding the complexity of the deceptive narratives are mixed. The possible reasons behind mixed findings in the linguistic cues of deception will be discussed in the next section.

1.2.1 Lexicon and Grammar of Lies

DePaulo and colleagues' 2003 metareview on deception cues revealed that in comparison to truth tellers, lie tellers look tenser with a less positive attitude, and narrate less compelling stories with fewer details (Markowitz and Hancock 2014; Ott et al. 2011; Sarzynska-Wawer et al. 2023; Xu et al. 2015), having a smaller number of ordinary mistakes and eccentric content. These cues were more pronounced when the lie tellers had higher motivations to deceive, especially when related to their identity (DePaulo et al. 2003). Similarly, an increased number of words, sense-related words (e.g., *seeing, smelling*, Hancock et al. 2007; negative language, and generalizing terms (e.g., *everyone, no one, always, never*); a decreased use of self-references (e.g., *I, me, my, mine*; Newman et al. 2003; Solà-Sales et al. 2023; Vrij 2008, 101); conditional forms (Meibauer 2018), and inconsistent uses of tenses (Porter and ten Brinke 2010) were identified as the markers of deceptive statements. Also, negative emotions (e.g., afraid, warn, fool), profanity (especially when the audience is more suspicious; Van Swol, Braun, and Malhotra 2012), negations (Hauch et al. 2015; Vrij 2008), and the total number of verbs were found as deceptive cues (Holtgraves and Jenkins 2020).

However, some studies attested to the previous findings and found mixed results that false statements had more positive words than negative ones (Sarzynska-Wawer et al. 2023) and first-person pronouns (along with all other person referents) were more frequent in real-life examples of deceptive text messages (Holtgraves and Jenkins 2020). Additionally, studies from forensic investigations reported that true statements were longer and had fewer adjective uses (Solà-Sales et al. 2023). In contrast, false confessions were imbued with a decreased number of adjectives than true confessions (Villar, Arciuli, and Paterson 2013). These differences in findings can be due to a number of factors, such as participants' proficiency in lying, their motivations, and the characteristics of the data that is being handled (e.g., naturalistic vs. laboratory investigations; Villar, Arciuli, and Paterson 2013), where language use heavily depends on the context.

While these dependencies can significantly influence the presentation of lies, one crucial aspect here is to pinpoint linguistic features that elicit deceptive content specific to each language. Hence, as discussed in this section, grammatical and lexical cues to lying vary across contexts and languages, highlighting the need for cross-linguistic comparison, which is a key aspect of this thesis. In the next section, markers of lying across languages will be examined, and then another grammatical category, *evidentiality*, as a potential deceptive marker in certain languages, will be introduced afterward.

1.2.2 Cultural Implications

Given that culture shapes cognition in various domains (DiMaggio 1997; Henrich, Heine, and Norenzayan 2010; Quinn and Holland 1987), and grammatical structures differ across languages, cultural investigations on truth assessment are fundamental (see Rubin 2014, for cross-cultural considerations in Asian languages). Although some studies observed similar linguistic cues of veracity comparing truthful and deceitful statements across languages (Hwang, Matsumoto, and Sandoval 2016; Matsumoto, Hwang, and Sandoval 2015*a,b*), others found varying results (Laing 2015; Taylor et al. 2017), emphasizing the differences in motivations, skills, and cultural codes (e.g, individualism vs. collectivism). Thus, it is important to distinguish language-specific indicators of lying for the differences between languages.

The vast majority of literature on lying and deception revolves around English-speaking communities, which sheds light on the path through deception research in language on the one hand, and limits the perception toward the universals of the topic on the other. Research in additional languages (for instance, see Arabic; Yousef 2025; Chinese; Zhou and Su 2018; Italian; Spence, Villar, and Arciuli 2012; Polish; Sarzynska-Wawer et al. 2023) is handful and mostly restricted to the characteristics predetermined by the findings in English. To achieve a more comprehensive picture, a brief review of the findings of deception in languages that were selected to compare will be provided: French, Turkish, and Japanese.

For instance, a study on COVID-19 fake news on Twitter examined the use of hedges and modality words in French. The researchers identified the modal *falloir* (should, need) as signalling fake news and *devoir* (must) to mark true and *pouvoir* (can) to be related to any of them (Chiu et al. 2025). The results were in line with the expectation that the fake news writers tend to escape responsibility and *save face* (Boncea 2013) by using uncertain language. Another study investigating the effect of lying in autobiographical narratives for French speakers observed that when people lie about their experiences, they provided general information in their narratives due to not having an actual memory (Fekete 2019). On the other hand, automation of deception detection techniques in Turkish has recently increased (e.g, Ekici et al. 2017; Eskin 2024; Taskin, Kucuksille, and Topal 2022). A study utilizing hotel reviews in Turkish found that authentic comments were longer, included more complex sentences with a balanced use of pronouns, and required cognitive effort to process when reading (Akkol and Gökşen 2024). Lastly, in congruent with previous literature, one study focusing on the deceptive cues in a denial condition of the Japanese sample found that deceitful participants spoke faster with a lower pitch, and exhibited expressions of uncertainty such as using unfinished sentences and nervousness

(Danielewicz-Betz and Ogasawara 2013). Together, these findings demonstrate that grammatical and lexical cues to lying vary across contexts and languages. However, the literature on deception focuses on English as the base language, highlighting the gap and need for cross-linguistic comparison.

1.3 Epistemicity Markers: Evidentiality Across Languages

How people acquire knowledge and portray its source in their talk has been an intriguing question. A type of epistemicity and stance marker, *Evidentiality* (see Nuckolls and Michael 2014; San Roque 2019; for cross-cultural and cross-linguistic discussions) is the linguistic marker of the information source, which can be either through grammatical (Aikhenvald 2004) or lexical forms (Boye and Harder 2009; Cornillie 2007; Lazard 2001; Mélac 2022). Evidentiality is broadly categorized as *sensory-direct*, *hearsay-indirect*, and *inferential-indirect* (Aikhenvald 2004; Willett 1988); however, it can be codified through a smaller or larger number of markers depending on the evidential lexicon and the grammar available in a language. Although it is a debated construct, languages with grammatical evidentiality as an obligatory category (Lazard 2001) are usually referred to as *evidential languages*, whereas languages without grammatical marking, where evidentiality is an optional category, are referred to as *non-evidential languages*. In this paper, I will follow this categorization and consider English and French as non-evidential languages (i.e., evidentiality use is optional), while Turkish and Japanese will be the evidential languages (i.e., evidentiality use is obligatory).

In English, evidentiality can be expressed through predicates such as *I saw*, *it seems*, *I heard*, and adverbs such as *apparently*, *obviously*, and *seemingly*, *reportedly*, and through other epistemic and pragmatic strategies (Tantucci 2016). Similarly, in French, *J'ai vu*, *J'ai entendu*, and *sans doute* can be used to mark the evidentiality of a statement, where the speakers can also use distinct verbs to denote through which sensory source they acquired the information (see Desclés 2018, for a review on *stancetaking* through linguistic expressions).

In Turkish, for instance, one can use lexical forms to mark evidentiality, yet in grammatical terms, evidentiality is a binary construct codified through past tense morphemes. The morpheme signaling the information as directly witnessed or known is the *direct evidential* suffix *-DI*, whereas the one signaling secondhand information, such as heard or inferred, is the *indirect evidential* suffix of *-mİş* (Aikhenvald 2004). The following examples summarize the evidential concept in Turkish:

Figure 1.1 Associated reliability of an information source and evidential marking in Turkish

	Indirect	<—————>	Direct
			<i>-mİş</i> (hearsay and inference) < <i>-DI</i> (direct)
(1)	<i>Kedi</i>	<i>sokak-ta</i>	<i>kal-dı</i>
	Cat-NOM	outside-LOC	stay-ed-DIR
	“The cat stayed outside.” [witnessed: I saw it/know it]		
(2)	<i>Kedi</i>	<i>sokak-ta</i>	<i>kal-mış</i>
	Cat-NOM	outside-LOC	stay-ed-INDIR
	“The cat stayed outside.” [reported: I heard/inferred it]		

Unlike Turkish, there are *roughly* three types of grammatical evidential markers in Japanese. A direct evidential form does not involve an additional inflection, whereas usually, hearsay information is provided by the use of *rashii* and *soo da*, and inference is marked by *yoo da/mitai da* in one’s statement. The nuances between *rashii*, *soo da*, and *yoo da* have been discussed extensively (see Ishida 2006, for a brief overview), yet the distinction is still vague, and all are referred to as the indirect evidential markers (Karlsson 2013). However, because *yoo da* is based on one’s personal inferences through direct observation, *rashii* and *soo da* imply less authority in the narrative than *yoo da* (Matsumura 2017). Figure 1.2 exemplifies Japanese categories of evidentiality.

As summarized in this section, while English and French adopt similar evidential strategies, Turkish and Japanese use comparable ones with slight differences. To apply the previous findings from deception detection research in English to other languages and examine whether non-evidential and evidential languages use analogous linguistic features in deception, this thesis study focuses on four different languages and compares them in a lying context.

Figure 1.2 Associated reliability of an information source and evidential marking in Japanese

Indirect <—————> **Direct**

soo da (hearsay) $\leq$ *rashii* (hearsay and inference) $\leq$ *yoo da* (inference) $<$ *no marker* (direct)

- (1) ゆきは 歌を 歌った
 Yuki-wa uta-o utatta
 Yuki-TOP song-ACC sang-DIR
 “Yuki sang a song.” [witnessed: I saw it/know it]

- (2) ゆきは 歌を 歌った-ようだ
 Yuki-wa uta-o utatta-yoo da
 Yuki-TOP song-ACC sang-INDIR/INFR
 “Yuki sang a song.” [reported: “it seems”, inferred it]

- (3) ゆきは 歌を 歌った-らしい
 Yuki-wa uta-o utatta-rashii
 Yuki-TOP song-ACC sang-INDIR
 “Yuki sang a song.” [reported: “it seems”, I heard it/inferred it]

- (4) ゆきは 歌を 歌った-そうだ
 Yuki-wa uta-o utatta-soo da
 Yuki-TOP song-ACC sang-INDIR/HR
 “Yuki sang a song.” [reported: “I hear”, I heard it]

1.3.1 Epistemic Modality

Before moving on to the pragmatic uses of evidentiality, here, a closely related concept and another type of epistemicity marker will be introduced: *epistemic modality*. Even though they co-exist with evidential markers in evidential languages, they can be considered as the counterparts of evidential grammatical markers (see Boye

2012; Matthewson, Davis, and Rullmann 2007, for a discussion) in non-evidential languages to codify the factuality of a proposition (Palmer 2001, 24), yet should not be taken as direct correspondences (see, for instance, Aikhenvald 2004; Aksu-Koç 2016; De Haan 1999; Speas 2008).

Epistemic modals become the means of denoting uncertainty and the speaker’s proposition to not fully commit to the information conveyed. Probability notions, such as *might* and *may*, imply the absence of sufficient knowledge on an assertion, whereas plausibility notions, such as *must* and *should*, remark the inference based on observations (Desclés 2018). Consider the following example: when someone says “Tom might be at the conference” (probability), they indicate that there is a chance of Tom being at the conference. However, in the example of “Tom must be at the conference” (plausibility), the speaker implies that they have some knowledge at hand (e.g., knowing that Tom had applied to the conference for that date), hence they can infer about Tom’s location at the time. However, epistemic modalities differentiate from evidential markers; while the former denotes the factuality and the probability of the information and provides “epistemic support”, the latter denotes the relationship between the knowledge and the speaker (Palmer 2001, 24) as “epistemic justification” (Carretero, Marín-Arrese, and Ruskan 2022). Also, they are expressed through separate markers (Mélac and Leclercq 2024).

In summary, even though epistemic modality is occasionally considered together with evidentiality as the same, it is a separate yet adjacent linguistic category. However, it is important to distinguish them earlier in the discussion to extend the proposition to the pragmatics of evidentiality, as they can be adopted with similar epistemic stances and purposes in speech, especially in non-evidential languages.

1.3.2 Pragmatics of Evidential Use

In pragmatic terms, evidential and non-evidential languages use similar evidential strategies. However, speakers of obligatory use of evidential marking are reported to be more sensitive to the information source in conceptualizing their statements (Mushin 2013). While direct evidential markers strengthen the claims and assert authority in knowledge (Bergqvist and Grzech 2023), indirect evidentials are extensively considered less reliable as indicating secondhand information (Jahiu 2022), which implies that the original message could be lost in *narration*, such as misinterpreted or distorted by the conveyor (see McGlone and Baryshevtsev 2018).

On the flip side, using an indirect evidential removes the responsibility for the authenticity of an assertion from the speaker, because it serves as the cue of unattested

information (AnderBois 2014). In fact, omitting the reported verbs in speech indexes the original author as credible, such that the speaker behaves as if they witnessed the information firsthand (Ishida 2006) and takes the risk of being interpreted as having the statement as their own (Jahiu 2022). Conversely, when the reported speech is employed or the author is mentioned, the speaker tends to shift the responsibility to the original source, resulting in a clearance in their account (Mushin 2001). As a result, one can manipulate evidential markers to tell a lie, either by indexing a false source (e.g., narrating a directly accessed event *watched from a video clip* with indirect evidentials, or narrating an indirectly accessed event *heard from an audio recording* with direct evidentials) with a true statement or by indexing a true source with a false statement (Aikhenvald 2004, 5; cited in Jahiu 2022), depending on their pragmatic needs.

1.3.2.1 Developmental trajectory

The developmental trajectory for the comprehension of the information source has been examined across the target languages. Particularly, this section was included to emphasize the contrasts between languages with varying forms of evidential marking, spanning from childhood to adulthood. Here, understanding the disparities in the utility and the perception of source marking throughout cognitive and cultural development can be instrumental in grasping why speakers of evidential languages can manipulate the evidential uses in their statements and exposing crosslinguistic variations.

In Turkish, evidential uses are employed to codify whether a proposition is reliable or not by the speaker (Aksu-Koç 2016; Arslan 2020). While Turkish children start implementing evidential markers in their language around the age of 3 (Aksu-Koç 1988, as cited in Aksu-Koç, Ögel-Balaban, and Alp 2009), they often fail to comprehend the exact meaning attributed to them until the age of 6 (Aksu-Koç, Ögel-Balaban, and Alp 2009; Ozturk and Papafragou 2008). However, Turkish-speaking children’s comprehension of a statement’s reliability was affected by the evidentials, and they prioritized utterances with direct evidentials as more reliable than those with indirect ones (Ozturk and Papafragou 2016), earlier than their English-speaking peers. For instance, Aydin (2011) revealed that Turkish-speaking children as early as 4 found sentences implicating direct sources of information more accountable than a hearsay source. In comparison, English-speaking children at the age of 4 were more susceptible to misinformation than their Turkish-speaking counterparts when the information source was secondhand, and they could differentiate the source of infor-

mation at around the age of 8 (Aydin 2011; Aydin and Ceci 2013). Notably, though, both Turkish and English speakers of 3-4 year-olds recalled perceived sources better than heard sources (Baer et al. 2025).

In Japanese context on the other hand, when participants were presented with two contradicting statements in direct and indirect evidential forms about the location of an hidden object and asked to indicate its location afterward, children above the age of 5 and adult Japanese speakers found the statements with direct evidential forms superior in accountability than indirect forms and based their judgments on that (Matsui, Yamamoto, and McCagg 2006).

Similarly, a difference in reliability attribution to direct and indirect sources was observed between adult speakers of evidential and non-evidential languages. Tosun, Vaid, and Geraci (2013) found that Turkish speakers recognized sentences and their sources better when they were presented with direct evidentials than indirect evidentials. English speakers, on the other hand, expressed no differences between direct and indirect information, such that their memory of indirect assertions was better than their Turkish counterparts (Tosun, Vaid, and Geraci 2013). In a similar fashion, Arslan et al.'s (2024) findings on sentence ratings and simultaneous eye tracking measures of Turkish speakers indicated that statements with less witnessable content, as well as indirect evidentials and mismatches between the evidential marker and the source, were perceived more frequently as deceptive than the opposite conditions.

In terms of evidential marking, although comprehension comes with age, the emphasis on the knowledge source remains stable over time if the spoken language and cultural imposition are the same. To exemplify, native English speakers of Japanese second language learners use a decreased number of evidential markers in their narratives than native Japanese speakers (Ishida 2006; Matsumura 2017). Likewise, speakers of evidential languages use more evidential strategies in their second language (e.g., Aikhenvald 2002; and Slobin 2016 as cited in Filipović, Brown, and Engelhardt 2023), indicating that transfer effects occur in evidential marking, which is more acknowledged among speakers of languages with evidentials than without evidentials. However, these transfer effects are not observed in heritage speakers of evidential languages (Arslan, De Kok, and Bastiaanse 2017; Tokaç-Scheffer, Nickels, and Arslan 2024); hence, constant exposure to and practice of a language is essential in preserving this cognitive perspective.

In a nutshell, these findings indicate that practicing source-citing through a mandatory use of evidential grammar in a language can lead to differences in developmental trajectory and cognition of source monitoring across languages. Speaking an evi-

dential language may therefore make the emphasis on the relationship between the source and the speaker a linguistic habit. However, what is not known is their productive use in active deception. It can be inferred from their early use as young as 2-3 years, though they may be utilized when lying in evidential languages. These differences can further enhance our understanding of the evidentiality across languages to distinguish the cognitive and pragmatic approaches to evidential marking.

1.4 The Present Study: Evidentiality and Deception

Building on these insights, this study investigates whether the speakers of evidential and non-evidential languages use distinct linguistic strategies when lying, considering that languages with different grammatical structures are equipped with a diverse set of linguistic tools for different epistemic strategies. The key reasons behind this research line are as follows: (1) the previous research on the examination of evidentiality and lying is handful, and (2) studies on evidentiality mostly cover a pair of languages with and without grammatical evidentiality (or focusing only on languages from one category). Therefore, this study focuses on two languages from each category (i.e., non-evidential: English and French, evidential: Turkish and Japanese) and examines the manipulation of evidential use in deception to extend the existing knowledge on the topic and obtain cross-linguistic as well as language-specific patterns. More specifically, this study examines whether Turkish and Japanese speakers deliberately manipulate the use of evidential inflections in their statements to deceive others. Additionally, as opposed to an abundance of research on deception detection in Western languages (i.e., Germanic and Romance), here, a cross-linguistic perspective was adopted, aiming to test the previous findings on linguistic markers of deception across languages.

Given that the speakers of evidential languages to be more concerning toward the knowledge source in their talks (Mushin 2013) and develop source awareness in speech younger than non-evidential language speakers (e.g., Aydin and Ceci 2013), it is plausible to think that the speakers of evidential languages are more sensitive to certainty and the source of information (i.e., *source monitoring*; Johnson, Hashtroudi, and Lindsay 1993). Furthermore, because the emergence of grammatical evidentiality is discussed to be closely tied to the social functions of specific cultures (Michael 2015; San Roque 2019) and the selection of a particular evidential marker depends on context and the motivations of the speaker (e.g., *politeness* in Japanese, Filipović, Brown, and Engelhardt 2023; Hoyer 2008), it can be asserted

that people can actively manipulate the use of evidential markers (Aikhenvald 2004; Xu 2022) to make their statements sound more reliable or shift responsibility to another source. Therefore, the question is, how can these findings be applied to deliberate lying conditions?

To this end, the present study employed a 2 x 2 within-participants design to test the hypotheses. The first dimension is the *Veracity*, with two conditions: *Truth* and *Lie* as the two perspectives in recounting a story; and the other one is the *Modality*, with the conditions of video (i.e., *Witnessed*) and audio (i.e., *Reported*) as the information sources. This design allowed for comparison of the linguistic features of true and false statements, and their interactions with the source of information (firsthand vs. secondhand).

Taken together, because in evidential languages there are specific grammatical markers to denote the source of the information, a compliance in the use of evidential markers according to the modality is expected. That is, in evidential languages:

(H1) Direct evidential markers will be used in the *Witnessed_Truth* condition.

(H2) Indirect evidential markers will be employed in the *Reported_Truth* condition.

Conversely, in lie conditions, participants can manipulate the evidentials in two alternative directions, such that using direct evidentials to sound more certain and credible, or indirect evidentials to put a distance between the statement and themselves to remove the responsibility—in the face of causing mismatches between source and evidential markers (i.e., using indirect evidentials in witnessed condition and direct evidentials in reported condition):

(H3a) Direct evidential markers will be used in the *Witnessed_Lie* and *Reported_Lie* conditions.

(H3b) Indirect evidential markers will be used in the *Witnessed_Lie* and *Reported_Lie* conditions.

Additionally, metalinguistic awareness is expected to decrease when conceptualizing a lie due to the high cognitive efforts needed. Hence, participants will not be able to track their narratives, leading to unmotivated switches in both the evidential markers and the tenses when lying (Porter and ten Brinke 2010):

(H4) Grammar hoppings (i.e., switches in both the evidential markers and the tenses) will be observed in *Witnessed_Lie* and *Reported_Lie* conditions for evidential languages.

(H5) Tense hoppings (i.e., switches in the used tenses) will be observed in *Wit-*

nessed_Lie and *Reported_Lie* conditions for non-evidential languages.

Similar to the expectation that Turkish and Japanese speakers will support their lies by direct observation markers, certainty markers (i.e., perception verbs, affirmative adverbs, and plausibility modals of *must have* or *devoir*) as the firsthand evidential tools will be tested for whether they will be resorted to more frequently in English and French when lying than when telling the truth. On the other hand, considering previous findings on the use of specific deceptive cues in Western languages, the frequencies of negation words and first-person pronouns were investigated across conditions for non-evidential languages. Hypotheses for the non-evidential languages are as follows:

(H6) Certainty markers will be utilized more frequently in the *Witnessed_Lie* and *Reported_Lie* conditions.

(H7) The prevalence of negations will be higher in the *Witnessed_Lie* and *Reported_Lie* conditions.

(H8) First-person pronouns will be utilized more in the *Witnessed_Lie* and *Reported_Lie* conditions.

2. MATERIALS AND METHODS

2.1 Participants

Participants were recruited from four language groups: English, French, Turkish, and Japanese. The sample size in each language group was greater than the minimum required ($n = 37$), which was calculated based on a power analysis for linear multiple regression analysis with two categorical variables (2×2 ; Veracity: Truth and Lie, Modality: Witnessed and Reported), $\alpha = .05$, power = .80, and a medium effect size $f^2 = 0.47$ (reliability assignment to evidentials in Turkish, from Karaaslan et al. 2018). Predetermined exclusion criteria were (I) having a non-native knowledge of the required first languages (English, French, Turkish, or Japanese), (II) having a low education status (i.e., lower than secondary education level), (III) having a low self-rated proficiency in their mother tongue (e.g., due to extensive stay abroad). Two English, five French, and two Japanese speakers were removed from the analysis due to not complying with the exclusion criteria (i.e., one Japanese participant being bilingual in English and living abroad, three learning the language at a later age, and five identifying being not native and learning the language at the ages between 0-5). All the recruited participants from the Turkish group were kept for the analysis.

Including three pilot studies, a total of 120 participants were recruited for the Turkish study. After each pilot data collection, refinements in the experimental design and study instructions were implemented for all languages. The pilot studies included 23, 17, and 29 participants, in the respective order. A total of 51 Turkish participants were included in the final analysis. The sample sizes for the other languages after the exclusions were as follows: 63 English speakers, 53 French speakers, and 51 Japanese speakers (see Table 2.1). Except for English and about half of the Japanese data ($n = 22$), where data collection was through crowdsourcing websites like Prolific (www.prolific.com; Palan and Schitter 2018), data from all languages were collected using convenience sampling (Japanese) or in return for course credit

Table 2.1 General characteristics of participant groups by languages

	English	French	Turkish	Japanese
Grammatical evidentiality	No	No	Yes	Yes
<i>N</i>	63	53	51	51
Mean age	38.98 (10.32)	23.61 (4.00)	22.84 (2.00)	34.84 (6.24)
<i>n</i> Female	32	39	39	36
Second language proficiency	2.48 (0.90)	2.79 (0.77)		2.61 (1.00)
How frequently do you lie in daily life?	2.26 (0.54)	2.77 (0.99)	2.33 (0.68)	2.80 (0.63)
How frequently do people lie in your culture?	3.24 (0.59)	3.49 (0.54)	3.61 (0.64)	3.12 (0.52)

Note. Results represent mean scores with *SD* scores in parentheses. Lying frequencies were rated on a 5-point Likert scale (*1: Never, 5: Always*). One English speaker who was included in the analysis did not provide their demographic information.

(Turkish and French), and data collection has stopped in accordance with the conducted a priori power analysis. The average age of participants was highest for the English (age range = 18.32-55.68) and second for the Japanese (age range = 19.32-45.53) groups, which were ~10-15 years older than the French (age range = 18.54-35.80) and Turkish (age range = 18.74-28.45) groups. Gender distribution was comparable across the groups, with the majority of the participants identifying as female, 51% of the English, 74% of the French, 76% of the Turkish, and 71% of the Japanese samples (2 Japanese speakers preferred not to answer the gender question). Most participants reported being right-handed: 56 in the English group, 49 in the French group, 50 in the Turkish group, and 47 in the Japanese group. Among French participants, 1, and among Japanese participants, 3 reported using both hands equally.

All the participants who were included in the analyses reported a native proficiency in the target languages. Of these, 3 participants from the English group and 3 Japanese speakers reported learning the target language between the ages of 0 to 5. Eight of the English participants reported being bi- or multilingual, while this number was 10 for French, none for Turkish, and 11 for the Japanese sample. The majority of the participants in each language indicated at least a limited knowledge (*1: "I cannot speak a second language", 4: "I'm bilingual or multilingual"*; see Appendix E) in additional languages except 10 people from the English, 2 from the French, and 8 from the Japanese group. All the Turkish participants indicated a knowledge of at least an additional language, yet they could not receive the proficiency rating question due to a presentation error during the experiment.

Each participants held at least a high school degree: while on a 6-point Likert scale (*1: Primary school, 6: Ph.D.*; see Appendix E) the mean education level for the

English group was 3.84 ($SD = 0.89$), 2.77 ($SD = 1.01$) for French, 2.78 ($SD = 1.05$) for Turkish, and 4.22 ($SD = 1.14$) for Japanese participants. Moreover, most of the French ($n = 38$) and Turkish participants ($n = 45$) were continuing their education, whereas among the English and Japanese speakers, only 4 people indicated being a student in each.

In order to rule out cultural differences in lying, participants were also asked to self-report their perceptions of lying behavior. On a 5-point Likert scale assessing how frequently they lie in daily life, all the language groups rated similarly, with French and Japanese speakers reporting a slightly higher mean score than the other groups. When asked about the frequency with which people lie in their culture, again, the average ratings were close, yet were highest for the Turkish and lowest for the Japanese speakers. The mean frequencies across languages suggest a balance between participants' self-reported lying behavior (with a slight moderation) and their perceptions of how common lying is within their cultural context. Note that although the Japanese participants received 6-point Likert scales for the lying behavior questions, two middle ratings were merged to allow for a better comparison across the groups. For similar reasons, although the French group received a 5-point Likert scale for the question asking about the education level, score points were adjusted according to the other languages.

2.2 Materials

2.2.1 Stories

Twelve brief narratives were developed in each target language using a default/neutral tense to avoid priming participants in the audio condition (please see Appendix C). Each scenario featured one central character performing four distinct transitive actions (e.g., *"Mary is sitting at the kitchen table with a photograph in front of her. She pours juice into a glass and drinks it. She suddenly tears the photo in two and throws the pieces into the bin next to her."*). The narratives were concise, ranging between 29-45 words ($M = 34.92$, $SD = 4.42$) for English, 26-41 words ($M = 33.67$, $SD = 4.36$) for French, 17-30 words ($M = 20.75$, $SD = 3.65$) for Turkish, 57-89 characters ($M = 70.33$, $SD = 10.47$) for Japanese, and were designed to exclude non-intentional or abstract actions.

Table 2.2 Mean (SD) frequencies of the norming study

	English	French	Turkish	Japanese
Grammatical evidentiality	No	No	Yes	Yes
N	24	45	27	30
Narrative length	34.92 (4.42)	33.67 (4.36)	20.75 (3.65)	70.33 (10.47)
Mean age	31.16 (12.61)	25.68 (9.16)	27.85 (9.33)	36.65 (16.22)
<i>The story sounded natural to me.</i>	5.78 (1.42)	5.54 (1.66)	5.49 (1.70)	4.74 (2.00)
<i>The story was easy to understand.</i>	6.55 (0.83)	6.10 (1.31)	6.00 (1.32)	5.58 (1.82)
<i>I can remember and retell the story without difficulty.</i>	6.23 (1.16)	6.31 (0.98)	6.08 (1.32)	5.16 (1.85)

Note. Norming questions were rated on a 7-point Likert scale (1: *Not at all*; 7: *Very much*). Narrative length represents the frequencies for character count of the stories in Japanese and word count of the stories in other languages.

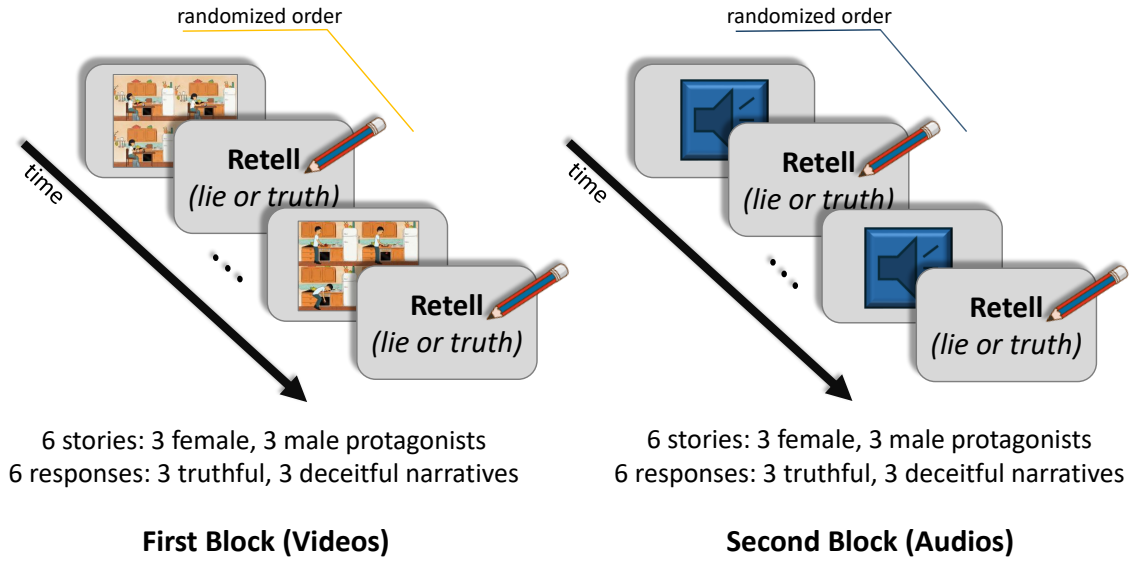
2.2.2 Auditory Stimuli

The event narratives were recorded as high-quality audio files by female native speakers (a male native speaker recorded the audios in French), who consistently voiced all stories in a clear and neutral tone. Each recording lasted approximately 15-20 seconds. While the gender of the protagonists varied across stories for balance, the same speaker was used throughout to maintain uniformity across the auditory materials.

2.2.3 Visual Stimuli

For the visual modality, each of the twelve stories was adapted into short, colored animated clips. These silent animations clearly depicted the four distinct actions performed by a central character in each story (e.g., pouring juice, tearing a photograph). The duration of the videos ranged from 13 to 16 seconds, corresponding to the structure and pacing of the auditory versions (see Figure 2.2).

Figure 2.1 Experimental conditions



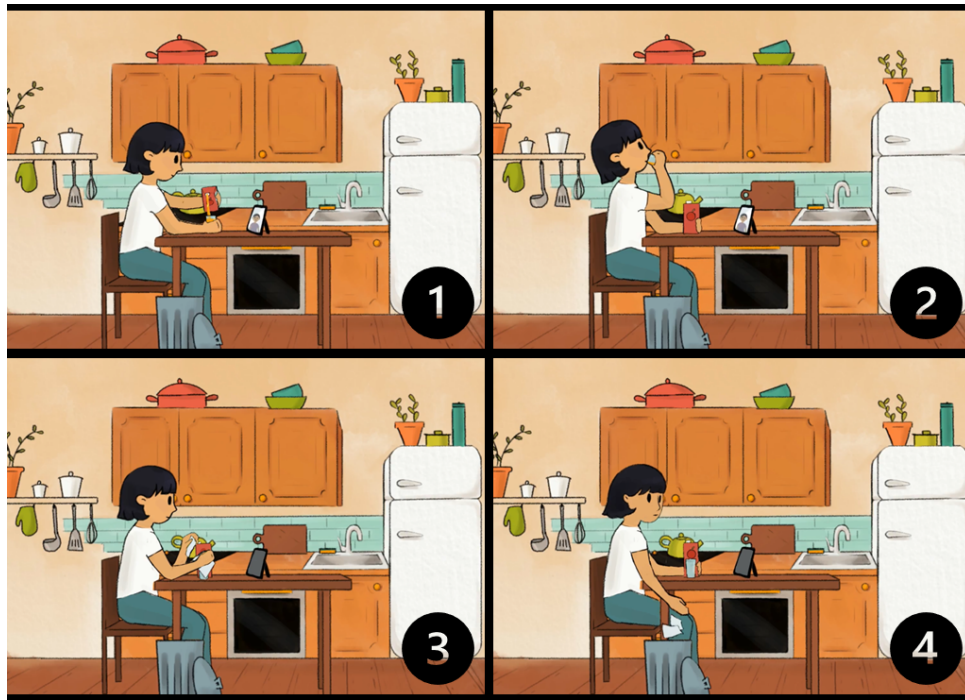
2.3 Procedure

2.3.1 Experimental Procedure

A cross-modal narrative production task was employed, requiring participants to recount short events either honestly or deceptively under four experimental conditions: (i) *Witnessed_Truth*, (ii) *Witnessed_Lie*, (iii) *Reported_Truth*, and (iv) *Reported_Lie*. A fully crossed, counterbalanced design was employed to ensure that each of the twelve narratives appeared in all experimental conditions with a randomized order across the participants. Participants were randomly assigned to one of two experimental groups, where each story appeared in a different combination of Modality (video-based or audio-based input) and Veracity (truthful or deceptive) across the groups. That is, if a specific story was encountered as a truthful witnessed event (video) in one group, it was presented as a deceptive reported event (audio) in the other group, and vice versa. In the witnessed conditions, participants viewed silent video animations for a random half of the stories; whereas in the reported conditions, they listened to an auditory narration of the other half of the stories without visual access (see Appendix F). This design allowed the modality of the information source (directly observed vs. indirectly reported) along with the veracity of the account (truthful vs. deceptive) to be systematically varied. The task was delivered via the Gorilla.sc experimental platform (Anwyl-Irvine et al. 2020).

The experimental task was structured in two consecutive blocks with a fixed block order to avoid priming participants in the video conditions with the linguistic struc-

Figure 2.2 Example arrays from the visually animated video clips for story 1



ture used in audio recordings (see Figure 2.1). Therefore, in the first block, half of the stories were shown as silent animations. In the second block, participants listened to audio recordings of the remaining six stories, now presented without visual support. Of these, three were randomly assigned to be recounted truthfully, and the remaining three required participants to fabricate a deceptive account in each block. Truth and lie assignments, as well as the presentation modality of the stories, were reversed for half of the participants to maintain balance across conditions. To control for potential gender effects, the protagonists were evenly split by gender (six male, six female) and systematically distributed across truth and lie trials. The gender assignments were cross-balanced to avoid confounding with either the modality or veracity conditions.

In the *Witnessed_Truth* and *Reported_Truth* conditions, once the participants watched/listened to the story, they were given the following instructions: “*You have just watched a clip showing/ listened to a recording about Mary. Now imagine that you are talking to Mary’s partner, shown in the photo. Her partner has no idea how serious the disagreement between them is for Mary. Therefore, you want Mary’s partner to know what Mary did today. In the text box below, tell the story shown to you in the video/ presented to you in the recording as accurately as possible.*” (see Appendix B). The participants were prompted to be as accurate as possible; the outputs were therefore elicited as ‘*truthful*’ retellings. In the *Witnessed_Lie* and *Reported_Lie* conditions, the participants were instructed in the

following way: “*You have just watched a clip showing Mary./ You have just listened to a recording about Mary. Now imagine that you are talking to Mary’s partner, shown in the photo. Her partner found out that Mary tore up the photo today because of their previous disagreement. In order to make things better between them, you must convince her partner that Mary did not tear up the photo. Tell a deceptive version of the story you just watched/listened to in the text box below by changing at least three of Mary’s actions shown in the video/ presented in the recording.*” The two lie conditions necessitated that participants deliberately deceive by changing at least three actions, rather than other random details, to allow for the observation of potential switches in verb forms, such as tense and evidentiality. Since the stories elicited low-stakes motivations for lying, participants were particularly instructed to convince their audience that their pre-held knowledge was in fact not true in lie conditions. In other words, they were expected to lie in *a denial scenario* to provide a stronger motivation for lying. To allow a better comparison, participants were also prompted with a contextual motivation grounded in the instructions appropriate for each mini-story in truth conditions. There was no time limit imposed on the participants to conceptualize and write their retellings.

2.3.2 Coding of the Retellings and Analysis

Participants’ written responses were recorded as raw text. Two native French-speaking and two native Turkish-speaking independent researchers manually coded the retellings in English. The same French-speaking researchers automated the count scores of the total number of words, verbs, embeddings, perception verbs, first-person pronouns, other pronouns, and negations in Microsoft Excel by creating functions, and coded the rest of the French data. The Turkish speaker researchers coded the Turkish data manually, except that the total word counts were computed in Excel. The Japanese data was automated altogether, again using Excel functions, and a native Japanese speaker coded 15% of the data for reliability purposes. Retellings in each language group was coded considering the study hypotheses and their unique linguistic features for most or some of the following categories (see Table 3.1 for all frequencies): (i) General characteristics of retellings including the total number of words, verbs, embeddings, adverbs (doubt adverbs, affirmation adverbs, *suddenly*), pronouns (first-person pronouns, other-person pronouns, pronoun hoppings), negations, and (ii) Tense/Evidentiality-specific outcomes including: the number of present and past tense forms (all tense forms that make present [simple present, present continuous, present perfect continuous] or past time reference [simple past, past perfect, past continuous, past perfect continuous, present perfect]), number of

tense hopping instances (i.e., unmotivated/non-pragmatic switches between tense markers), number of direct and indirect evidential forms, number of evidentiality hopping instances (i.e., unmotivated switches between evidential forms), perception verbs (i.e., “*I saw*”, “*J’ai vu*”). Similar to the direct evidential markers, as certainty epistemic modals (i.e., plausibility modals) are perceived as more believable and taken more seriously—even in reported speech—(Jahiu 2022), their frequencies in non-evidential languages were also coded (i.e., *must have* for English and *devoir* for French). All variables were count scores, except tense/evidential/grammar hopping and *must have* scores were based on a binary categorization (i.e., no switch vs. at least one switch). Due to scope reasons, only the frequencies of the grammatical categories that are related to the hypotheses will be discussed in this paper.

The inter-rater reliability was calculated over 15% of the data in all languages. The percent agreement between the coders was found to be .88 for English, .97 for Turkish, and .99 for Japanese. Inter-rater reliability could not be assessed for the French dataset, as a second coder fluent in French was not available during the coding process. In the cases where the automated codes and the raters did not agree, human raters were prioritized. When the codes of the first and the second human raters were not consistent, corrections were applied until agreement or codes of the original rater were used in disagreements. Responses with no text input due to technical issues (i.e., audio/video not played) were removed from the analysis. The removed data were 0.26% for the English, 1.89% for the French (i.e., the whole data of a participant), 2.6% for the Turkish, and 0.82% for the Japanese study.

3. RESULTS

3.1 Data Analytics Strategy

A 2×2 within-participants design was adopted in testing the hypotheses. Factor variables were *Veracity* (*Truth* $\times$ *Lie*) and the presentation *Modality* (*Witnessed* $\times$ *Reported*); the outcome variables were the frequency occurrences of the selected grammatical markers. Since each participant recounted events in each trial across conditions as a repeated measure, trials ($N = 2580$) were nested within participants ($N = 217$); hence, the data had a multilevel structure. Therefore, a hierarchical linear modelling was employed in examining frequency occurrences across conditions to successfully consider within- and between-subject variability as well as the variability across items simultaneously without losing any participants due to missing data points and avoiding a Type 1 error (Brown 2021; Peugh 2010).

The data were analyzed using generalized linear mixed-effects regression models in Jamovi version 2.6.26 (The jamovi project 2025) with the GAMLj3 module, which is based on the lme4 (Bates et al. 2015) and nlme packages (Pinheiro and Bates 2023) in R software. The within-participant Level 1 predictors Veracity (Truth vs. Lie) and presentation Modality (Witnessed vs. Reported) were treated as fixed effects, and the frequency occurrences of grammatical markers as the outcome variables were at Level 1. Because Level 1 predictors were categorical variables held constant throughout, their random slopes were not entered into the model. For categorical variables, treatment coding was applied, and they were dummy coded with *Truth* and *Reported* conditions being at the reference level (Veracity: Truth = 0, Lie = 1; Modality: Reported = 0, Witnessed = 1).

To determine the variance due to participants and stories, intercept-only models were run. The between-participant variability accounted for a range of 15.90-100% (M of ICCs = 51.60%, $ps < .05$) of the data variance in all dependent variables (except affirmation adverbs in English, tense hopping and plausibility modals in

French), whereas between-story variability did not explain an additional variance in the analyses (M of the significant 7 models $ICCs \approx 9.93$). Moreover, individual lying behavior was entered into the models, allowing for a random slope to control for the potential effects on the outcome, yet the explained variance by the slopes did not reach the significance level. Therefore, individuals were kept as the cluster variable, allowing for a random intercept, while the stories and the lying behavior were removed from the following analyses. Further post-hoc between-condition comparisons were computed using Bonferroni correction, which obtained the exact same results as the computations without any correction.

All the models were run separately for each language and each dependent variable combination, considering the hypotheses; however, the analytical approach was the same across analyses. Binary categorical dependent variables (i.e., tense hopping, evidential hopping, grammar hopping, *must have* as plausibility modal in English) were tested using the Binomial family with a logistic link function. All other variables that had non-binary count scores were analyzed with the Poisson family and a logarithmic link function. The dependent variables that were not explained by the random effects of the participants (i.e., between-participant variability was not significant) were analyzed in generalized linear models (i.e., affirmation adverbs in English, tense hopping and plausibility modals in French). Please see Appendix G for the formulas illustrating sample generalized multilevel linear regression models that were applied.

For the evidential languages, due to the low number of observations in both hopping types of tenses and evidentials, the models could not perform successfully. Therefore, these two observations were merged together as a composite *grammar hopping* score with binary categories (no observation = 0, observation = 1). The data and all the analysis scripts are available at the Open Science Framework (OSF) repository: https://osf.io/uc6pr/?view_only=f362463a363d44a0a4a2a832f10ad835

3.2 Narrative Characteristics

To begin with, the distribution of the narrative characteristics across conditions and languages was examined (see Table 3.1 for frequency results). Before conducting any mixed-model analyses, the distributions of the dependent variables were checked separately for each language to determine whether the assumption for normality of residuals was violated. The descriptive statistics and Kolmogorov-Smirnov tests ($ps < .05$) revealed that all of the variables were highly positively skewed. Based

on the distribution type (Binomial vs. Poisson), the best-fitting models were designated.

3.3 Examination of the Deceptive Cues Across Conditions

In this section, the suggested grammatical markers of deception are examined across four languages: English, French, Turkish, and Japanese. Although separate models were run for each language, the results are presented by grouping the languages into two categories—evidential and non-evidential—to allow for a more structured comparison. The findings for the two languages within each category are reported together to highlight potential contrasts between language types. Before presenting the results, here the grammatical markers relevant to the main hypotheses will be briefly restated to clarify the order of presentation: For evidential languages, the analysis focuses on the frequencies of direct and indirect evidential markers and the presence of grammar-hopping instances (i.e., the composite scores for tense and evidentiality hoppings). For non-evidential languages, the analysis addresses the presence of tense hopping instances and the frequency of certainty markers (i.e., perception verbs, plausibility modals, and affirmation adverbs), negation words, and first-person pronouns. Even though previous literature points out that the total word and verb counts can be indicators of deception, they were not included in the present analysis. The rationale behind this decision is that, due to the instructions asking to change the action verbs in the lying conditions, participants were expected to produce longer narratives with an increased number of verbs.

3.3.1 Evidential Languages

For evidential languages, the distribution of direct and indirect evidential marker rates and grammar hopping observations was tested across conditions. Table 3.2 demonstrates mixed effects predictions across evidential languages.

3.3.1.1 Direct evidential markers

Following the hypotheses, the first analysis was on whether the direct evidential rates differed across conditions of evidential languages. A generalized mixed model analysis for repeated measures indicated that both of the manipulations held signif-

Table 3.1 Mean (SD) frequencies of the narrative characteristics by conditions and languages

	Witnessed Truth	Witnessed Lie	Reported Truth	Reported Lie
English				
Words	31.08 (13.16)	45.55 (21.31)	31.22 (12.46)	45.59 (20.10)
Verbs	6.29 (2.93)	8.63 (4.17)	5.75 (2.71)	8.96 (4.08)
Past tense	4.38 (2.16)	6.47 (3.33)	3.73 (2.13)	5.99 (3.11)
Other tense	0.41 (1.32)	0.36 (0.91)	0.87 (1.59)	0.71 (1.67)
Tense hopping (n)	10	6	9	14
Perception verbs	0.13 (0.41)	0.04 (0.23)	0.05 (0.25)	0.05 (0.29)
Plausibility modals (n)	0	11	4	3
Affirmation adverbs	0.02 (0.13)	0.01 (0.07)	0.01 (0.07)	0.02 (0.13)
Doubt adverbs	0.03 (0.16)	0.03 (0.16)	0.01 (0.10)	0.01 (0.10)
First pronoun	0.17 (0.57)	0.14 (0.61)	0.11 (0.39)	0.10 (0.39)
Other pronoun	1.48 (1.24)	2.39 (1.88)	1.72 (1.30)	2.52 (1.69)
Negation	0.21 (0.51)	0.45 (0.79)	0.16 (0.43)	0.41 (0.73)
French				
Words	49.16 (26.97)	58.22 (34.31)	30.49 (12.26)	41.00 (22.80)
Verbs	9.17 (4.61)	11.10 (5.95)	5.93 (2.51)	7.95 (4.43)
Finite verbs	5.74 (2.88)	7.26 (4.09)	4.24 (1.75)	5.20 (2.88)
Non-finite verbs	3.43 (2.50)	3.83 (2.54)	1.69 (1.55)	2.75 (2.20)
Past tense	4.68 (2.29)	6.10 (3.49)	3.41 (2.02)	4.09 (2.64)
Other tense	1.06 (2.01)	1.16 (1.87)	0.83 (1.58)	1.12 (1.81)
Tense hopping (n)	13	12	14	12
Perception verbs	0.07 (0.26)	0.04 (0.24)	0.03 (0.18)	0.03 (0.18)
Plausibility modals	0.06 (0.27)	0.12 (0.35)	0.03 (0.21)	0.06 (0.25)
Affirmation adverbs	0.06 (0.23)	0.20 (0.43)	0.03 (0.16)	0.13 (0.39)
Doubt adverbs	0.04 (0.24)	0.05 (0.22)	0.01 (0.08)	0.01 (0.08)
First pronoun	0.27 (0.83)	0.32 (1.05)	0.10 (0.39)	0.30 (0.97)
Other pronoun	3.93 (3.10)	4.58 (3.34)	2.64 (1.65)	3.26 (2.46)
Pronoun hopping (n)	0	1	0	1
Negation	0.22 (0.50)	0.43 (0.76)	0.08 (0.30)	0.22 (0.48)
Turkish				
Words	21.14 (9.48)	32.73 (16.81)	18.88 (7.53)	28.93 (12.51)
Verbs	5.15 (2.26)	7.97 (3.82)	4.77 (1.88)	7.40 (3.40)
Finite verbs	3.43 (1.54)	4.75 (2.31)	3.07 (1.28)	4.26 (1.89)
Non-finite verbs	1.80 (1.45)	3.15 (2.23)	1.57 (1.28)	3.08 (2.35)
Past tense	2.74 (1.96)	4.21 (2.59)	1.27 (1.65)	1.89 (2.36)
Present tense	0.61 (1.28)	0.59 (1.68)	1.97 (2.07)	2.41 (2.49)
Tense hopping (n)	5	5	3	7
Direct evidential	2.70 (1.97)	3.92 (2.62)	1.31 (1.65)	1.60 (2.22)
Indirect evidential	0.01 (0.08)	0.30 (1.23)	0 (0)	0.26 (1.16)
Evidential hopping (n)	0	9	0	4
Grammar hopping (n)	5	14	3	10
Japanese				
Characters	67.69 (34.49)	94.33 (47.82)	63.27 (27.52)	82.802 (36.42)
Verbs	7.69 (3.64)	9.74 (4.69)	6.71 (3.06)	9.41 (4.12)
Past tense	3.68 (2.10)	5.38 (3.15)	3.40 (2.25)	4.97 (2.79)
Other tense	2.11 (2.30)	2.96 (2.90)	1.22 (1.60)	2.43 (2.54)
Tense hopping (n)	0	0	6	2
Direct evidential	3.68 (2.10)	5.37 (3.16)	3.38 (2.26)	4.96 (2.80)
Indirect evidential	0.17 (0.44)	0.16 (0.44)	0.12 (0.42)	0.19 (0.50)
Evidential hopping (n)	2	1	0	0
Grammar hopping	2	1	6	2
Plausibility modals (n)	0	1	3	0
First pronoun	0.07 (0.30)	0.12 (0.39)	0.03 (0.21)	0.14 (0.46)
Other pronoun	0.33 (0.64)	0.48 (0.89)	0.36 (0.71)	0.47 (0.88)

Note. Scores in the tables for plausibility modals in French and Japanese samples; tense hopping, evidential hopping, grammar hopping, and pronoun hopping instances in all groups represent the total number of trials that included them.

ificant results for the Turkish sample (Veracity: $\beta = 0.35$, Incidence Rate Ratio, $IRR = 1.41$, $p < .001$, 95% CI [1.25, 1.61]; Modality: $\beta = -0.72$, $IRR = 0.49$, $p < .001$, 95% CI [0.41, 0.57]), without an interaction effect ($\beta = -0.14$, $IRR = 0.87$, $p = .23$, 95% CI [0.70, 1.09]). Complying with the main hypotheses, the Lie condition ($M = 2.05$, $SE = 0.22$) was more likely to bring higher rates of direct evidentials than the Truth condition ($M = 1.55$, $SE = 0.17$). However, the Witnessed condition was associated with a lower likelihood of having direct evidential markers ($M = 1.20$, $SE = 0.13$) than the Reported condition ($M = 2.65$, $SE = 0.27$). The model fit analysis suggested that 18% of the estimation for the direct evidentials was coming from the two predictors of Veracity and Modality. The model fit was high with a total estimation of 65% when both the fixed and random effects were included (Conditional R^2).

On the other hand, only the Veracity ($\beta = 0.38$, $IRR = 1.46$, $p < .001$, 95% CI [1.31, 1.62]) was significant in estimating the direct evidential rates for the Japanese sample (Modality: $\beta = -0.09$, $IRR = 0.91$, $p = .14$, 95% CI [0.81, 1.03]; interaction effect: $\beta = 0.01$, $IRR = 1.01$, $p = .91$, 95% CI [0.86, 1.18]), where producing evidential markers was less probable while telling the truth ($M = 3.28$, $SE = 0.20$) than narrating a deceptive story ($M = 4.804$, $SE = 0.29$). In this model, when the fixed effect predictors of Veracity and Modality were entered, they explained an additional 9.8% variance in the direct evidentials (Marginal R^2). This portion was increased to 47.2% with the addition of random effects (Conditional R^2).

3.3.1.2 Indirect evidential markers

Next, the distribution of indirect evidential rates was explored in the Turkish dataset. The analysis indicated that Veracity ($\beta = 3.84$, $IRR = 46.7$, $p < .001$, 95% CI [6.41, 339]) had an effect on the retellings, yet Modality ($\beta = -13.80$, $IRR = 1.01e-6$, $p = .79$, 95% CI [2.93e-51, 3.49e+38]) and the interaction effect were not significant ($\beta = 13.63$, $IRR = 827532$, $p = .80$, 95% CI [2.40e-39, 2.86e+50]). However, since the confidence interval was unconventionally large for the estimation of Veracity, a post-hoc analysis was conducted. The significance of Veracity on the indirect evidential markers could not survive after the post-hoc comparisons (Ratio = 2.36e-5, $SE = 7.92e-4$, $p = .75$). The model fit analyses revealed a towering results such that the explained variances was 63% for the fixed effects (Marginal R^2), and 100% with the addition of random effects (Conditional R^2). These results indicated poor model estimation due to the variance entirely stemming from the participants ($ICC = 1.00$).

Similar to the Turkish sample, indirect evidential rates did not differ across conditions for the Japanese sample (Veracity: $\beta = -0.04$, $IRR = 0.96$, $p = .89$, 95% CI [0.56, 1.67]; Modality: $\beta = -0.35$, $IRR = 0.71$, $p = .26$, 95% CI [0.39, 1.29]; interaction: $\beta = 0.50$, $IRR = 1.65$, $p = .23$, 95% CI [0.74, 3.68]). The model estimation for the total explained variances was 35.7% (Conditional R^2), while only 0.9% of the variance was due to the fixed effect factors of Veracity and Modality (Marginal R^2).

3.3.1.3 Grammar hoppings

As the final hypothesis regarding the evidential languages was on the grammar-hopping incidences, their presence across conditions was examined. Only Veracity ($\beta = 1.23$, $IRR = 3.43$, $p = .03$, 95% CI [1.16, 10.13]) was effective on the distribution of the grammar hopping instances in the retellings of the Turkish group, whereas neither the estimation of Modality ($\beta = -0.52$, $IRR = 0.59$, $p = .49$, 95% CI [0.14, 2.59]), nor the interaction effect was not significant ($\beta = 0.07$, $IRR = 1.07$, $p = .94$, 95% CI [0.19, 5.96]). Thus, as expected, the presence of the switches in the evidential forms or the tenses was anticipated with a greater likelihood in lying ($M = 0.05$, $SE = 0.02$) than telling the truth ($M = 0.02$, $SE = 0.01$). The model fit analysis results indicated that inclusion of Veracity and Modality manipulations accounted for an additional 9.2% of the grammar hopping instances (Marginal R^2), whereas the model explained a total of 33.9% variance, including both fixed and random effects (Conditional R^2). For the Japanese group however, the model did not estimate any significant effects (Veracity: $\beta = -0.73$, $IRR = 0.48$, $p = .56$, 95% CI [0.04, 5.63]; Modality: $\beta = 1.25$, $IRR = 3.49$, $p = .15$, 95% CI [0.64, 18.96]; interaction: $\beta = -0.52$, $IRR = 0.60$, $p = .73$, 95% CI [0.03, 11.72]). When the Veracity and Modality were included as fixed effects, the model predicted a change of 6.2% in the explained variances (Marginal R^2). The total model estimation was 59.1% with the addition of random effects (Conditional R^2).

Table 3.2 Statistical outputs from mixed-effects models in evidential languages

	Turkish				Japanese			
	β	SE	z	p	β	SE	z	p
Direct evidential count								
Intercept	0.80	0.11	7.34	< .001	1.23	0.07	17.98	< .001
Veracity	0.35	0.06	5.36	< .001	0.38	0.05		< .001
Modality	-0.72	0.09	-8.40	< .001	-0.09	0.06	-1.48	.14
Veracity*Modality	-0.14	0.12	-1.19	.23	0.01	0.08	0.11	.91
Indirect evidential count								
Intercept	-10.83	1.82	-5.94	< .001	-2.30	0.28	-8.25	< .001
Veracity	3.84	1.01	3.80	< .001	-0.04	0.28	-0.14	.89
Modality	-13.80	52.32	-0.26	.79	-0.35	0.31	-1.13	.26
Veracity*Modality	13.63	52.32	0.26	.80	0.50	0.41	1.21	.23
Presence of grammar hopping								
Intercept	-3.93	0.57	-6.87	< .001	-5.89	1.63	-3.62	< .001
Veracity	1.23	0.55	2.23	.03	-0.73	1.25	-0.58	.56
Modality	-0.52	0.75	-0.70	.49	1.25	0.86	1.45	.15
Veracity*Modality	0.07	0.88	0.08	.94	-0.52	1.52	-0.34	.73

3.3.2 Non-Evidential Languages

For the non-evidential languages, tense hopping instances, perception verbs of “I saw” and “J’ai vu”, plausibility modals of “must have” and “devoir”, affirmation adverbs, negation words, and first-person pronouns were examined in the respective order. Table 4.3 presents the regression results across non-evidential languages.

3.3.2.1 Tense hoppings

When the distribution of tense hopping instances across conditions was examined, the model demonstrated that, in contrast to the expectation, the probability of switching the tenses was similar across conditions of both Veracity ($\beta = -0.53$, $IRR = 0.59$, $p = .33$, 95% CI [0.20, 1.72]) and Modality ($\beta = -0.12$, $IRR = 0.89$, $p = .81$, 95% CI [0.34, 2.35]) for the English sample with an insignificant interaction effect ($\beta = 1.06$, $IRR = 2.88$, $p = .14$, 95% CI [0.70, 11.82]). The model fit analysis results revealed that the fixed effects predictors of Veracity and Modality explained a further 2.2% variance in the tense hopping instances (Marginal R^2), while the

total portion of the variances in the model was 34.3% including the random effects (Conditional R^2). Similar results were viable for the French sample, indicating no differences across conditions (Veracity: $\beta = -0.09$, $IRR = 0.92$, $p = .84$, 95% CI [0.40, 2.08]; Modality: $\beta = 0.08$, $IRR = 1.08$, $p = .84$, 95% CI [0.49, 2.39]; interaction: $\beta = -0.08$, $IRR = 0.92$, $p = .89$, 95% CI [0.29, 2.91]). The model fit was considerably low in predicting variances of the tense hopping instances for the French group, explained 6.59e-4 marginally.

3.3.2.2 Perception verbs

After the examination of the tense hopping instances, certainty markers were tested in separate models. The first model estimated whether the probability of the perception verb observations was different across conditions. For the English group, model indicated that the experimental conditions were effective (Veracity: $\beta = -1.14$, $IRR = 0.32$, $p = .005$, 95% CI [0.14, 0.71]; Modality: $\beta = -0.91$, $IRR = 0.40$, $p = .02$, 95% CI [0.19, 0.84]), except for the interaction effect ($\beta = 1.13$, $IRR = 3.09$, $p = .06$, 95% CI [0.95, 10.09]). The results indicated that the likelihood of occurrence of perception verb rates was significantly higher in truth ($M = 0.010$, $SE = 0.009$) and reported conditions ($M = 0.009$, $SE = 0.008$) than in lie ($M = 0.006$, $SE = 0.005$) and witnessed conditions ($M = 0.007$, $SE = 0.005$). Inclusion of the two fixed effects predictors accounted for 2.4% of the variance in the model (Marginal R^2). The model estimated a total of 74.1% variance in the perception verb counts, combining both fixed and random effects (Conditional R^2).

Unlike the English sample, the model failed to estimate any differences in perception verb rates between the levels of the conditions for the the French sample (Veracity: $\beta = -0.45$, $IRR = 0.64$, $p = .35$, 95% CI [0.25, 1.64]; Modality: $\beta = -0.79$, $IRR = 0.46$, $p = .14$, 95% CI [0.16, 1.31]; interaction: $\beta = 0.45$, $IRR = 1.57$, $p = .57$, 95% CI [0.33, 7.48]). When the fixed effects of Veracity and Modality were included in the model, only 0.4% of the explained variance was accounted for (Marginal R^2). The model estimated a total of 100% of the variance in perception counts with the addition of random effects (Conditional R^2).

3.3.2.3 Plausibility modals

Next, the likelihood of plausibility modal occurrences was examined. The results for plausibility modal use comparisons between the two conditions of the model

predictors were not significant for both the English-speaking participants (Veracity: $\beta = 18.50$, $IRR = 1.11e+8$, $p = .76$, 95% CI [3.04e-45, 4.07e+60]; Modality: $\beta = 17.30$, $IRR = 3.33e+7$, $p = .78$, 95% CI [9.10e-46, 1.22e+60]; interaction: $\beta = -18.8$, $IRR = 6.53e-9$, $p = .76$, 95% CI [1.78e-61, 2.40e+44]) and for the French speakers (Veracity: $\beta = 0.64$, $IRR = 1.90$, $p = .10$, 95% CI [0.88, 4.09]; Modality: $\beta = -0.69$, $IRR = 0.50$, $p = .21$, 95% CI [0.17, 1.46]; interaction: $\beta = 0.05$, $IRR = 1.05$, $p = .94$, 95% CI [0.28, 3.94]). The analysis of the model fit indicated that the Veracity and Modality manipulations as fixed effects predicted a change of 88.5% in the explained variances in plausibility modals (Marginal R^2), while these portions increased to 95% with the addition of the random effects (Conditional R^2) for the analyses of the English group. The model estimated 3.75% of the variance for French.

3.3.2.4 Affirmation adverbs

Similarly, when the affirmation adverb rates were tested across conditions, for the English sample, the results indicated that there were no differences across conditions (Veracity: $\beta = -1.09$, $IRR = 0.34$, $p = .34$, 95% CI [0.03, 3.22]; Modality: $\beta = -1.09$, $IRR = 0.34$, $p = .34$, 95% CI [0.03, 3.22]; interaction: $\beta = 2.19$, $IRR = 8.91$, $p = .18$, 95% CI [0.36, 218.59]). In predicting the affirmation adverbs, the model accounted for a 2.85% estimation in the explained variances.

For the French group, however, there was a significant main effect of Veracity on affirmation adverb rates, where the lie condition ($M = 0.13$, $SE = 0.03$) yielded a higher likelihood of affirmation adverb rates than the truth ($M = 0.03$, $SE = 0.01$) condition (Veracity: $\beta = 1.24$, $IRR = 3.44$, $p = .001$, 95% CI [1.64, 7.23]). Again, both the estimation of Modality and the interaction effect did not reach a significant level (Modality: $\beta = -0.81$, $IRR = 0.44$, $p = .18$, 95% CI [0.14, 1.44]; interaction: $\beta = 0.42$, $IRR = 1.52$, $p = .53$, 95% CI [0.41, 5.60]). In total, the model estimated 33% of the variance in affirmation adverb counts for the French sample (Conditional R^2), whereas this percentage was 17.8% for the variance due to fixed effects only (Marginal R^2).

3.3.2.5 Negation words

As to the negation word rate, analyses revealed significant effects of conditions (Veracity: $\beta = 0.76$, $IRR = 2.13$, $p < .001$, 95% CI [2.13, 2.14]; Modality: $\beta = -$

0.25, $IRR = 0.78$, $p < .001$, 95% CI [0.78, 0.78]) and an interaction effect ($\beta = 0.16$, $IRR = 1.18$, $p < .001$, 95% CI [1.17, 1.18]) for the English sample. These results indicated that more negation words were produced in the Lie ($M = 0.36$, $SE = 7.03e-4$) condition than in the Truth condition ($M = 0.15$, $SE = 2.15e-4$). Again, participants used negation words when reporting a story ($M = 0.26$, $SE = 3.56e-4$) than when they recounted after watching them in a video ($M = 0.22$, $SE = 4.24e-4$). However, because the model performance was not optimal and the convergence was not guaranteed, the interaction effect did not hold true in the simple effect comparisons. The model fit results indicated that 9.2% of the model estimation in the variance of the negation words was due to the fixed effects (Marginal R^2). When the model included the random effects in addition to the fixed effects, this model explained a total of 28.8% of the variance (Conditional R^2). Since the results might be biased for this model, these findings will not be discussed in detail in the following sections.

For the French sample, both Veracity and Modality affected the negation word rates (Veracity: $\beta = 0.65$, $IRR = 1.91$, $p = .002$, 95% CI [1.27, 2.88]; Modality: $\beta = -0.99$, $IRR = 0.37$, $p = .002$, 95% CI [0.20, 0.70]), with no observation of an interaction effect ($\beta = 0.34$, $IRR = 1.41$, $p = .38$, 95% CI [0.66, 3.00]). These results indicated that French participants used negation words more likely when lying ($M = 0.27$, $SE = 0.04$) compared to when they were telling the truth ($M = 0.12$, $SE = 0.02$), and when they retold the stories after listening ($M = 0.27$, $SE = 0.04$) than watching ($M = 0.12$, $SE = 0.02$). In this model, fixed effect predictors, Veracity and Modality, accounted for a marginal estimation of 15% in the explained variances (Marginal R^2). The total estimation of the model was increased to 28.5% when the random effects were added (Conditional R^2).

3.3.2.6 First-person pronoun

Finally, the distribution of first-person pronouns was tested across conditions. For English, no significant results were obtained (Veracity: $\beta = -0.20$, $IRR = 0.82$, $p = .44$, 95% CI [0.49, 1.36]; Modality: $\beta = -0.45$, $IRR = 0.64$, $p = .11$, 95% CI [0.37, 1.11]; interaction: $\beta = 0.04$, $IRR = 1.04$, $p = .92$, 95% CI [0.46, 2.34]). The fixed effect predictors could only explain 0.9% of the variance when the model estimated first-person pronouns (Marginal R^2), whereas the portion of the random effects in the explained variances was large, resulting in a total estimation of 68.9% variance (Conditional R^2).

Akin to the English speakers, Veracity did not account for a significant effect for

the French group (Veracity: $\beta = 0.17$, $IRR = 1.19$, $p = .41$, 95% CI [0.79, 1.79]). However, both the Modality and the interaction effect successfully predicted the model (Modality: $\beta = -1.03$, $IRR = 0.36$, $p < .001$, 95% CI [0.20, 0.64]; interaction: $\beta = 0.97$, $IRR = 2.63$, $p = .008$, 95% CI [1.29, 5.36]), revealing that participants produced first-person pronouns in the Reported condition ($M = 0.03$, $SE = 0.02$) with a greater likelihood than the Witnessed condition ($M = 0.02$, $SE = 0.01$). To pinpoint the exact effect of the interaction, a series of follow-up tests was conducted. The simple effect of Veracity in estimating the first person pronoun count was not found significant at the Reported condition ($\beta = 0.17$, $IRR = 1.19$, $p = .41$, 95% CI [0.79, 1.79]); and significant in the Witnessed condition ($\beta = 1.14$, $IRR = 0.30$, $p < .001$, 95% CI [1.75, 5.60]), where the lying manipulation led participants to use more first person pronouns in their retellings more frequently ($M = 0.03$, $SE = 0.02$) than the truth condition ($M = 0.01$, $SE = 0.01$). Moreover, the simple effect of Modality was significant when the condition was Truth ($\beta = -1.03$, $IRR = 0.38$, $p < .001$, 95% CI [0.20, 0.64]; yet not at the level of Lie ($\beta = -0.06$, $IRR = 0.94$, $p = .76$, 95% CI [0.63, 1.40]), meaning that first person pronoun use was more likely for truthful recounting of a story from an audio ($M = 0.03$, $SE = 0.02$) than a video ($M = 0.01$, $SE = 0.01$). This model estimated 85.5% of the variance in first-person pronoun counts, including both fixed and random effects (Conditional R^2). The margin of the fixed effects, Veracity and Modality, in the model estimation was 3.2% of the explained variance (Marginal R^2).

Table 3.3 Statistical outputs from mixed-effects models in non-evidential languages

	English				French			
	β	SE	z	p	β	SE	z	p
Presence of tense hopping								
Intercept	-3.57	0.45	-7.86	< .001	-2.40	0.29	-8.28	< .001
Veracity	-0.53	0.55	-0.97	.33	-0.09	0.42	-0.21	.84
Modality	-0.12	0.50	-0.24	.81	0.08	0.40	0.20	.84
Veracity*Modality	1.06	0.72	1.47	.14	-0.08	0.59	-0.14	.89
Perception verb count								
Intercept	-4.12	0.84	-4.93	< .001	-7.76	1.72	-4.51	< .001
Veracity	-1.14	0.41	-2.80	.005	-0.45	0.48	-0.94	.35
Modality	-0.91	0.37	-2.42	.02	-0.79	0.54	-1.46	.14
Veracity*Modality	1.13	0.60	1.87	.06	0.45	0.80	0.57	.57
Plausibility modals								
Intercept	-22.8	61.8	-0.37	.71	-2.75	0.32	-8.69	< .001
Veracity	18.5	61.8	0.30	.76	0.64	0.39	1.64	.10
Modality	17.3	61.8	0.28	.78	-0.69	0.55	-1.27	.21
Veracity*Modality	-18.8	61.8	-0.31	.76	0.05	0.67	0.08	.94
Affirmation adverb count								
Intercept	-4.14	0.58	-7.18	< .001	-3.11	0.37	-8.40	< .001
Veracity	-1.09	1.16	-0.95	.34	1.24	0.38	3.27	.001
Modality	-1.09	1.16	-0.95	.34	-0.81	0.60	-1.35	.18
Veracity*Modality	2.19	1.63	1.34	.18	-0.42	0.66	0.64	.53
Negation word count								
Intercept	-1.74	0.001	-1399	< .001	-1.65	0.20	-8.42	< .001
Veracity	0.76	0.001	608	< .001	0.65	0.21	3.11	.002
Modality	-0.25	0.001	-202	< .001	-0.99	0.33	-3.05	.002
Veracity*Modality	0.16	0.001	130	< .001	0.34	0.398	0.88	0.38
First-person pronoun count								
Intercept	-3.48	0.49	-7.09	< .001	-3.67	0.61	-6.02	< .001
Veracity	-0.20	0.26	-0.77	.44	0.17	0.21	0.83	.41
Modality	-0.45	0.28	-1.60	.11	-1.03	0.30	-3.43	< .001
Veracity*Modality	0.04	0.41	0.10	.92	0.97	0.36	2.67	.008

Note. For plausibility modal analyses, a logistic regression was conducted for English and a logarithmic regression for French. The results for affirmation adverbs in English, tense hoppings and plausibility modals in French are the outputs from generalized linear model analyses; the outputs for the rest of the variables are from generalized mixed-model analyses.

4. DISCUSSION

The present study sought to determine whether speakers of languages with and without an obligatory category of grammatical evidential markers employ different grammatical strategies to manipulate their verbal narratives when lying. In addition, previously reported patterns regarding the increased use of specific grammatical markers in deceptive speech—compared to truthful speech—were tested across languages. To this end, the presentation modality and the narrative type were manipulated, where stories were accessed in audio or video format, then the participants were asked to recount them either deceptively or exactly as it was. Confirming the main hypothesis, participants used direct evidential forms (i.e., *-DI* in Turkish, and no marker in Japanese) more frequently during deceptive narratives compared to their truthful retellings, regardless of the language spoken. However, contrary to the predictions, the direct evidential markers were preferred to a lesser extent when the event was witnessed compared to reported for the Turkish group, and the indirect evidential markers were produced similarly across conditions in both evidential languages. Overall, the evidence regarding the distribution of other grammatical structures was not uniform across languages, indicating the adoption of different strategies beyond a single cross-linguistic pattern.

The findings for the increased use of direct evidentials in deception may stem from the participants’ motivation to possess more authority over their deceptive assertions to sound more certain and reliable (Aksu-Koç 2016; Arslan 2020; Bergqvist and Grzech 2023; Ishida 2006). For the purpose of achieving a successful implantation of the deceitful information to the receiver, one needs to present their lie in the most convincing manner possible; and because direct evidential markers are considered more reliable from childhood onward (Aydin and Ceci 2013; Matsui, Yamamoto, and McCagg 2006; Ozturk and Papafragou 2016), people can actively manipulate the use of evidential markers in their statements for pragmatic reasons, especially when the face value is crucial, such as in diplomatic contexts (Xu 2022). The exploitation of direct evidential markers in lying conditions would be spontaneous thereof. This result is especially informative in the Japanese context, as direct evidential form is

commonly avoided to claim less assertion and express politeness in speech (Filipović, Brown, and Engelhardt 2023; Matsumura 2017; Ohta 1991). Yet, the results for the Turkish sample may merely reflect the direct evidential markers being the default grammar type for past tense uses in Turkish (Johanson 2008). Also, the increased use of direct evidential markers in the reported condition compared to the witnessed condition indicates a violation of the grammatical rules of Turkish. This result was unforeseen and needs further examination of the tradeoff between direct and indirect evidential uses in the witnessed events. Notably, though, the manipulation of reported speech as a secondhand access to the information may not be successful in this design, causing the predictions regarding the contrasts between the modality conditions to fail. Since the participants heard the stories directly from a narrator in an experimental setting, they might have considered them firsthand information and a reliable source, as if like listening to a news report.

On the other hand, the finding of no difference between conditions for the indirect evidential use could be due to several reasons. Contrary to Turkish, indirect evidentials can be considered as the default use in Japanese daily speech, and they may have been used in other pragmatic functions. In this respect, although *rashii* and *soo da* both represent inference in speech, their meaning can be pragmatically manipulated in the discursive context. While *rashii* is preferred when the speaker wants to stand aloof from the information presented, *yoo da* is employed to convey information euphemistically (Karlsson 2013). Therefore, the first function could have resulted in an increased use of indirect evidentials in lie conditions, whereas the latter can explain the overall pervasive use across conditions for the conventional function of politeness (Filipović, Brown, and Engelhardt 2023). For the Turkish group, however, the indirect evidential marker *-mİş* is the grammatical standard in reported speech; thus, its rate was expected to be higher in the reported condition than the witnessed condition. However, because the brief stories in the experiment included small misdemeanors that needed to be vindicated in both truth and lie, participants might have resorted to indirect evidentials to shift the responsibility of their accounts to another source, inflating the rates in all conditions.

Again, contrary to the hypothesis, tense and grammar hopping observations were indistinguishable across conditions for all language groups except Turkish. The results indicated that the presence of the grammar-hopping incidences was more likely to occur when lying versus telling the truth among the Turkish speakers. The direction of this effect in Turkish was consistent with the previous evidence (Porter and ten Brinke 2010). Here, the assumption was that due to the high cognitive demands of lying (Adams-Quackenbush 2015; Gombos 2006), when the condition was to deceive, participants would be more prone to the frequent switches in tense and evidential

markers in their statements. While the grammar hopping was a rare instance across the narratives, it is conceivable to observe that it did not reach a significance level. However, the failure in this assumption may be explained by the employed experimental design: Since the participants were asked to provide written statements, they had enough resources to track and edit their assertions, whereas Porter and ten Brinke (2010) exemplify tense hoppings in oral narratives with high-stakes deception. Therefore, the expected cognitive effort put into deceptive narratives might have been alleviated in this sample.

Given that certainty markers of perception verbs, plausibility modals (Desclés 2018; Jahiu 2022), and affirmation adverbs can serve to denote direct evidentiality in speech (Tantucci 2016), they were also expected to be exploited when lying. When certainty markers were analyzed, however, different patterns were observed. Contrary to previous research (Hancock et al. 2007), perception verbs were associated with true and reported retellings than the opposite conditions for the English sample; yet, they were not different between the conditions for the French sample. Likewise, plausibility modals showed no significant difference between the experimental conditions, contrary to previous findings in French, which indicated an increased use of plausibility modals (i.e., *devoir*) in truthful statements (Chiu et al. 2025). However, although there was no significant effect of conditions in the English group, affirmation adverbs were more prevalent in lies than in true narratives of the French speakers. Together, these findings suggest that English and French speakers utilize certainty markers with different epistemic functions. While French-speaking participants were similar to the speakers of evidential languages in terms of their manipulation of the lexical forms that pertain to a more reliable image, English participants exhibited a reverse strategy. One possible reason behind this result is that English speakers might have used such strategies in truth conditions more frequently, as they are less sensitive to evidential marking in assigning reliability to their statements than evidential language speakers (Mushin 2013; Tosun, Vaid, and Geraci 2013). Here, the model estimation was low for perception verbs in English; hence, their findings should be interpreted cautiously. Both being non-evidential languages, however, this disparity emphasizes the importance of cross-linguistic examinations.

On the other hand, replicating the previous findings, the lying condition was associated with increased use of negation words with a stronger link than the truth condition (Hauch et al. 2015; Vrij 2008) and reported than the witnessed condition for both participant groups. Inflated negation use in deceptive speech was discussed in terms of the negative attitude arising from the defensive tone in a denial of wrongdoing (Hauch et al. 2015). However, it should be noted that the findings for the

English sample may be biased due to non-optimal model fit.

Finally, the distribution of the first-person pronouns was similar across the conditions for the English sample, while French participants used fewer first-person pronouns after watching a video compared to listening to a story from an audio recording. Furthermore, when French speakers narrated a witnessed event, they were more likely to use first-person pronouns in lying as opposed to telling the truth. Also, between the true narratives, first-person pronouns were more prevalent in reported speech than in the narratives of direct access. These findings did not comply with the previous literature that suggested a decreased use of first-person pronouns in deceptive speech (Solà-Sales et al. 2023; Vrij 2008). As previous research stressed, employed methods and the context can be influential on the use of first-person pronouns such that deceptive statements may contain an increased number of pronouns (both first-person and other person) than true statements (Holtgraves and Jenkins 2020). Therefore, this incompatibility might be explained by the experimental design adopted in this study. The contextual information given for the narratives did not directly address the participants, where their involvement was low. Meaning that, they were rather observers in the events, which might have caused a decline in overall first-person pronoun use in their narratives.

4.1 Limitations and Future Directions

Although some of the findings replicated previous studies, there are some limitations imposed on this study. Here, it should be noted that linguistic observations can heavily depend on the contextual differences; hence, the findings should be considered within the scope and boundaries of this study. Furthermore, as the attained effect sizes were relatively small in some analysis models, the results should be interpreted cautiously.

First and foremost, the employed experimental design may not allow the findings to be applied to a broader category of spontaneous lying situations. For example, because this study was conducted on an online platform, the relative control over the participants was low during the experiment. While the design restricted the participants' responses to only written production, their liberty to read and review their narratives might have led the participants to change their initial statements. Hence, the expected effects of the manipulations on the production of the specified grammar types, especially the grammar/tense hopping instances, might have vanished after being revised or corrected for grammatical rules. For this reason,

future designs should employ controlled experimental procedures allowing for oral production of the deceptive statements to ensure that the cognitive demands of lying behavior would be reflected in the narratives.

In a similar vein, while this study put its participants into a roleplay scenario, where they were assumed to be the acquaintances of the protagonists from the given stories, this assumption would naturally not be as strong as real-life experiences, and the urge to lie may not be as comparable to a real-life event. Even though participants were motivated to lie by the provision of some contextual information, this study elicited low-stakes motivations for lying; thereby, expected contrasts between truth and deception conditions in the linguistic structures might have been moderated. As previous research highlighted, the motivation types, especially those that are related to self, were associated with diverse linguistic observations, stronger than in low-stakes scenarios (DePaulo et al. 2003). Thereby, to support this study’s findings and apply them in courtrooms, forensic science should examine the evidential uses in eyewitness testimonies and defendant statements from countries speaking evidential languages. Similarly, future research should consider high-stakes lying scenarios to make their participants more motivated to craft their lies more carefully and use intricate linguistic forms.

In addition to that, given that exposure to and practising a second language for an extended period of time is influential in evidential marking in speech (Arslan, De Kok, and Bastiaanse 2017; Tokaç-Scheffer, Nickels, and Arslan 2024), this study can be extended to heritage language speakers. To ascertain whether the disparity in the employment of direct evidential uses in deception across non-evidential and evidential languages is due to practice effects (i.e., evidential uses would be altered) or differences in established cognitive perspectives (i.e., transfer effects in a second language would remain), future research should investigate evidentiality in heritage languages.

Another interesting approach would be to compare the adoption of evidential uses by children who speak a language with or without grammatical evidential markers in their attempt to deceive others. Although the proper use of evidentials is a challenging concept for young children (Ozturk and Papafragou 2008), the utilization of evidentials was observed among preschoolers speaking evidential languages (Kandemirci et al. 2023), who also make their reliability assessments based on the evidential markers in speech, differently from and earlier than their non-evidential speaker peers (Aydin 2011; Aydin and Ceci 2013). Additionally, evidential marking was discussed to facilitate the online reasoning of Theory of Mind (ToM; Özoran 2009) and false belief understanding through source monitoring skills among Turkish-speaking children (Kandemirci et al. 2023), while the development of ToM

and lying was also found to be related (see Lee and Imuta 2021, and Sai et al. 2021 for meta-analysis on lying and ToM). Hence, it is plausible to think that evidential uses can be exploited even by young children for telling lies. This research line needs further examination to establish the differences in the developmental trajectory of manipulating evidentials in lying across languages.

Finally, to ensure that the results regarding the frequencies of the grammatical markers do not simply reflect an inflated narrative length in this study, future research should use normalized scores by computing ratios relative to the number of other grammatical categories when necessary. Here, because the distributions of grammatical uses were highly skewed, and there were an increased number of meaningful zeros in the data, the best fitting models were determined as non-normal distributions (Mullahy and Norton 2024), which only permitted the use of integers as the input data instead of ratios. However, although this study was properly powered, future studies should encourage participants to produce longer narratives with an elevated number of grammatical structures to obtain normal distributions and provide better insights.

4.2 Conclusion

All in all, the literature highlights that there are multiple dimensions to consider while evaluating the veracity of a statement, ranging from context, prosodic and cognitive implications, interpersonal dynamics (e.g., motivations, audience, social cues) to modality (i.e., oral vs. written) and the means of communication and language. However, studies on lie detection rarely go beyond certain linguistic categories in certain languages and look into rather *niche* grammatical markers. For this reason, this study examined four target languages in their diverse and unique approaches to evidential marking in deception. The present study obtained confirming results regarding the main hypothesis, such that speakers of evidential languages used an increased amount of direct evidentials when lying than when telling the truth, whereas non-evidential speakers resorted to diverse strategies. Evidential speakers were discussed to manipulate their narratives to become more reliable by asserting first-hand knowledge and certainty through the constant practice and emerging need for source marking. However, the exact reason behind this observation was not pinpointed within the scope of this research, leaving it an open question. This study was, to my knowledge, the first to investigate the intentional manipulation of evidential uses in lying. Moreover, considering the differences in available pragmatic

tools across languages, this study employed a well-rounded approach by investigating four distinct languages and a range of linguistic structures, elaborating on the prior knowledge in language use in deception conditions. Considering the increased rate of attention in lie detection research and the demand for access to truthful content online in recent years, this study has particular importance in applied science for the current needs of an information age.

BIBLIOGRAPHY

- Adams-Quackenbush, Nicole M. 2015. "The effects of cognitive load and lying types on deception cues."
- Aikhenvald, Aleksandra Y. 2002. *Language contact in Amazonia*. Oxford University Press.
- Aikhenvald, Alexandra Y. 2004. *Evidentiality*. OUP Oxford.
- Akkol, Ekin, and Yılmaz Gökşen. 2024. "Creating a Comprehensive Data Set for Deception Detection Studies in Turkish Texts." *Research Journal of Business and Management* 11(2): 138–145.
- Aksu-Koç, Ayhan. 1988. *The acquisition of aspect and modality*. Cambridge University Press.
- Aksu-Koç, Ayhan. 2016. "The interface of evidentials and epistemics in Turkish: Perspectives from acquisition." In *Exploring the Turkish Linguistic Landscape*. John Benjamins Publishing Company pp. 143–156.
- Aksu-Koç, Ayhan, Hale Ögel-Balaban, and I Ercan Alp. 2009. "Evidentials and source knowledge in Turkish." *New Directions for Child and Adolescent Development* 2009(125): 13–28.
- Almela, Ángela. 2021. "A corpus-based study of linguistic deception in Spanish." *Applied Sciences* 11(19): 8817.
- AnderBois, Scott. 2014. On the exceptional status of reportative evidentials. In *Semantics and Linguistic Theory*. pp. 234–254.
- Anolli, Luigi, Michela Balconi, and Rita Ciceri. 2003. "Linguistic styles in deceptive communication: Dubitative ambiguity and elliptic eluding in packaged lies." *Social Behavior and Personality: An international journal* 31(7): 687–710.
- Antomo, Mailin. 2025. "Lying with Gestures." *Advances in Experimental Philosophy of Lying* p. 169.
- Anwyl-Irvine, Alexander L, Jessica Massonnié, Adam Flitton, Natasha Kirkham, and Jo K Evershed. 2020. "Gorilla in our midst: An online behavioral experiment builder." *Behavior Research Methods* 52(1): 388–407.
- Arciuli, Joanne, David Mallard, and Gina Villar. 2010. "'Um, I can tell you're lying': Linguistic markers of deception versus truth-telling in speech." *Applied Psycholinguistics* 31(3): 397–411.
- Arslan, Seçkin. 2020. "When the owner of information is unsure: Epistemic uncertainty influences evidentiality processing in Turkish." *Lingua* 247: 102989.

- Arslan, Seçkin, Doerte De Kok, and Roelien Bastiaanse. 2017. “Processing grammatical evidentiality and time reference in Turkish heritage and monolingual speakers.” *Bilingualism: Language and Cognition* 20(3): 457–472.
- Arslan, Seçkin, Elif Tutku Tunalı, Yağmur Çetin, and Özgür Aydın. 2024. “Eyes do not lie but words do: Evidence from eye-movement monitoring during reading that misuse of evidentiality marking in Turkish is interpreted as deceptive.” *Functions of Language* 31(1): 90–108.
- Aydin, Cagla. 2011. “Language And Suggestibility: Cross-Linguistic Evidence On Children’S Assessment Of Source Knowledge.”
- Aydin, Cagla, and Stephen J Ceci. 2013. “The role of culture and language in avoiding misinformation: Pilot findings.” *Behavioral Sciences & the Law* 31(5): 559–573.
- Bachenko, Joan, Eileen Fitzpatrick, and Michael Schonwetter. 2008. Verification and implementation of language-based deception indicators in civil and criminal narratives. In *Proceedings of the 22nd International Conference on Computational Linguistics (COLING 2008)*. pp. 41–48.
- Baer, Carolyn, Antonia Frederike Langenhoff, Dilara Keşşafoglu, Winuss Mohtezebade, Celeste Kidd, Aylin C Küntay, Jan Engelmann, and Bahar Köymen. 2025. “Anticipating disagreement enhances source memory in English-and Turkish-speaking preschool children.” *Developmental Psychology* .
- Bates, Douglas, Martin Mächler, Ben Bolker, and Steve Walker. 2015. “Fitting linear mixed-effects models using lme4.” *Journal of Statistical Software* 67: 1–48.
- Bergqvist, Henrik, and Karolina Grzech. 2023. “The role of pragmatics in the definition of evidentiality.” *STUF-Language Typology and Universals* 76(1): 1–30.
- Boncea, Irina Janina. 2013. “Hedging patterns used as mitigation and politeness strategies.” *Annals of the University of Craiova* 2(1): 7.
- Boye, Kasper. 2012. *Epistemic meaning: A crosslinguistic and functional-cognitive study*. Vol. 43 Walter de Gruyter.
- Boye, Kasper, and Peter Harder. 2009. “Evidentiality: Linguistic categories and grammaticalization.” *Functions of Language* 16(1): 9–43.
- Brown, Violet A. 2021. “An introduction to linear mixed-effects modeling in R.” *Advances in Methods and Practices in Psychological Science* 4(1): 2515245920960351.
- Carretero, Marta, Juana I Marín-Arrese, and Anna Ruskan. 2022. “Epistemicity and stance in English and other European languages: Discourse-pragmatic perspectives.”
- Chiu, Ming Ming, Alex Morakhovski, Zhan Wang, and Jeong-Nam Kim. 2025. “Detecting Covid-19 Fake News on Twitter/X in French: Deceptive Writing Strategies.” *Media and Communication* 13.

- Conway III, Lucian Gideon, Felix Thoemmes, Amy M Allison, Kirsten Hands Towgood, Michael J Wagner, Kathleen Davey, Amanda Salcido, Amanda N Stovall, Daniel P Dodds, Kate Bongard et al. 2008. "Two ways to be complex and why they matter: Implications for attitude strength and lying." *Journal of Personality and Social Psychology* 95(5): 1029.
- Cornillie, Bert. 2007. "The continuum between lexical and grammatical evidentiality: a functional analysis of Spanish parecer." *Rivista di linguistica* 19(1): 109–128.
- Danielewicz-Betz, A, and N Ogasawara. 2013. "Truth statements in denial context: A study of prosodic, paralinguistic, and discursive cues in Japanese speakers." *Journal of Forensic Research* 11.
- De Haan, Ferdinand. 1999. "Evidentiality and epistemic modality: Setting boundaries." *Southwest Journal of Linguistics* 18(1): 83–101.
- DePaulo, Bella M, James J Lindsay, Brian E Malone, Laura Muhlenbruck, Kelly Charlton, and Harris Cooper. 2003. "Cues to deception." *Psychological Bulletin* 129(1): 74.
- Desclés, Jean-Pierre. 2018. "Epistemic modality and evidentiality from an enunciative perspective." *Guentchéva, Z., Epistemic modalities in cross-linguistic perspective, and evidentiality* pp. 383–401.
- DiMaggio, Paul. 1997. "Culture and cognition." *Annual Review of Sociology* 23(1): 263–287.
- Dor, Daniel. 2017. "The role of the lie in the evolution of human language." *Language Sciences* 63: 44–59.
- Ekici, Sami, Turgut Kavas, Yaman Akbulut, and Abdulkadir Sengur. 2017. Deception detection from speech signals. In *International Conference on Advances and Innovations in Engineering (ICAIE)*.
- Elbatanouny, Hagar, Noora Al Roken, Abir Hussain, Wasiq Khan, Bilal Khan, and Eqab Almajali. 2025. "A comprehensive analysis of deception detection techniques leveraging machine learning." *Expert Systems with Applications* 283: 127601.
- Eskin, Maria Raluca. 2024. Analysis of Speech Content and Voice for Deceit Detection. Master's thesis Bilkent Universitesi (Turkey).
- Fekete, Gabriella. 2019. "information-processing in french adults Practicing Deception." *Psychology in Russia: State of the art* 12(1): 30–47.
- Filipović, Luna, Mika Brown, and Paul E Engelhardt. 2023. "Evidential meanings in native and learner Japanese and English: Implications for the assessment of speaker certainty." *Pragmatics and Society* 14(3): 484–508.
- García-Carpintero, Manuel. 2023. "Lying versus misleading, with language and pictures: The adverbial account." *Linguistics and Philosophy* 46(3): 509–532.

- Gödert, Heinz Werner, Hans-Georg Rill, and Gerhard Vossel. 2001. "Psychophysiological differentiation of deception: the effects of electrodermal lability and mode of responding on skin conductance and heart rate." *International Journal of Psychophysiology* 40(1): 61–75.
- Gombos, Victor A. 2006. "The cognition of deception: The role of executive processes in producing lies." *Genetic, Social, and General Psychology Monographs* 132(3): 197–214.
- Goupil, Louise, Emmanuel Ponsot, Daniel Richardson, Gabriel Reyes, and Jean-Julien Aucouturier. 2021. "Listeners' perceptions of the certainty and honesty of a speaker are associated with a common prosodic signature." *Nature Communications* 12(1): 861.
- Gröndahl, Tommi, and N Asokan. 2019. "Text analysis in adversarial settings: Does deception leave a stylistic trace?" *ACM Computing Surveys (CSUR)* 52(3): 1–36.
- Hancock, Jeffrey T, Lauren E Curry, Saurabh Goorha, and Michael Woodworth. 2007. "On lying and being lied to: A linguistic analysis of deception in computer-mediated communication." *Discourse Processes* 45(1): 1–23.
- Hauch, Valerie, Iris Blandón-Gitlin, Jaume Masip, and Siegfried L Sporer. 2015. "Are computers effective lie detectors? A meta-analysis of linguistic cues to deception." *Personality and Social Psychology Review* 19(4): 307–342.
- Henrich, Joseph, Steven J Heine, and Ara Norenzayan. 2010. "The weirdest people in the world?" *Behavioral and Brain Sciences* 33(2-3): 61–83.
- Holtgraves, Thomas, and Elizabeth Jenkins. 2020. "Texting and the language of everyday deception." *Discourse Processes* 57(7): 535–550.
- Hoye, Leo Francis. 2008. "Evidentiality in discourse: A pragmatic and empirical account." *Pragmatics and corpus linguistics: A mutualistic entente* .
- Hwang, Hyisung C, David Matsumoto, and Vincent Sandoval. 2016. "Linguistic cues of deception across multiple language groups in a mock crime context." *Journal of Investigative Psychology and Offender Profiling* 13(1): 56–69.
- Ishida, Kazutoh. 2006. "How can you be so certain? The use of hearsay evidentials by English-speaking learners of Japanese." *Journal of Pragmatics* 38(8): 1281–1304.
- Jahiu, Edona. 2022. "The Evidentiality through reported speech: Pragma-linguistic factors affecting its reliability." *Studies in Pragmatics and Discourse Analysis* 3(1): 33–50.
- jamovi project, The. 2025. "jamovi (Version 2.6) [Computer Software].".
- Johanson, Lars. 2008. "Evidentiality in Turkic." In *Studies in evidentiality*. John Benjamins Publishing Company pp. 273–290.
- Johnson, Marcia K, Shahin Hashtroudi, and D Stephen Lindsay. 1993. "Source monitoring." *Psychological Bulletin* 114(1): 3.

- Kandemirci, Birsu, Anna Theakston, Ditte Boeg Thomsen, and Silke Brandt. 2023. "Does evidentiality support source monitoring and false belief understanding? A cross-linguistic study with Turkish-and English-speaking children." *Child Development* 94(4): 889–904.
- Karaaslan, Hatice, Annette Hohenberger, Hilmi Demir, Simon Hall, and Mike Oaksford. 2018. "Cross-cultural differences in informal argumentation: Norms, inductive biases and evidentiality." *Journal of Cognition and Culture* 18(3-4): 358–389.
- Karlsson, Simon. 2013. "Yooda & Rashii: Pragmatic motivations and the perception of information."
- Ketterer, Holly R, and Michael Hoerger. 2014. *Lying in daily life*. Vol. 2 Sage Publications.
- Laing, Brent Logan. 2015. *The language and cross-cultural perceptions of deception*. Brigham Young University.
- Lakshan, Imasha, Lumini Wickramasinghe, Sandaru Disala, Smirithika Chandrasegar, and Prasanna S Haddela. 2019. Real time deception detection for criminal investigation. In *2019 National information technology conference (NITC)*. IEEE pp. 90–96.
- Lazard, Gilbert. 2001. "On the grammaticalization of evidentiality." *Journal of Pragmatics* 33(3): 359–367.
- Le, Marina Tam. 2016. Language of low-stakes and high-stakes deception: Differences within individuals PhD thesis University of British Columbia.
- Lee, Jia Ying Sarah, and Kana Imuta. 2021. "Lying and theory of mind: A meta-analysis." *Child Development* 92(2): 536–553.
- Levitan, Sarah Ita, Angel Maredia, and Julia Hirschberg. 2018. Linguistic cues to deception and perceived deception in interview dialogues. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*. pp. 1941–1950.
- Markowitz, David M, and Jeffrey T Hancock. 2014. "Linguistic traces of a scientific fraud: The case of Diederik Stapel." *PloS One* 9(8): e105937.
- Matsui, Tomoko, Taeko Yamamoto, and Peter McCagg. 2006. "On the role of language in children's early understanding of others as epistemic beings." *Cognitive Development* 21(2): 158–173.
- Matsumoto, David, Hyisung C Hwang, and Vincent A Sandoval. 2015a. "Cross-language applicability of linguistic features associated with veracity and deception." *Journal of Police and Criminal Psychology* 30(4): 229–241.
- Matsumoto, David, Hyisung C Hwang, and Vincent A Sandoval. 2015b. "Ethnic similarities and differences in linguistic indicators of veracity and lying in a moderately high stakes scenario." *Journal of Police and Criminal Psychology* 30(1): 15–26.

- Matsumura, Tomomi. 2017. “The use of evidentials in hearsay contexts in Japanese and English.”
- Matthewson, Lisa, Henry Davis, and Hotze Rullmann. 2007. “Evidentials as epistemic modals: Evidence from St’át’imcets.” *Linguistic Variation Yearbook* 7(1): 201–254.
- Mbaziira, A, and J Jones. 2016. A text-based deception detection model for cyber-crime. In *Int. Conf. Technol. Manag.* pp. 1–8.
- McGlone, Matthew, and Max Baryshevtsev. 2018. “Lying and Quotation.”
- Meibauer, Jörg. 2005. “Lying and falsely implicating.” *Journal of Pragmatics* 37(9): 1373–1399.
- Meibauer, Jörg. 2014. *Lying at the semantics-pragmatics interface*. Vol. 14 Walter de Gruyter GmbH & Co KG.
- Meibauer, Jörg. 2016. “Understanding bald-faced lies: an experimental approach.” *International Review of Pragmatics* 8(2): 247–270.
- Meibauer, Jörg. 2018. “The linguistics of lying.” *Annual Review of Linguistics* 4(1): 357–375.
- Mélac, Eric. 2022. “The grammaticalization of evidentiality in English.” *English Language & Linguistics* 26(2): 331–359.
- Mélac, Eric, and Pascale Leclercq. 2024. “The functions of evidentiality.”
- Merzagora, Anna Caterina, Scott Bunce, Meltem Izzetoglu, and Banu Onaral. 2006. Wavelet analysis for EEG feature extraction in deception detection. In *2006 International Conference of the IEEE Engineering in Medicine and Biology Society*. IEEE pp. 2434–2437.
- Michael, Lev. 2015. “Chapter 5. The cultural bases of linguistic form: The development of Nanti quotative evidentials.” In *Language structure and environment: Social, cultural, and natural factors*. John Benjamins Publishing Company pp. 99–130.
- Mullahy, John, and Edward C Norton. 2024. “Why transform Y? The pitfalls of transformed regressions with a mass at zero.” *Oxford Bulletin of Economics and Statistics* 86(2): 417–447.
- Mushin, Ilana. 2001. “Japanese reportive evidentiality and the pragmatics of retelling.” *Journal of Pragmatics* 33(9): 1361–1390.
- Mushin, Ilana. 2013. “Making knowledge visible in discourse: Implications for the study of linguistic evidentiality.” *Discourse Studies* 15(5): 627–645.
- Newman, Matthew L, James W Pennebaker, Diane S Berry, and Jane M Richards. 2003. “Lying words: Predicting deception from linguistic styles.” *Personality and Social Psychology Bulletin* 29(5): 665–675.

- Nuckolls, Janis B, and Lev Michael. 2014. "Evidentials and evidential strategies in interactional and socio-cultural context." In *Evidentiality in Interaction*. John Benjamins Publishing Company pp. 13–20.
- Ohta, Amy S. 1991. "Evidentiality and politeness in Japanese." *Issues in Applied Linguistics* 2(2).
- Ott, Myle, Yejin Choi, Claire Cardie, and Jeffrey T Hancock. 2011. "Finding deceptive opinion spam by any stretch of the imagination." *arXiv preprint arXiv:1107.4557*.
- Özoran, Dincer. 2009. Cognitive development of Turkish children on the relation of evidentiality and Theory of Mind. Master's thesis Middle East Technical University (Turkey).
- Ozturk, Ozge, and Anna Papafragou. 2008. "The acquisition of evidentiality in Turkish."
- Ozturk, Ozge, and Anna Papafragou. 2016. "The acquisition of evidentiality and source monitoring." *Language Learning and Development* 12(2): 199–230.
- Palan, Stefan, and Christian Schitter. 2018. "Prolific. ac—A subject pool for online experiments." *Journal of Behavioral and experimental finance* 17: 22–27.
- Palmer, Frank Robert. 2001. *Mood and modality*. Cambridge University Press.
- Peugh, James L. 2010. "A practical guide to multilevel modeling." *Journal of School Psychology* 48(1): 85–112.
- Pinheiro, J, and D Bates. 2023. "Core Team (2022). nlme: Linear and nonlinear mixed effects models (R package version 3.1-157).".
- Porter, Stephen, and Leanne ten Brinke. 2010. "The truth about lies: What works in detecting high-stakes deception?" *Legal and Criminological Psychology* 15(1): 57–75.
- Quinn, Naomi, and Dorothy Holland. 1987. "Culture and cognition¹." *Cultural Models in Language and Thought* p. 1.
- Repke, Meredith A, Lucian Gideon Conway III, and Shannon C Houck. 2018. "The strategic manipulation of linguistic complexity: A test of two models of lying." *Journal of Language and Social Psychology* 37(1): 74–92.
- Rogers, Todd, Richard Zeckhauser, Francesca Gino, Michael I Norton, and Maurice E Schweitzer. 2017. "Artful paltering: The risks and rewards of using truthful statements to mislead others." *Journal of Personality and Social Psychology* 112(3): 456.
- Rubin, Victoria L. 2014. "Pragmatic and cultural considerations for deception detection in Asian languages."
- Sai, Liyang, Siyuan Shang, Cleo Tay, Xingchen Liu, Tingwen Sheng, Genyue Fu, Xiao Pan Ding, and Kang Lee. 2021. "Theory of mind, executive function, and lying in children: A meta-analysis." *Developmental Science* 24(5): e13096.

- San Roque, Lila. 2019. “Evidentiality.” *Annual Review of Anthropology* 48(1): 353–370.
- Sarzynska-Wawer, Justyna, Aleksandra Pawlak, Julia Szymanowska, Krzysztof Hanusz, and Aleksander Wawer. 2023. “Truth or lie: Exploring the language of deception.” *Plos One* 18(2): e0281179.
- Simpson, David. 1992. “Lying, liars and language.” *Philosophy and Phenomenological Research* 52(3): 623–639.
- Slobin, Dan I. 2016. “Thinking for speaking and the construction of evidentiality in language contact.” *Exploring the Turkish Linguistic Landscape: Essays in honor of Eser Erguvanlı-Taylan* 175: 105–120.
- Solà-Sales, Sara, Chiara Alzetta, Carmen Moret-Tatay, and Felice Dell’Orletta. 2023. “Analysing deception in witness memory through linguistic styles in spontaneous language.” *Brain Sciences* 13(2): 317.
- Speas, Peggy. 2008. “On the syntax and semantics of evidentials.” *Language and Linguistics Compass* 2(5): 940–965.
- Spence, Katelyn, Gina Villar, and Joanne Arciuli. 2012. “Markers of deception in Italian speech.” *Frontiers in Psychology* 3: 453.
- Sperber, Dan, Fabrice Clément, Christophe Heintz, Olivier Mascaro, Hugo Mercier, Gloria Origgi, and Deirdre Wilson. 2010. “Epistemic vigilance.” *Mind & Language* 25(4): 359–393.
- Stokke, Andreas. 2013. “Lying and asserting.” *The Journal of Philosophy* 110(1): 33–60.
- Tantucci, Vittorio. 2016. “Toward a typology of constative speech acts: Actions beyond evidentiality, epistemic modality, and factuality.” *Intercultural Pragmatics* 13(2): 181–209.
- Taskin, Suleyman Gokhan, Ecir Ugur Kucuksille, and Kamil Topal. 2022. “Detection of Turkish fake news in Twitter with machine learning algorithms.” *Arabian Journal for Science and Engineering* 47(2): 2359–2379.
- Taylor, Paul J, Samuel Lerner, Stacey M Conchie, and Tarek Menacere. 2017. “Culture moderates changes in linguistic self-presentation and detail provision when deceiving others.” *Royal Society Open Science* 4(6): 170128.
- Tokaç-Scheffer, Suzan D, Lyndsey Nickels, and Seçkin Arslan. 2024. “One suitcase, two grammars: What can we conclude about Australian Turkish heritage speakers’ divergent processing of evidentiality?” *Linguistics Vanguard* 10(s2): 125–138.
- Tosun, Sümeýra, Jyotsna Vaid, and Lisa Geraci. 2013. “Does obligatory linguistic marking of source of evidence affect source memory? A Turkish/English investigation.” *Journal of Memory and Language* 69(2): 121–134.
- Van Swol, Lyn M, Michael T Braun, and Deepak Malhotra. 2012. “Evidence for the Pinocchio effect: Linguistic differences between lies, deception by omissions, and truths.” *Discourse Processes* 49(2): 79–106.

- Villar, Gina, and Paola Castillo. 2017. "The Presence of 'Um' as a Marker of Truthfulness in the Speech of TV Personalities." *Psychiatry, Psychology and Law* 24(4): 549–560.
- Villar, Gina, Joanne Arciuli, and David Mallard. 2012. "Use of "um" in the deceptive speech of a convicted murderer." *Applied Psycholinguistics* 33(1): 83–95.
- Villar, Gina, Joanne Arciuli, and Helen Paterson. 2013. "Linguistic indicators of a false confession." *Psychiatry, Psychology and Law* 20(4): 504–518.
- Vogler, Nikolai, and Lisa Pearl. 2020. "Using linguistically defined specific details to detect deception across domains." *Natural Language Engineering* 26(3): 349–373.
- Vrij, Aldert. 2008. *Detecting lies and deceit: Pitfalls and opportunities*. John Wiley & Sons.
- Wang, Di, Chongxiang Wang, Xingyu Yi, Liyang Sai, Genyue Fu, and Xiaohong Allison Lin. 2022. "Detecting concealed information using functional near-infrared spectroscopy (fNIRS) combined with skin conductance, heart rate, and behavioral measures." *Psychophysiology* 59(8): e14029.
- Willett, Thomas. 1988. "A cross-linguistic survey of the grammaticization of evidentiality." *Studies in Language. International Journal sponsored by the Foundation "Foundations of Language"* 12(1): 51–97.
- Xu, Yinqing, Bei Shi, Wentao Tian, and Wai Lam. 2015. A unified model for unsupervised opinion spamming detection incorporating text generality. In *Proceedings of the 24th International Conference on Artificial Intelligence*. pp. 725–731.
- Xu, Zhongyi. 2022. "Pragmatic functions of evidentiality in diplomatic discourse: Toward a new analytical framework." *Frontiers in Psychology* 13: 1019359.
- Yousef, Ghaida M. 2025. "Prosodic Features of Lying: A View From Jordanian Colloquial Arabic." *Theory and Practice in Language Studies* 15(5): 1462–1470.
- Zhou, Yingsu, and Qi Su. 2018. "Exploring distinction between deception and truth in Chinese: an analysis based on linguistic features." *Int. J. Knowl. Lang. Process* 9: 1–16.

APPENDIX A

Consent Form

Tell me lies: A study on linguistics and lying

What is this study about?

We are doing a study to examine the way people lie, focusing primarily on the language structure that they employ when doing so. We created 12 short stories to investigate this idea, which you will be interacting with throughout the experiment.

Participation is voluntary!

You have the right to choose to either partake or not in this experiment. If you participate, you have the right to withdraw from the study at any stage without giving a reason, and this will not have any consequences for you. No financial compensation will be provided for your participation. However, this experiment will help us further research in deception in linguistics, so we hope you choose to participate!

Privacy

The personal data collected from this survey is limited to age, gender, profession, education, native language, and country of residence, with the sole purpose of helping the researchers properly interpret the experiment data. Your data will be kept closed to all access other than the researcher, and no individual-level analysis will be performed or published. Don't worry, your responses will remain anonymous, and your personal data will not be traceable!

Questions?

Any questions or concerns can be directed to this email address:
selmaberfin@sabanciuniv.edu

I understand:

- Participating in the study is on a voluntary basis.
- I can change my mind/stop the experiment at any point.
- Participating in this study does not involve any risks.

- I have to complete the study in one session.

___ I have read and understood the information above.

___ I agree to participate in this research voluntarily.

APPENDIX B

Study Instructions

Introduction:

“Hello!

In this experiment, you will be presented with a series of short stories in silent clips and audio recordings that you will be asked to recount just afterwards. Sometimes, you will be asked to recount the story **as accurately as possible**, and other times you will be asked to **change a few details** in the story.

Note: Some of these stories will be in audio form, so make sure your volume is turned up on your device and listen with headphones if possible.”

Video Cue:

“**Video Task**

On the next page, you will watch a **silent clip**. Please watch this clip carefully. When you feel ready, press *Next* to continue.”

Video Instructions:

1. “You have just watched a clip of Mary. Now imagine that you are talking to Mary’s partner in the photo. Her partner found out that Mary tore up the photo today because of their previous disagreement. However, in order to make things better between them, you must convince her partner that Mary did not tear the photo. Tell a deceptive version of the story you just watched in the text box below by changing at least three of Mary’s actions described in the video.”
2. “You have just watched a clip of George. Now imagine that you are talking to George’s mother. His mother cares about food safety and expects her son to eat safely as well. Therefore, you want his mother to know what George did today. In the text box below, tell the story presented to you in the video as accurately as possible.”
3. “You have just watched a clip of Sarah. Now imagine that you are talking to Sarah’s grandmother. Her grandmother found out that Sarah broke her antique Chinese vase today. However, in order to prevent Sarah from being punished, you

must convince her grandmother that Sarah did not break the vase. Tell a deceptive version of the story you just watched in the text box below by changing at least three of Sarah's actions described in the video."

4. "You have just watched a clip of Ethan. Now imagine that you are talking to Ethan's father. His father thinks his son is studying in his room. Therefore, you want his father to know what Ethan did today. In the text box below, tell the story presented to you in the video as accurately as possible."

5. "You have just watched a clip of Lily. Now imagine that you are talking to Lily's older sister. Her sister found out that the gold necklace that was sent as a gift to her was received by Lily today. But in order to avoid a rift between them, you must convince her sister that Lily did not take the necklace. Tell a deceptive version of the story you just watched in the text box below by changing at least three of Lily's actions described in the video."

6. "You have just watched a clip of Trevor. Now imagine that you are talking to Trevor's roommate. His roommate thinks he broke the door handle because the handle came off on his hands. Therefore, you want his roommate to know what Trevor did today. In the text box below, tell the story presented to you in the video as accurately as possible."

7. "You have just watched a clip of Angela. Now imagine that you are talking to Angela's next-door neighbor. Her neighbor found out that Angela broke their window today. However, since Angela cannot afford the glass repair costs, you must convince her neighbor that Angela did not break the window. Tell a deceptive version of the story you just watched in the text box below by changing at least three of Angela's actions described in the video."

8. "You have just watched a clip of Luke. Now imagine that you are talking to Luke's uncle. His uncle wonders why Luke is late in visiting him. Therefore, you want his uncle to know what Luke did today. In the text box below, tell the story presented to you in the video as accurately as possible."

9. "You have just watched a clip of Olivia. Now imagine that you are talking to Olivia's teacher. Her teacher saw that Olivia did not do her homework and built a sand castle on the beach today. However, to prevent Olivia from failing the course, you must convince her teacher that Olivia was not playing around the beach. Tell a deceptive version of the story you just watched in the text box below by changing at least three of Olivia's actions described in the video."

10. "You have just watched a clip of Tom. Now imagine that you are talking to the parents of Tom's child patient. The parents give a lot of importance that their

child does not consume sugar and eats healthily. Therefore, you want them to know what Tom did today. In the text box below, tell the story presented to you in the video as accurately as possible.”

11. “You have just watched a clip of Betty. Now imagine that you are talking to a forester on the lakeside. The forester found out that Betty lit a fire by the lake today, although it was forbidden. However, in order to avoid an official penalty for Betty, you must convince the forester that Betty did not light a fire. Tell a deceptive version of the story you just watched in the text box below by changing at least three of Betty’s actions described in the video.”

12. “You have just watched a clip of Gary. Now imagine that you are talking to Gary’s little brother. His brother thinks that the spoon that Gary used is clean, thus he thinks of using it. Therefore, you want his brother to know what Gary did today. In the text box below, tell the story presented to you in the video as accurately as possible.”

Audio Cue:

“Listening Task

On the next page, you will listen to an **audio recording**. **Make sure your device’s volume is turned up.** Please listen to this recording carefully. When you feel ready, press *Next* to continue.”

Audio Instructions:

1. “You have just listened to a recording about Angela. Now imagine that you are talking to Angela’s next-door neighbor. The neighbor found out that Angela broke their window today. Since Angela cannot afford the glass repair costs, you must convince her neighbor that Angela did not break the window. Tell a deceptive version of the story you just heard in the text box below by changing at least three of Angela’s actions described in the recording.” 2. “You have just listened to a recording about Luke. Now imagine that you are talking to Luke’s uncle. The uncle wonders why Luke is late in visiting him. Therefore, you want his uncle to know what Luke did today. In the text box below, tell the story presented to you in the recording as accurately as possible.”

3. “You have just listened to a recording about Olivia. Now imagine that you are talking to Olivia’s teacher. Her teacher saw that Olivia did not do her homework and instead built a sand castle on the beach today. In order to prevent Olivia from failing the course, you must convince her teacher that Olivia was not playing on the beach. Tell a deceptive version of the story you just heard in the text box below by changing at least three of Olivia’s actions described in the recording.”

4. “You have just heard a recording about Tom. Now imagine that you are talking to the parents of Tom’s patient. For the parents, it is important that their child not consume sugar and eat healthily. Therefore, you want them to know what Tom did today. In the text box below, tell the story presented to you in the recording as accurately as possible.”
5. “You have just listened to a recording about Betty. Now imagine that you are talking to a forest ranger on the side of the lake. The forest ranger found out that Betty lit a fire by the lake today, although it was forbidden. In order to avoid an official penalty for Betty, you must convince the forest ranger that Betty did not light a fire. Tell a deceptive version of the story you just heard in the text box below by changing at least three of Betty’s actions described in the recording.”
6. “You have just listened to a recording about Gary. Now imagine that you are talking to Gary’s little brother. His brother thinks that the spoon that Gary used is clean, thus he thinks of using it himself. Therefore, you want his brother to know what Gary did today. In the text box below, tell the story presented to you in the recording as accurately as possible.”
7. “You have just listened to a recording about Mary. Now imagine that you are talking to Mary’s partner, shown in the photo. Her partner has no idea how serious the disagreement between them is for Mary. Therefore, you want Mary’s partner to know what Mary did today. In the text box below, tell the story presented to you in the recording as accurately as possible.”
8. “You have just listened to a recording about George. Now imagine that you are talking to George’s mother. His mother cares about food hygiene, and she found out that her son ate an apple that fell on the floor today. In order to prevent George from being scolded, you must convince his mother that George did not eat a contaminated apple slice. Tell a deceptive version of the story you just heard in the text box below by changing at least three of George’s actions described in the recording.”
9. “You have just listened to a recording about Sarah. Now imagine that you are talking to Sarah’s grandmother. The grandmother was very upset when she saw that her antique Chinese vase was broken today. Therefore, you want her grandmother to know what Sarah did today. In the text box below, tell the story presented to you in the recording as accurately as possible.”
10. “You have just listened to a recording about Ethan. Now imagine that you are talking to Ethan’s father. His father saw that Ethan did not study and instead played with a paper airplane in his room today. In order to prevent Ethan from

being punished, you must convince his father that Ethan was not playing in his room today. Tell a deceptive version of the story you just heard in the text box below by changing at least three of Ethan's actions described in the recording."

11. "You have just listened to a recording about Lily. Now imagine that you are talking to Lily's older sister. The sister is looking for the gold necklace that was sent to her as a gift today. Therefore, you want her sister to know what Lily did today. In the text box below, tell the story presented to you in the recording as accurately as possible."

12. "You have just listened to a recording about Trevor. Now imagine that you are talking to Trevor's roommate. The roommate found out that Trevor broke the door handle today. In order to avoid a rift between them, you must convince his roommate that Trevor did not break the door handle. Tell a deceptive version of the story you just heard in the text box below by changing at least three of Trevor's actions described in the recording."

APPENDIX C

The List of the Storylines

Story Number	English (original) WC = 29-45	Turkish WC = 17-30	French WC = 26-41	Japanese Characters = 57-89
1	Mary is sitting at the kitchen table with a photograph in front of her. She pours juice into a glass and drinks it. She suddenly tears the photo in two and throws the pieces into the bin next to her. (Word count = 40)	Zeynep mutfak masasında önünde bir fotoğrafla oturur. Bardağa meyve suyu koyar ve içer. Zeynep birden fotoğrafı ikiye ayırır ve parçaları yanındaki çöp kutusuna fırlatır. (Word count = 24)	Elsa est assise à la table de la cuisine devant une photo. Elle se verse un verre de jus de fruit et le boit. Soudain, elle déchire la photo en deux et jette les morceaux dans la poubelle à côté d'elle. (Word count = 41)	あさみは、台所のテーブルで、写真を目の前にして座っている。彼女はジュースをグラスに注ぎ、飲む。突然写真を二つに破り、その破片を隣のゴミ箱に投げ捨てる。 (77 characters)
2	George is in the kitchen. He cuts a big apple into slices on the counter and then accidentally drops a slice on the floor. He quickly picks it up and he eats it. (Word count = 33)	Yusuf mutfaktadır. Tezgahta büyük bir elmayı dilimler ve yanlışlıkla bir dilimi yere düşürür. Hızlıca dilimi yerden alır ve yer. (Word count = 19)	Thierry est dans la cuisine. Il coupe une pomme en quartiers sur le comptoir et sans faire exprès, il en fait tomber un morceau par terre. Il le ramasse vite et le mange. (Word count = 33)	ゆうじは、台所にいる。彼はキッチンカウンターの上で大きなリンゴを切り、切ったリンゴを間違って床に落としてしまう。彼はすぐにリンゴを拾い、食べる。 (72 characters)

3	Sarah is in her living room. She pulls the couch closer to the TV. She accidentally bumps into the coffee table and she knocks over a vase. She carefully collects the broken pieces. (Word count = 33)	Seher oturma odasındadır. Kanepeyi televizyonun yakınına çeker. Yanlışlıkla kahve sehпасına çarpar ve bir vazoyu düşürür. Titizlikle kırılan parçaları toplar. (Word count = 19)	Julie est dans son salon. Elle déplace le canapé pour le rapprocher de la télévision. Sans faire exprès, elle se cogne contre la table basse et fait tomber un vase. Avec prudence, elle ramasse les morceaux. (Word count = 36)	さゆりは、リビングにいる。彼女はソファをテーブルの近くまで引っ張る。サイドテーブルにぶつかって、花瓶を落としてしまう。割れた破片を丁寧に拾い集める (82 characters)
4	Ethan is in his bedroom. He folds a piece of paper to make it into an airplane. He throws it into the air, and it immediately crashes into the ground. He stomps on the airplane. (Word count = 35)	Erhan odasındadır. Uçak yapmak için bir kağıt parçasını katlar. Uçağı havaya fırlatır ve uçak hemen yere çakılır. Erhan uçağı ayaklarıyla ezer. (Word count = 21)	Richard est dans sa chambre. Il plie une feuille de papier pour en faire un avion. Il le lance, mais aussitôt l'avion tombe par terre. Il le piétine. (Word count = 28)	しょうたは、寝室にいる。彼は飛行機を作るために、紙を折る。飛行機を飛ばすが、すぐに地面に落ちてしまう。彼は飛行機を踏みつぶす (63 characters)

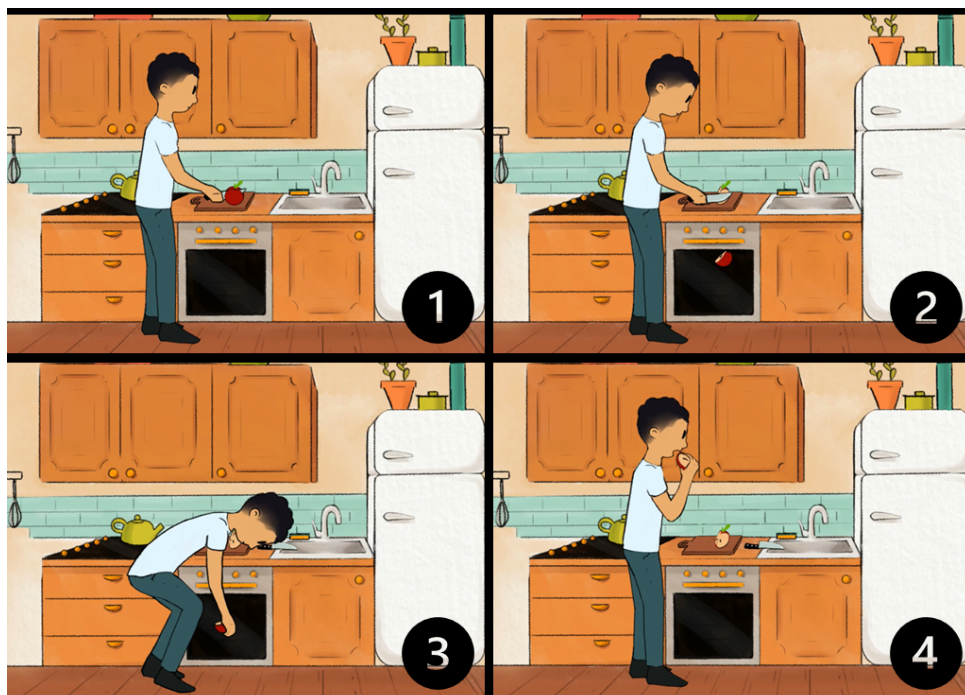
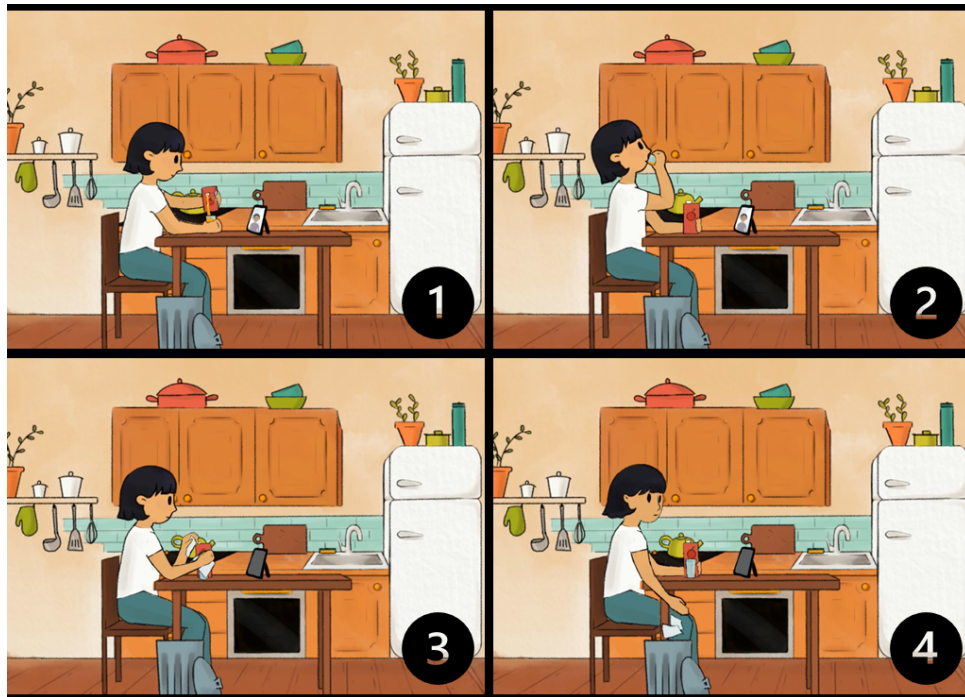
5	Lily is in front of her house. She finds a small box on her doorstep. She shakes it hesitantly then she opens it. Inside the box there is a gold necklace. She puts it on. (Word count = 33)	Leyla evinin önündedir. Kapısının eşiğinde küçük bir kutu bulur. Tereddütle kutuyu sallar ve açar. Kutunun içinde altın bir kolye vardır. Leyla kolyeyi takar. (Word count = 23)	Louna est devant chez elle. Elle découvre un petit paquet sur le pas de sa porte. Elle le secoue délicatement et l'ouvre. A l'intérieur se trouve un collier en or. Elle le met. (Word count = 33)	りょうこは、家の前にいる。ドアの前に小さな箱を見つける。彼女は戸惑いながらも箱を振り、開ける。中に入っていたのは金のネックレスだ。彼女はそのネックレスをつける。 (78 characters)
6	Trevor is in the bathroom and he breaks the doorknob off of the door. He pushes the door with all his force. He bangs on the door with his fists, and then he kicks the door open. (Word count = 37)	Tolga banyodadır. Kapının kolunu kırar. Tüm gücüyle kapıyı iter. Yumruklarıyla kapıya vurur ve tekme atarak kapıyı açar. (Word count = 17)	Antoine est dans la salle de bains et casse la poignée de la porte. Il pousse sur la porte de toutes ses forces. Il tape dessus avec les poings, puis l'ouvre d'un coup de pied. (Word count = 35)	たかしは、洗面所において、ドアノブを壊してしまふ。彼は、力いっぱいドアを押す。拳でドアを叩き始め、それからドアを蹴破る。 (60 characters)
7	Angela is in the garden with her dog. She waters the plants and then she pets her dog. She throws a ball for the dog and breaks a window. (Word count = 29)	Özge köpeğiyle birlikte bahçededir. Bitkileri sular ve ardından köpeğini oksar. Topu köpeğine doğru fırlatır ve bir cam kırar. (Word count = 18)	Angela est dans le jardin avec son chien. Elle arrose les plantes puis caresse le chien. Elle lance une balle au chien et casse une vitre. (Word count = 26)	あかりは、犬と一緒に庭にいる。彼女は植物に水をやり、犬をなでる。彼女は犬に向かってボールを投げ、窓を割ってしまう。 (57 characters)

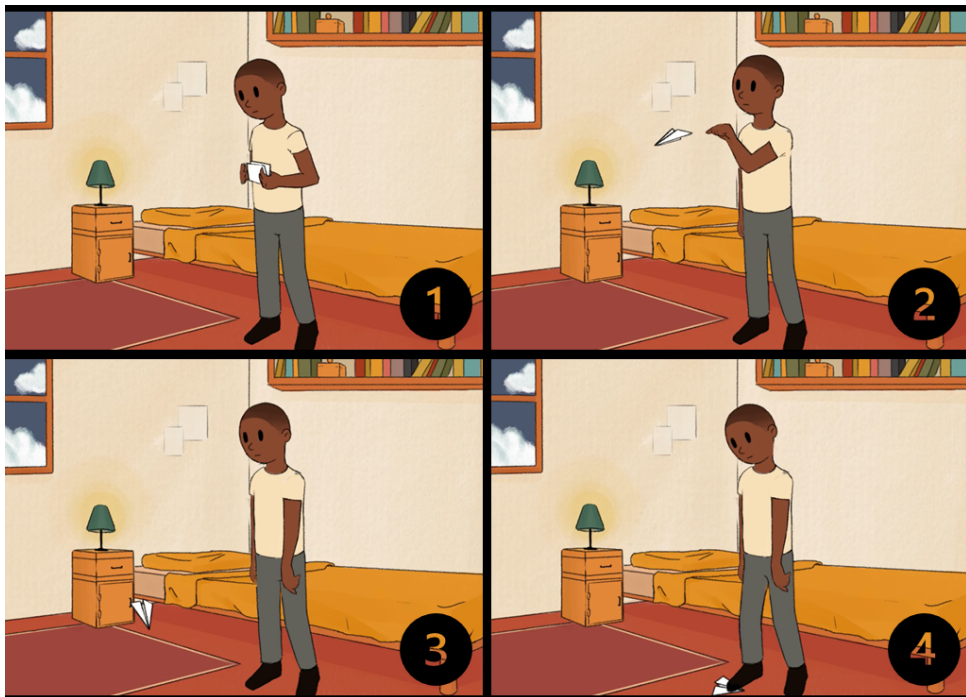
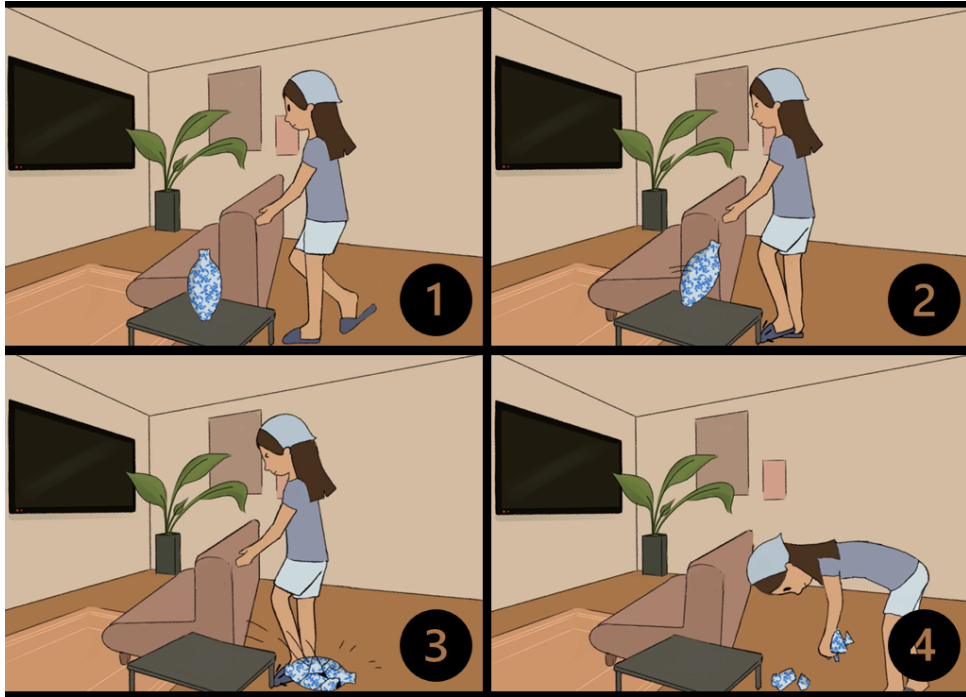
8	Luke is at an airport food stall. He buys a sandwich, and he accidentally drops his passport. At the airport security checkpoint, he searches his pockets for the passport. He follows an officer to the lost and found. (Word count = 38)	Murat, havalimanının yemek bölümündedir. Bir sandviç alır ve yanlışlıkla pasaportunu düşürür. Havalimanı kontrol noktasında ceplerinde pasaportu arar. Görevlinin peşinden kayıp bürosuna gider. (Word count = 22)	Thomas est dans le hall de restauration à l'aéroport. Il achète un sandwich et sans faire exprès fait tomber son passeport. Au contrôle de sécurité, il cherche son passeport dans ses poches. Il suit un employé aux objets trouvés. (Word count = 39)	りゅうたは、空港の売店にいる。サンドイッチを買い、誤ってパスポートを落としてしまう。彼は空港の検問所で、ポケットの中にあるはずのパスポートを探す。落とし物の預かり遺失物取扱所に行くためを探して、警備員の後を追いかける。 (89 characters)
9	Olivia is at the beach. She fills a bucket with sand, and turns it over. She lifts the bucket off, and then she carefully places seashells on the top of the sandcastle. (Word count = 32)	Duru plajdadır. Bir kovayı kum ile doldurur ve ters çevirir. Kovayı kaldırır ve dikkatlice deniz kabuklarını kumdan kalenin üzerine yerleştirir. (Word count = 20)	Olivia est à la plage. Elle remplit un seau avec du sable et le renverse. Elle retire le seau, puis elle dispose délicatement des coquillages sur le haut du château de sable. (Word count = 32)	いずみは、海岸にいる。彼女はバケツを砂で満たして、ひっくり返す。彼女はバケツを持ち上げ、できた砂のお城の上に注意深く貝殻を置く。 (61 characters)

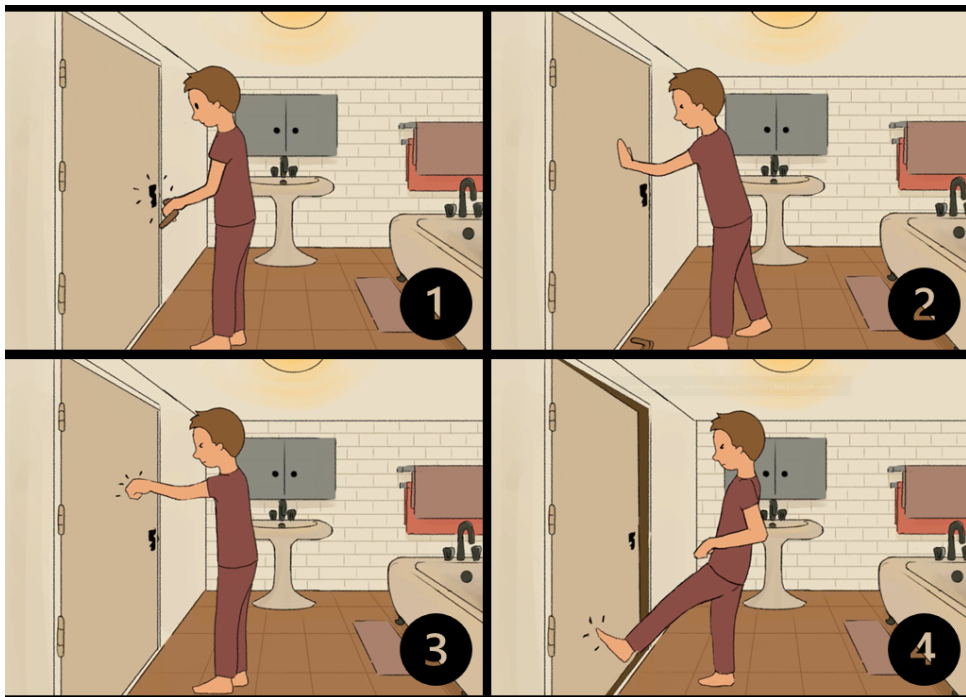
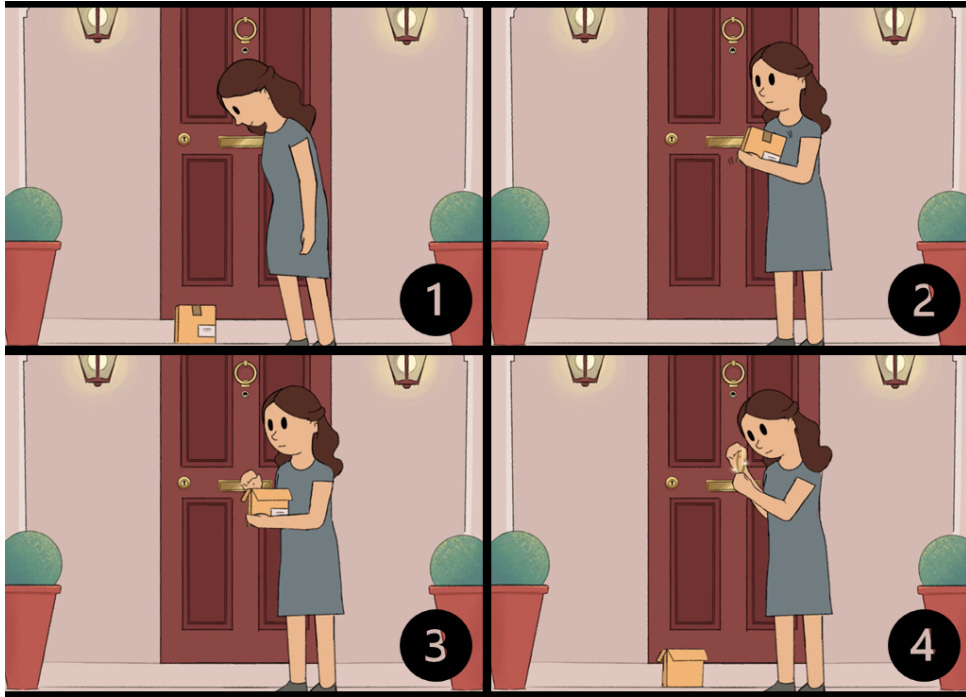
10	Tom is a nurse at the hospital, with a little child who has a broken arm. He lifts the child up onto the table and examines the plaster cast. He removes the cast with a medical saw and gives the child a piece of candy. (Word count = 45)	Tamer bir hastanede hemşiredir ve yanında kolu kırık küçük bir çocuk vardır. Çocuğu masanın üzerine taşır ve alçıyı kontrol eder. Tıbbi bir testereyle alçıyı çıkarır. Çocuğa bir tane şeker verir. (Word count = 30)	Simon est infirmier à l'hôpital avec un jeune enfant qui a le bras cassé. Il porte l'enfant sur la table d'examen. Il regarde le plâtre. Avec une scie médicale, il l'enlève et donne à l'enfant un bonbon. (Word count = 37)	とあるは、病院の看護師で、腕を骨折した小さな子供と一緒にいる。子供を台机の上へ上げ、ギプスを検査する。彼は医療用のはさみを使ってギプスを外す。し、そして、子供に飴をあげる。 (79 characters)
11	Betty is at a campground by a lake. She catches a fish with a fishing rod. She collects wood and she lights a fire. She grills the fish over the fire. (Word count = 31)	Betül göl kenarında bir kamp yerindedir. Oltayla balık tutar. Odun toplar ve ateş yakar. Balığı ızgara ateşinde pişirir. (Word count = 18)	Léa est dans un camping au bord d'un lac. Elle attrape un poisson avec une canne à pêche. Elle ramasse du bois et allume un feu. Elle fait griller le poisson sur le feu. (Word count = 34)	れいこは、湖のほとりのキャンプ場にいる。彼女は、釣り竿を使って魚を捕まえる。彼女は木を集めて、火をつける。彼女はその火で魚を焼く。 (61 characters)
12	Gary is in the kitchen. He mixes the ingredients together in a bowl. He then pours the batter into a baking pan. He licks the spoon and puts the cake in the oven. (Word count = 33)	Cem mutfaktır. Bir kasede malzemeleri çırpar. Daha sonra kek harcını kalıba döker. Kaşığı yalar ve keki fırına verir. (Word count = 18)	Jérôme est dans la cuisine. Il mélange des ingrédients dans un saladier. Il verse la pâte dans le moule. Il lèche la cuillère et met le gâteau dans le four. (Word count = 30)	こうじは、台所にいる。彼は、ボウルの中で材料を混ぜる。それから型に生地を注ぐ。彼は、スプーンを舐めてから、オーブンにケーキを入れる。 (65 characters)

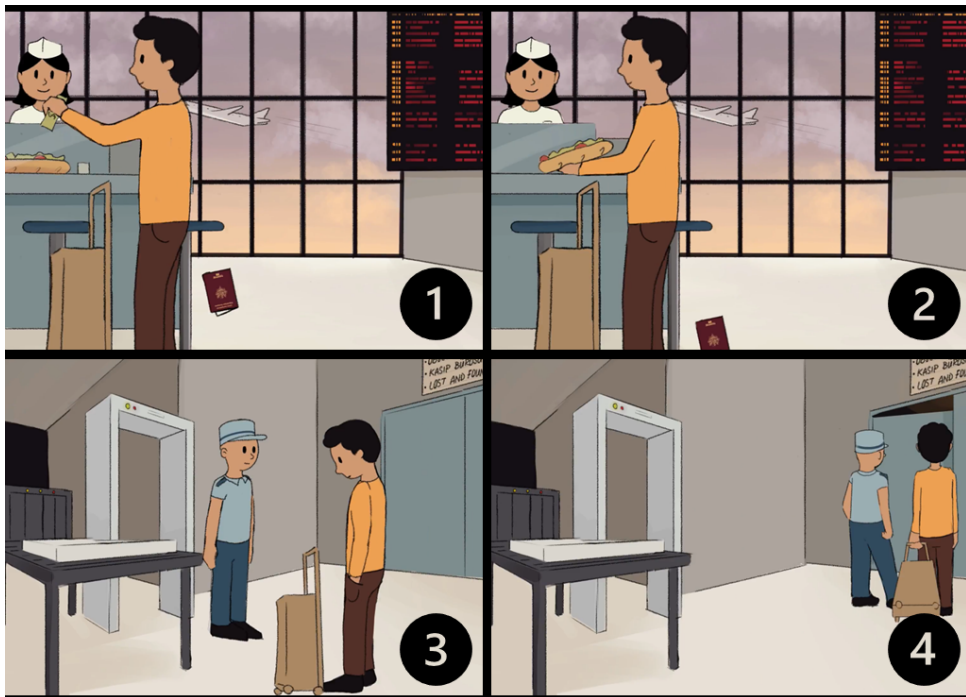
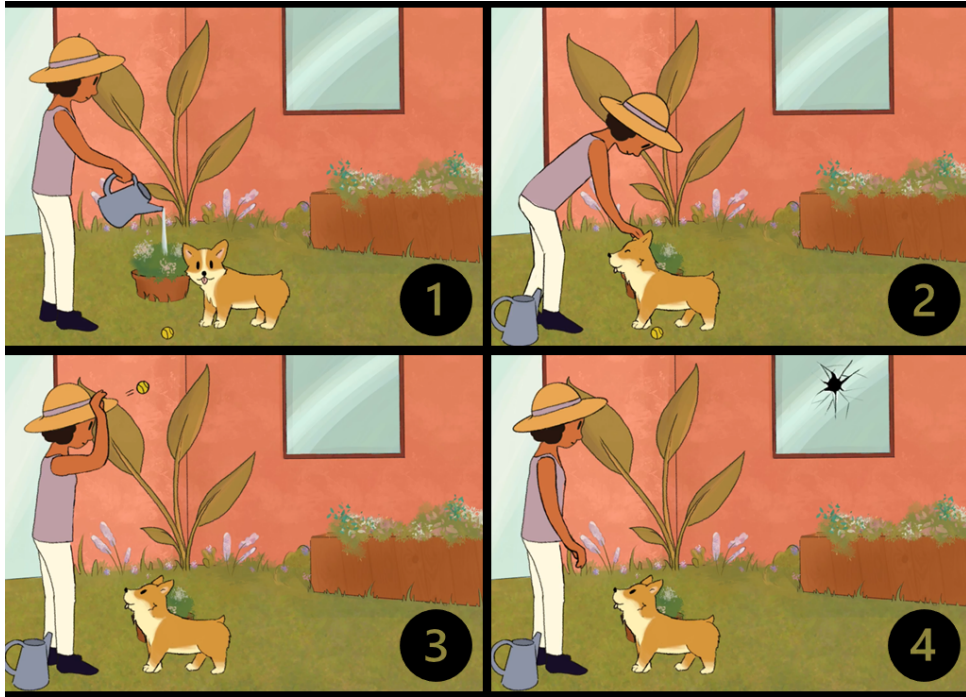
APPENDIX D

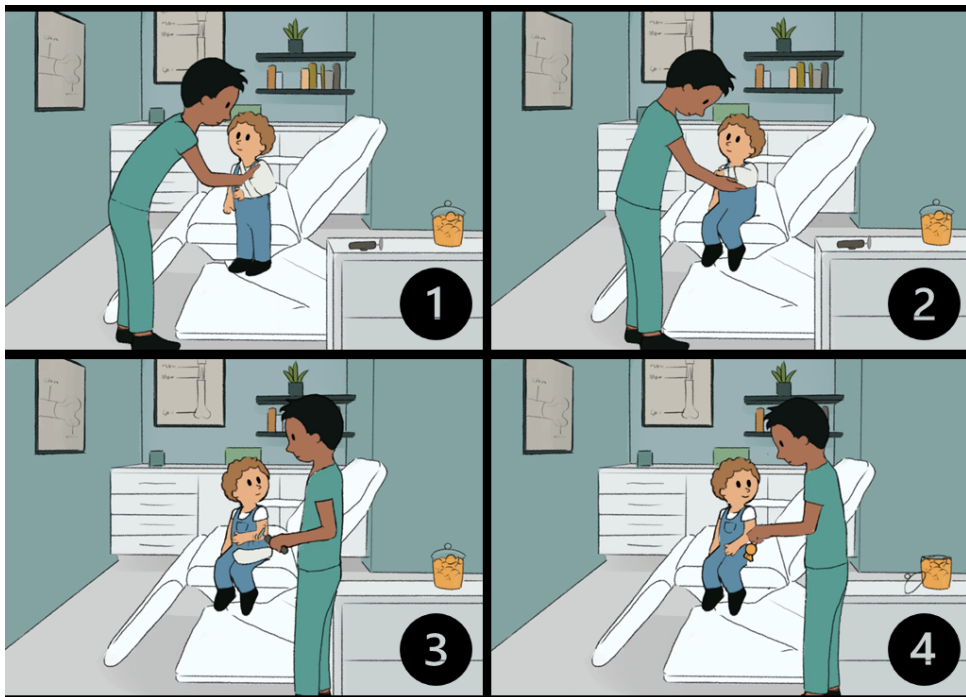
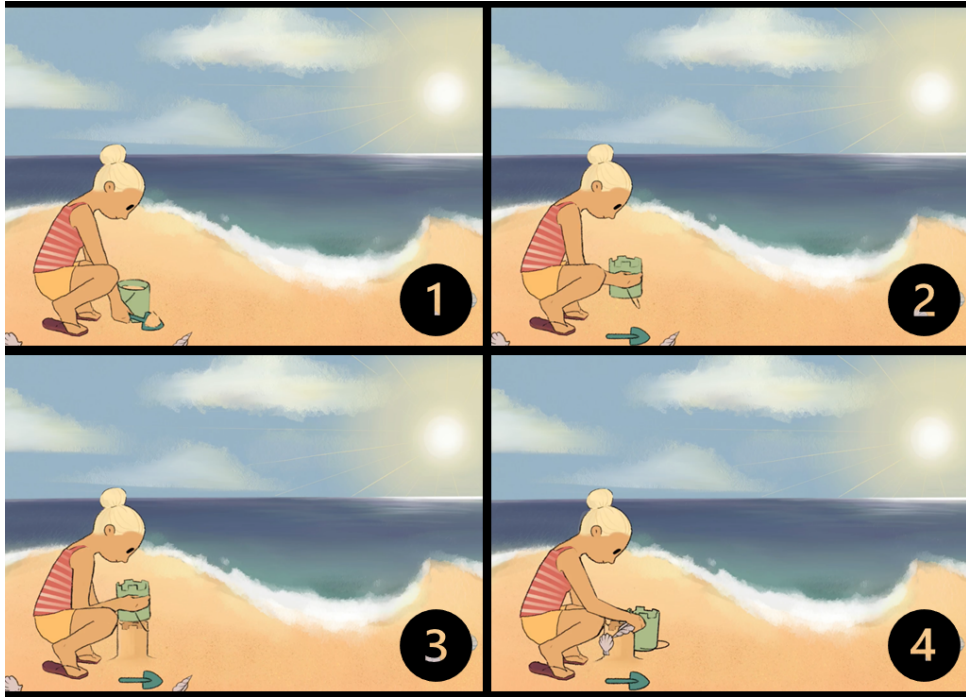
Example Arrays from the Visually Animated Video Clips

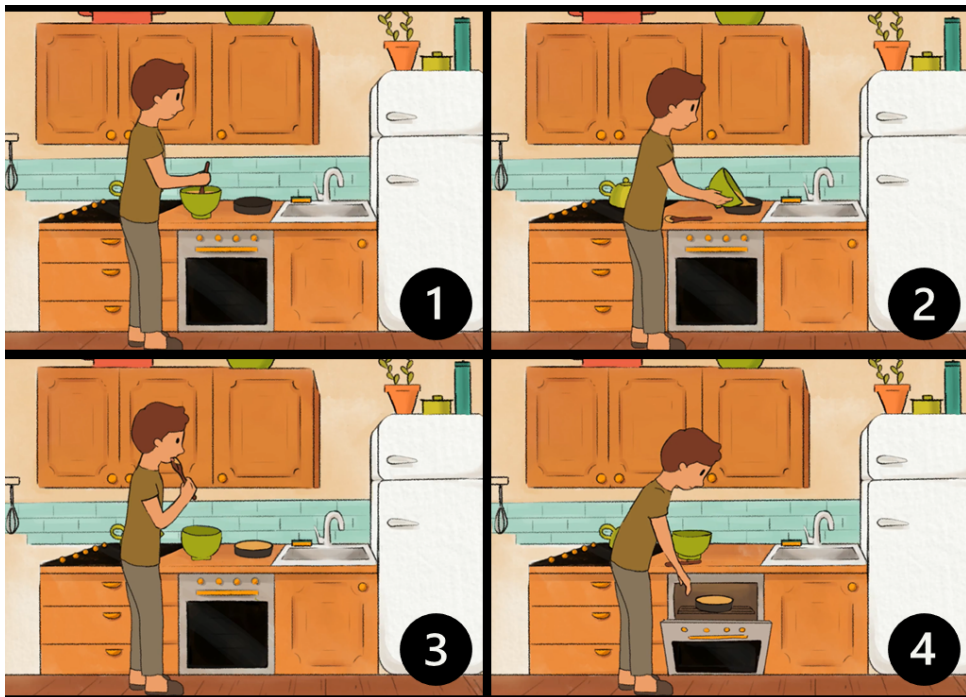
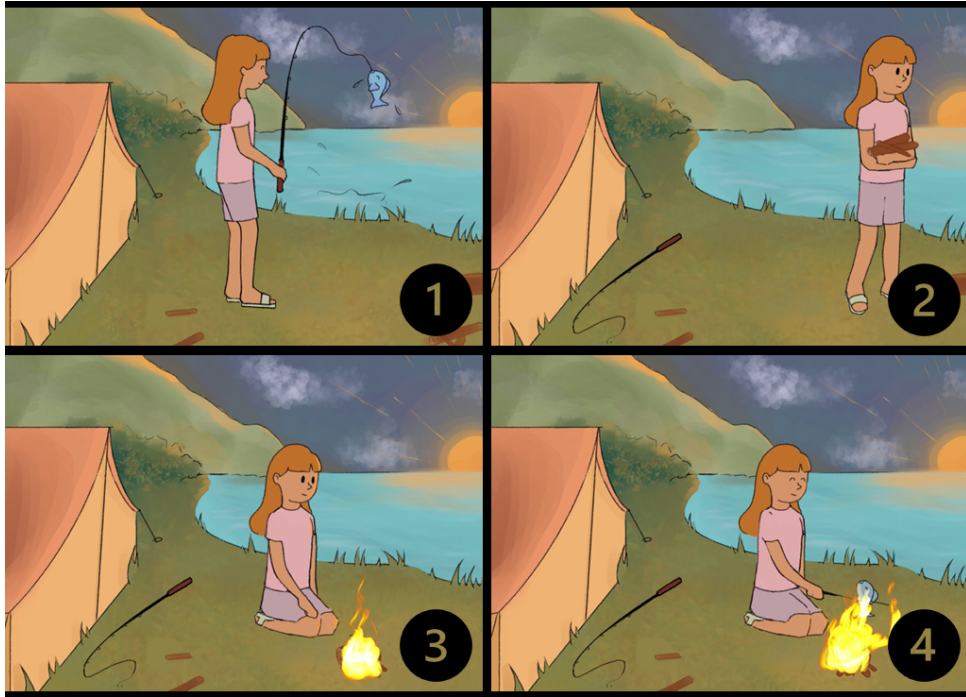












APPENDIX E

Demographics Questionnaire

1. Please enter your birth date in day, month, year format: (dd/mm/yy)
2. Gender: (Female / Male / Non-binary / I do not want to specify / Other (please specify))
3. The last educational degree you received: (Primary school / High school / Community college / Bachelor's degree / Master's degree / Ph.D.)
4. Profession:
5. Where do you currently live (city, state, country):
6. Native language: (English / Other (please specify))
7. At what age did you learn English? (English is my native language / 0-5 years old / 5-10 years old / 10-15 years old / 15+ years old)
8. Please choose the option that describes you the best: (I cannot speak a second language / I have learned a second language at school only / I'm proficient in my second language / I'm bilingual or multilingual)
9. Please list all the languages you can speak other than English, separated by commas (e.g., German, Italian)
10. Your dominant hand: (Right / Left / I use both of my hands equally)
11. How often do you lie in your daily life? (Never / Rarely / Sometimes / Often / Always)
12. How often do people in your culture tell lies? (Never / Rarely / Sometimes / Often / Always)

APPENDIX F

Combinations of the Experimental Conditions by Participant Groups

Group 1			Group 2		
First block					
Story 1 (female)	lie	video	Story 7 (female)	truth	video
Story 2 (male)	truth	video	Story 8 (male)	lie	video
Story 3 (female)	lie	video	Story 9 (female)	truth	video
Story 4 (male)	truth	video	Story 10 (male)	lie	video
Story 5 (female)	lie	video	Story 11 (female)	truth	video
Story 6 (male)	truth	video	Story 12 (male)	lie	video
Second block					
Story 7 (female)	lie	audio	Story 1 (female)	truth	audio
Story 8 (male)	truth	audio	Story 2 (male)	lie	audio
Story 9 (female)	lie	audio	Story 3 (female)	truth	audio
Story 10 (male)	truth	audio	Story 4 (male)	lie	audio
Story 11 (female)	lie	audio	Story 5 (female)	truth	audio
Story 12 (male)	truth	audio	Story 6 (male)	lie	audio

APPENDIX G

Sample Formulas

Sample formula for Binomial mixed model regression with a logistic link function:

Level 1 (trials):

$$\text{logit}(\hat{\text{Grammar}}_{ij}) = \log\left(\frac{\text{Grammar}}{1 - \text{Grammar}}\right)_{ij} = \beta_{0j} + \beta_{1j} \cdot (\text{Veracity}) + \beta_{2j} \cdot (\text{Modality})$$

Level 2 (participants):

$$B_{0j} = \gamma_{00} + u_{0j}$$

$$B_{1j} = \gamma_{10}$$

$$B_{2j} = \gamma_{20}$$

Mixed Model:

$$\text{logit}(\hat{\text{Grammar}}_{ij}) = \gamma_{00} + \gamma_{10} \cdot (\text{Veracity}) + \gamma_{20} \cdot (\text{Modality}) + u_{0j}$$

Sample formula for Poisson mixed model regression with a logarithmic link function:

Level 1 (trials):

$$\log(\hat{\text{Grammar}}_{ij}) = \beta_{0j} + \beta_{1j} \cdot (\text{Veracity}) + \beta_{2j} \cdot (\text{Modality})$$

Level 2 (participants):

$$B_{0j} = \gamma_{00} + u_{0j}$$

$$B_{1j} = \gamma_{10}$$

$$B_{2j} = \gamma_{20}$$

Mixed Model:

$$\log(\hat{\text{Grammar}}_{ij}) = \gamma_{00} + \gamma_{10} \cdot (\text{Veracity}) + \gamma_{20} \cdot (\text{Modality}) + u_{0j}$$