

Why are implicit causes predictable?1

The predictability of implicit causes: testing frequency and topicality explanations

Shuang Guan

Department of Linguistics, Swarthmore College, Swarthmore

and

Jennifer E. Arnold

Department of Psychology and Neuroscience, University of North Carolina at Chapel Hill,

Chapel Hill

Keywords: Implicit Causality, predictability, reference, discourse

ABSTRACT

In discourses involving implicit causality, the implicit cause of the event is referentially predictable; i.e. it is likely to be re-mentioned. But it is unclear how referential predictability is calculated. We test two possible explanations: 1) The frequency account suggests that people learn that implicit causes are predictable through experience with the most frequent patterns of reference in natural language; 2) The topicality account asks whether implicit causes tend to play topical roles in the discourse, which itself may lead to the perception of discourse accessibility. With two text analyses, we show that implicit causes are frequently re-mentioned, but only if we consider a narrow set of discourse circumstances, which would require comprehenders to track contingent frequencies. We found no evidence for the topicality account: in two experiments, implicit causality affected predictability but not topicality, and in a corpus of natural speech, implicit causes tended to not occupy topical positions.

Introduction

Understanding language tends to be faster and easier when information is **predictable**, meaning that it is redundant in part with the preceding context. Of particular relevance to the current study is predictability at the discourse level¹; reference comprehension is facilitated when the referent can be anticipated from the context (Altmann & Kamide, 1999; Arnold, Hudson Kam, & Tanenhaus, 2007; Arnold & Lao, 2008; Brocher, Chiriacescu, & von Heusinger, 2016; Delogu et al., 2020; Kehler, Kertz, Rohde, & Elman, 2008; Lowder & Ferreira, 2016; Stevenson, Crawley, & Kleinman, 1994; Tily & Piantadosi, 2009), where such activation is assumed to be partial and probabilistic. For example, following *Alice is my sister. I love...*, we might expect the referent Alice to be re-mentioned. This expectation is independent of linguistic form, given that the speaker could say *Alice*, *my sister*, or *her*. Yet there are many open questions about what it means for a referent to be predictable, and precisely how contextual features translate into representations of referential predictability.

In this paper we specifically test how referential predictability emerges from **implicit causality**, which refers to the tendency for people to assume that one character is more closely related to the cause of an event than the other. For example, in *Ana admired Liz* or *Matt impressed Will*, people tend to assume that Liz and Matt were the causes of those events (inter alia, Au, 1986; Bott & Solstad, 2014; Brown & Fish, 1983; Caramazza, Grober, Garvey, & Yates, 1977; Crinean & Garnham, 2006; Garvey & Caramazza, 1974; Solstad & Bott, 2017; Stevenson et al., 1994). Implicit causality has drawn the attention of researchers because causal inferences are a fundamental part of language comprehension (e.g., Kaiser, 2019; Kehler & Rohde, 2017;

¹ It is also well established that comprehension is facilitated by the predictive activation of sounds, words, and syntactic structures (inter alia Coulson, Federmeier, Van Petten, & Kutas, 2005; Falkauskas & Kuperman, 2015; Federmeier & Kutas, 2001; Kochari & Flecken, 2019; Kowalski & Huang, 2017; Kutas & Hillyard, 1984; Levy, 2008; Pickering & Garrod, 2007; Ryskin, Mimnaugh, Brown-Schmidt, & Federmeier, 2019; Smith & Levy, 2013; Viebahn, Ernestus, & McQueen, 2015).

Why are implicit causes predictable?⁴

Magliano, Baggett, Johnson, & Graesser, 1993), and some verbs elicit causal inferences more than others (Solstad & Bott, 2017). In addition, scholars debate whether people impute implicit causes based on real-world knowledge (Pickering & Majid, 2007; Majid, Sanford, & Pickering, 2007) or primarily from lexical representations (Hartshorne, 2014; Bott & Solstad, 2014).

The current research ignores the debates about where implicit causality biases come from, and instead examines how these biases relate to linguistic expectations. We build on widespread agreement that people can and do draw inferences about the likely cause of an event, and that these implicit causality judgments are related to both an expectation for an explanation (Bott & Solstad, 2014; Kehler & Rohde, 2017; Solstad & Bott, 2017) and specifically an expectation for reference to the person assumed to be the cause (Kehler et al., 2008; Kehler & Rohde, 2013, 2019; Rohde & Kehler, 2014; Stevenson et al., 1994). Our question is: how do implicit causality biases affect referential expectations?

Much of the work on implicit causality comes from work showing that implicit causality affects both judgments about who will be mentioned next and the interpretation of ambiguous pronouns. For example, in *Ana amazed Liz because she...* most people interpret the ambiguous pronoun *she* to be Ana, the implicit cause of the interpersonal event; Ana probably did something to cause Liz to be amazed. On the other hand, in *Ana admired Liz because she....* most people interpret *she* to be Liz; Liz probably did something to make Ana admire her. Verbs like *amaze* are subject-biased because they impute causality to the entity in subject position, while verbs like *admire* are object-biased because they impute causality to the entity in object position (e.g., Kehler & Rohde, 2013). It is well established that implicit causality guides the final interpretation of ambiguous pronouns (inter alia, Garnham, Oakhill, & Cruttenden, 1992; Garvey & Caramazza, 1974; Hartshorne, O'Donnell, & Tenenbaum, 2015; Kehler et al., 2008; Kehler &

Why are implicit causes predictable?5

Rohde, 2013; Rohde & Kehler, 2014; Stevenson et al., 1994). It also speeds the interpretation of sentences where the gender of the pronoun matches the implicit cause bias (Koornneef & Van Berkum, 2006).

The explanation for the observed effects on pronoun comprehension draws on the idea that implicit causes are generally expected to be mentioned again (e.g. Bott & Solstad, 2014; Kehler et al., 2008; Kehler & Rohde, 2013, 2019; Johnson & Arnold, 2021; Rohde & Kehler, 2014;), and that implicit causes may be focused in mental representations of the discourse (Cozijn, Commandeur, Vonk, & Noordman, 2011; Koornneef & Van Berkum, 2006; McDonald & MacWhinney, 1995; McKoon, Greene, & Ratcliff, 1993; but for a different view see Garnham et al., 1996; Stewart, Pickering, & Sanford, 2000). The idea that some referents are anticipated is not specific to implicit causes, and references can become expected for other reasons (e.g., Kehler et al., 2008; Langlois, Zerkle, & Arnold, under review; Stevenson et al., 1994). When the prior semantic context creates referential expectations, as with implicit causality, the effects are tied to the **coherence relation** between the two clauses. For example, implicit causes are expected when people expect that the following clause will specify an explanation of the previous event. Connector words like “so” or “because” can be used to modulate the expectation of the upcoming relation, and the presence of “because” supports the strongest effects of implicit causality on pronoun interpretation and referential (Crinean & Garnham, 2006; Ehrlich, 1980; Järvikivi, van Gompel, & Hyönä, 2017; Kehler et al., 2008; Koornneef & Sanders, 2013; Koornneef, Dotlačil, van den Broek, & Sanders, 2016; Mak, Tribushinina, & Andreiushina, 2013; Pyykkönen and Järvikivi, 2010; Stevenson et al., 1994; Stevenson, Knott, Oberlander, & McDonald, 2000).

Thus, understanding the relation between implicit causality and referential expectations

Why are implicit causes predictable?⁶

provides a window onto general processes of discourse comprehension, and in particular the relationship between causal judgments and other discourse representations. Yet relatively little work has explicitly addressed the question of how next-mention expectations stem from judgments about the likely cause of an event. We consider three existing proposals in the literature: 1) The semantic inference account, 2) The frequency account, and 3) the discourse status account. These accounts are not mutually exclusive, but each offers a different explanation for the origin of next-mention expectations. Our study then focuses on testing proposals (2) and (3).²

Our view critically assumes that people have separate knowledge about events (why did event x happen?) and about language (what is the speaker likely to say next?) Empirically, implicit cause biases tend to be measured through linguistic tasks, for example by examining who is mentioned in a passage-completion task where people finish sentences like *Ana admired Liz because....* Pickering and Majid (2007) argue that the implicit causality bias is merely “an abstraction of the type of reason that is most likely to be provided for the event.” (p. 785) -- that is, implicit causality judgments are essentially linguistic judgments. Similarly, Solstad and Bott (2017) explain implicit causality in terms of language, proposing that “IC verbs...trigger expectations for specific explanation types,” (p. 21). While we acknowledge that most studies have used language to show implicit causality, in principle, we expect that judgments about causes and next-mention expectations can be represented separately. Here we examine three possible ways that causality could affect next-mention biases.

² For a similar question based on the role of connector words, see Mak et al., 2013.

Why are implicit causes predictable?7

The semantic inference account

One hypothesis is that language comprehension results in representations of events, and these event representations lead to judgments about what information the speaker is likely to mention next. Sentences like *Ana admired Liz because...* contain enough semantic information to judge that the most likely cause of admiration was Liz, as well as an expectation for an explanation, and thus, the expectation that Liz will be mentioned. This idea is implicit in Bayesian models, which suggest that pronoun comprehension is a function of a) the probability of a referent being mentioned, and b) the likelihood that a pronoun would have been used to refer to that referent, where both these terms are divided by the sum of these calculations for all contextual referents (e.g., Kehler et al., 2008; Kehler & Rohde, 2013, 2017; see also Frank & Goodman, 2010 for a related model of noun phrase modification). Of particular importance to the current study is component (a), which models likelihood of reference. On their model, this calculation is driven by the semantics of the prior discourse, including assumptions about the coherence relation between utterances (see Stevenson et al., 1994 for a similar idea).

The importance of semantics is made explicit in Hartshorne et al.'s (2015) Bayesian model, which is similar but instead models the probability of the utterance instead of the referent. Hartshorne et al.'s model. On their model, referential probabilities are incrementally calculated based on the semantics of the utterance. For example, after a fragment like *Archie angered Bart because he...*, the comprehender generates probabilities about the meaning of the upcoming utterance, based on the preceding context. The authors argue that the upcoming content is more likely to involve Archie than Bart, given the meaning of *anger*, and the implicit causality bias toward Archie. Thus, if the explanation is more likely to mention Archie, the pronoun *he* is more likely to refer to Archie than Bart.

Why are implicit causes predictable?8

Bott and Solstad (2014; Solstad & Bott, 2017) offer a different semantic account of implicit causality. On their view, Implicit Causality verbs explicitly represent the expectation for an explanation of the event. Thus, their model views implicit causality as an expectation for what will be said, as opposed to an assumption about causes in the world. They offer a typology of different types of explanations, and suggest that different verbs elicit expectations for particular types of causes, which are linked to either the subject or object referent.

In sum, the semantic inference account includes a class of diverse approaches in the literature, all of which depend on semantic representations of the verb and/or event.

The frequency-based account

The frequency-based account suggests that referential expectations in general (and not just for implicit causality) stem from experience with the most common patterns of reference. This idea was proposed as the Expectancy hypothesis (Arnold, 2010; Arnold et al., 2007; see also Arnold 1998, 2001). On this account, people track the frequency of referential patterns in discourse, based on common linguistic categories like syntactic or semantic role. To illustrate, if speakers frequently re-mention the entities that were in subject position, comprehenders can learn that in novel situations, the subject entity has a high probability of re-mention, making it referentially predictable. Thus, if speakers frequently re-mention implicit causes, comprehenders should expect re-mention of implicit causes. In the Expectancy hypothesis, real-world reference production becomes the input for reference prediction (for a similar idea for syntactic comprehension, see MacDonald, 2013).

We know that the frequency account is plausible for explaining the predictability of subjects, due to corpus data. Arnold, Strangmann, Hwang, & Zerkle (2018) found that in

Why are implicit causes predictable?⁹

sentences with two or more references, the subject was mentioned again (37%) more often than other referents (20%; see also Arnold, 1998, 2010). We also know that the frequency account could potentially explain the predictability of goal arguments in transfer verbs. For example, in *John passed the comic to Bill* and *Bill seized the comic from John*, John is the source, and Bill is the goal (Stevenson et al., 1994). This matches a tendency for people to talk more about goals than sources, as shown in corpus analyses (Arnold, 2001; Arnold, unpublished). In fact, Arnold (unpublished) found that the likelihood of re-mention for goals was robust even when coherence relation was controlled.

However, we don't know whether frequency could also explain the predictability of implicit causes, because there is little existing data on whether speakers tend to frequently refer to implicit causes or not. There is preliminary evidence that it might, from a corpus analysis reported by Long and de Ley (2000). They examined anaphoric subjects, and found that they are somewhat more likely to refer to implicit causes than non-causes, but only for object-biased verbs. We ask a different but related question: given an event that evokes implicit causality biases, is the implicit cause more likely to be re-mentioned? If so, it might provide the necessary input for listeners to learn that reference to the implicit cause is frequent, and therefore predictable.

Another open question is whether people keep track of frequency separately for different contexts. A **generalized frequency account** would suggest that if, for example, a language user encounters 1000 tokens of non-IC verbs where the subject is re-mentioned, and 10 tokens of an object-biased IC verb where the object is re-mentioned, they might generalize the observed frequency and expect subject re-mention across the board. A priori this account would not

Why are implicit causes predictable?¹⁰

explain implicit cause expectation, because only some verbs elicit strong implicit cause judgments.

More relevant to the current study, a **semantically-based frequency account** would suggest that people track semantic patterns of re-mention separately for different verb classes, for example noting that while transfer verbs tend to be followed by re-mention of the goal argument (Arnold, 2001), verbs with high implicit causality might exhibit high re-mention of the implicit cause. This account would hold if implicit causes are mentioned frequently in natural language. However, it is not a forgone conclusion that such a pattern exists. Experiments tend to use single-sentence contexts, but natural language includes richer contexts that may make the implicit explanation redundant. E.g., in *Will is a great dancer. Matt admired Will because...*, Will's dancing skills are unlikely to be mentioned again (Bott & Solstad, 2014).

The semantically-based frequency account could further occur in different levels of granularity. In a **verb-based version**, people might note the frequency of referent re-mention for verb classes, regardless of coherence relation. This might hold if, for example, sentences with emotion verbs tended to be frequently followed by explanations, or if implicit causes tended to be mentioned frequently for multiple different coherence relations. By contrast, a **coherence-contingent version** would be supported if people track the frequency of referent re-mention as a function of both verb class and coherence relation. We know that in sentence completion tasks, people tend to mention the implicit cause (y) more after *x admires y because...* than after *x admires y so...* (Stevenson et al., 1994). Does this mean that people categorize referential frequencies separately according to the coherence relation?

Why are implicit causes predictable? 11

In sum, the fundamental question is whether referential expectations follow from observed referential frequencies. We use two corpus analyses to test whether observed frequencies are in line with known implicit causality expectations.

The discourse status (topicality) account

A third idea stems from proposals in the literature that implicit causes are mentally focused (McKoon et al., 1993; Stevenson et al., 1994). This idea is rooted in the widespread assumption that language comprehension takes place in the context of a mental representation of the discourse context, often called a mental model or situation model (van Dijk & Kintsch, 1983; Johnson-Laird, 1983; Bower & Morrow, 1990). Here we ask whether implicit causes are focused because they are perceived to have a topical discourse status. Several theories suggest that some information is privileged in the mental model, in particular, the sentence topic (Ariel, 1990; Chafe, 1976; Reinhart, 1982). Yet other evidence suggests that implicit causes may not be perceived as topical, raising questions about the role of topicality in judgments of referential predictability.

One reason to examine the relationship between implicit causality and topicality is that topical things tend to be predictable. While the topic is notoriously hard to define, intuitively it represents what the sentence is “about”. It is also associated with syntactically prominent categories like the subject (Brennan, Friedman, & Pollard, 1987; Grosz, Joshi, & Weinstein, 1995; Hendriks, 2016; van Rij, van Rijn, & Hendriks, 2013); subjects are also referentially predictable in that they tend to be re-mentioned (Arnold, 1998; Arnold et al., 2018). Topicality and predictability also overlap in theoretical accounts. For example, Givón (1983) proposed that topicality falls on a continuum on which every discourse entity could be evaluated. On this

Why are implicit causes predictable?¹²

view, one measure of topicality is “persistence”, which reflects how long the entity will remain in the discourse.

The idea that topicality overlaps with predictability is also supported by evidence from other verb types. Zerkle and Arnold (2019) tested constructions like *Ana is cleaning up with Liz*, where both characters have the same semantic role. In this case, syntactic position is the only distinguishing characteristic, where the subject character (here, Ana) is more topical. They found that with this construction, people judged the subject character to be both more topical and more predictable, using the same methods as those used in our experiments. This finding corresponds to corpus evidence that subjects tend to be frequently re-mentioned (Arnold, 1998; Arnold et al., 2018). In addition, Arnold’s (2017) corpus analysis with transfer verbs found that goals tended to occur in discourse-prominent positions. Goals tended to be given, pronominal, animate, and 1st/2nd person more than sources. Thus, semantic predictability for transfer verbs is associated with discourse topicality and frequent re-mention.

Despite the frequent overlap of topicality and predictability, implicit causality verbs like *admire* or *impress* present a situation where they do not always align. If the subject is considered the topic, then the implicit cause is also the topic for *impress*, but not *admire*. In fact, this dissociation is central to the claim that pronoun production is influenced by only topicality (e.g., subjecthood) but not predictability (Kehler and Rohde, 2013; Kehler et al. 2008; Fukumura & van Gompel, 2010). This idea is supported by data from implicit causality contexts, where speakers produce pronouns more for subjects than objects, but not as a function of implicit causality (e.g., Fukumura & van Gompel, 2010; Rohde & Kehler, 2013, but for conflicting evidence see Weatherford & Arnold, 2021). However, this does not rule out the possibility that

Why are implicit causes predictable?13

predictability might also influence perceptions of topicality at least partially. We test whether it does in both a corpus analysis and experiment.

Current study

This paper provides a first step toward understanding how next-mention judgments stem from implicit causality. Given that the semantics of the verb are inherently present in all implicit causality scenarios, it is likely that they play some role. In two text analyses and two experiments, we directly test the frequency and topicality accounts of next-mention expectation. We focus on interpersonal verbs like *admire*, *like*, *impress*, *amaze*, which are known to elicit strong implicit causality judgments, and have been the focus of much of the work on implicit causality in language processing. Critically, the implicit cause is in subject position for some verbs and in object position for others, which allows us to test how predictability is influenced by both implicit causality and grammatical role. See Supplement for the criteria for verb selection.

Our first question is whether the frequency account is at all plausible. Could listeners learn from experience that speakers tend to re-mention implicit causes? If so, we would expect this pattern in natural data. To test this, we perform two small-scale text analyses (**Text Analysis 1** and **Text Analysis 2**), testing whether implicit causes do tend to be re-mentioned. If they are, it would be consistent with the frequency hypothesis, but it would not rule out either the semantic integration or discourse status hypotheses. If they are not, it would falsify the frequency hypothesis. Results from these analyses also begin to probe the question of the granularity with which frequencies might be learned from experience.

Our second question is whether implicit causality aligns with discourse status in naturally-occurring language. Even though several models make the simplifying assumption that the grammatical subject is the topic (Kehler & Rohde, 2013; van Rij et al., 2013), there are

Why are implicit causes predictable?¹⁴

several discourse properties that are associated with topicality, such as givenness and pronominalization. If predictability and topicality are jointly associated with implicit causality, we should see these properties correlating in our text analysis (**Text Analysis 2**).

As a direct test of the topicality question, in **Experiments 1 and 2** we asked participants to judge predictability and topicality for short discourses about interpersonal events like *admire* and *impress*. These experiments use a metacognitive task in which participants read a short story fragment and answer questions about either predictability (which character is most likely to be mentioned next?) or topicality (which character is this story about?). If implicit causality affects both predictability and topicality in the same way, it would support the topicality hypothesis.

Text Analysis 1: Google

Goals of analysis

Our goal was to understand the frequency with which speakers refer to entities appearing in implicit cause and non-cause roles, so we selected a sample of both subject-biased and object-biased verbs occurring in the transitive frame. For this analysis, our sample was designed to mimic the sorts of sentences typically used in implicit causality experiments, namely clauses with emotion verbs with two animate characters, and where the connector word “because” signaled a causal relationship with the following clause. We used Google to search the internet for tokens, which gave us easy access to a large sample of language, with a mix of spoken and written, and in different genres. For details on selection criteria and examples, see Appendix A.

Methods

Sample and Selection Criteria

We collected a sample of tokens where both implicit cause and non-cause roles were animate and matched for salience and animacy. To target tokens with following explanations, we

Why are implicit causes predictable?15

limited our sample to utterances with the explicit connective *because*. To do this, we searched for strings that included two animate pronouns and each implicit causality verb (e.g., *I admired you because, she admired you because, he admired me because*, etc.). To access a variety of language styles, we used Google for the search, which returned mainly examples of written language. For a full description of sample selection, see Appendix A

Coding of dependent measure: reference continuation

After the following clause was selected, we coded for whether or not the subject or object were mentioned in the following clause. For example, in (1a), the implicit cause *he* is mentioned again in the next clause, while in (1b), the non-cause *he* is re-mentioned.

(1a) I used to talk to Coach Brown all the time, and he amazed me because he has been to a few places

(1b) I surprised him because he thought my list would closely mirror his.

Analysis

Our statistical analyses used SAS proc glimmix to perform a logistic regression, given that our dependent measures were binary. One analysis examined subject references as the dependent measure (is the subject rementioned in the next clause vs. not), and a second analysis examined object references as the dependent measure (is the object rementioned vs. not.) Verb type (centered) was the only predictor in both models (subject-biased vs. object-biased verbs). Models also included random intercepts by verb.

Results: Are implicit causes or non-causes re-mentioned?

Out of all tokens, 40.0% re-mentioned both referents, 57.7% re-mentioned one or the other, and 2.4% re-mentioned neither. Our results revealed a strong tendency to re-mention the implicit cause: people mentioned the subject 87% of the time for subject- biased verbs and 45% of the time for object- biased verbs. In addition, participants mentioned the object 90% of the time for object- biased verbs and 53% of the time for subject-biased verbs. For a breakdown of these patterns by verb, see Supplement. Both of these differences were supported by a main effect of verb type in our models (see Table 1a and Table 1b), suggesting that both subjects and objects are more likely to be re-mentioned when they are the implicit cause. We can also infer both subjects and objects have a likelihood to be re-mentioned above chance (significant intercepts).

Table 1a. Model with subject continuation as dependent measure.

Effect	Estimate (SE)	t value	p value
Intercept	0.75 (0.16)	4.59	0.0003
Verb type	2.10 (0.33)	6.28	<.0001

* Odds ratio for verbttype effect = 8.1

Table 1b. Model with object continuation as dependent measure.

Effect	Estimate (SE)	t value	p value
Intercept	1.28 (0.19)	6.84	<.0001
Verb type	-2.18 (0.37)	-5.96	<.0001

* odds ratio for verbttype effect = 0.11

Text Analysis 2: Fisher Corpus

Goals of analysis

The Google analysis demonstrated that implicit causes tend to be rementioned more than non-causes, but this sample was restricted to sentences with two animate characters and the causal word “because”. This means we can’t tell whether implicit cause-mention is frequent only in this context, or for all instances of these verbs. Our second text analysis tested a less restrictive

Why are implicit causes predictable?¹⁷

sample, including both animate and inanimate referents, and not restricted by the presence of “because”. This analysis used a small sample, because our questions required a high degree of human decision-making and hand coding, but it had the advantage of additionally allowing us to assess the types of referents that typically occur in implicit cause and non-cause roles.

Methods

Sample and Selection Criteria

We analyzed naturally-occurring speech from the Fisher Corpus (Cieri, Graff, Kimball, Miller, & Walker, 2004, 2005), which is a collection of over 16,000 telephone conversations. Our final sample included 198 tokens.³ Similar to the Switchboard corpus, participants in this project were asked to speak about randomly generated topics from a list. Our goal was to understand the frequency with which speakers refer to entities appearing in implicit cause and non-cause roles, so we selected a sample of both subject-biased and object-biased verbs occurring in the transitive frame. For details on selection criteria and examples, see Appendix A.

Goals of Analysis and Coding

We asked two questions. First, to test the frequency account, is the implicit cause more likely to be mentioned again than the non-cause in natural language? We examined this by counting the frequency with which each argument was mentioned in the immediately following clause. For example, in (2a), the implicit cause *people* is mentioned again in the next clause, while in (1b), the non-cause *I* is re-mentioned.

(2a) people just kind of **annoyed me** 'cause they'd be up all the time making noise

(2b) I really **hated it** but then I got older

³ This is comparable to the analysis in Arnold (2001), which included 174 tokens.

Why are implicit causes predictable?18

Given previous evidence that referent predictability is influenced by coherence relation (e.g., Kehler & Rohde, 2013; 2017), we also coded the coherence relation between the utterance describing an implicit cause event and the subsequent utterance as either an **explanation** relation (denoting a causal relationship) or some other, non-explanation relation.

Our second purpose was to test the topicality account by assessing whether implicit causes and non-causes in natural language tend to be associated with other indicators of topicality, namely givenness, pronominalization, prominence on a person hierarchy (1st/2nd vs. 3rd), and animacy (Arnold, 1998, 2010; Arnold, Kaiser, Kahn, & Kim, 2013; Givón, 1983; Prince, 1981). If implicit causes tend to be topical in the discourse, we would expect those referents to score more highly on these measures than the non-cause entities.

In sum, for each token in the database, we hand-coded three sets of information. First, our primary dependent measure was repeated reference. In the utterance following the implicit cause verb, we coded two dependent measures: 1) whether the subject was rementioned or not, and 2) whether the object was rementioned or not. Second, we coded the discourse properties of the subject and object of the implicit causality verb. Third, we coded the coherence relation between the two clauses. For details on coding, see Appendix B.

Statistical analysis

Following the same procedure as for the Fisher analysis, we used SAS proc glimmix to perform a logistic regression. The predictor verb type (subject-biased vs. object-biased) was centered. The verb was the random intercept in all models, except where the model estimated it to be zero, in which case it was excluded. We ran two models, one with subject re-mention as the dependent measure and one with object re-mention as the dependent measure.

Results and Discussion

Question 1. Who gets mentioned again?

The first question is whether implicit cause are more likely than implicit non-cause entities to be mentioned in the immediately following clause. Figure 1 shows that this is not the case. There was a small numerical preference for the opposite pattern, i.e. for non-cause roles to be re-mentioned for both subject and object positions. However, this was not significant in the logistic regressions. Verb type (subject biased vs. object biased) had no effect on the rate of re-mention of either the subject ($\beta = -0.71$ (0.4), $t=-1.77$, $p=0.1$, odds ratio = 0.49) or the object ($\beta = 0.2$ (0.29), $t=0.69$, $p=0.49$, odds ratio = 1.2). For a breakdown of these patterns by verb, see Supplement.

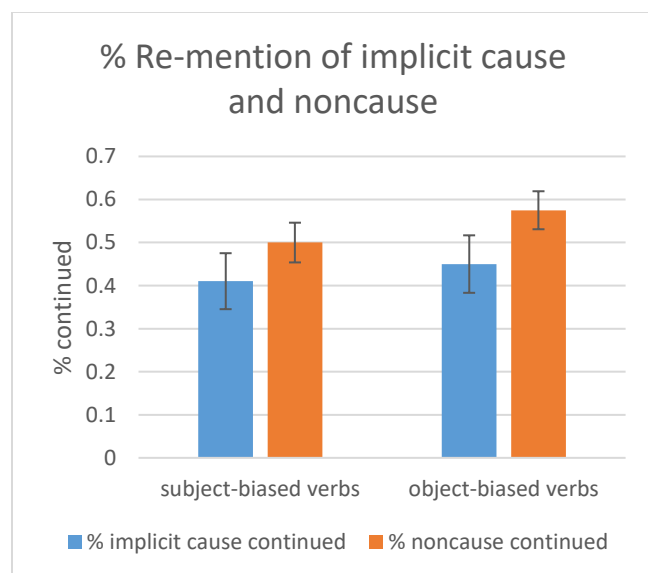


Figure 1. Fisher corpus results: Rate of continued mention of implicit cause and non-cause entities. Error bars represent the standard error of continuation rates for each verb. The percentage continued reflects the percentage of tokens in which the cause or non-cause is mentioned in the immediately next clause, either directly or indirectly. Note that each percentage is calculated out of the total number of tokens for that verb. Because the speaker could mention

both, none, or one of the arguments in each sentence, the bars for cause and non-cause do not add up to 100%.

Given the starkly different pattern between this analysis and the Google analysis, we performed a post-hoc comparison by analyzing the two samples together, including predictors verb type (centered), experiment (centered), and the interaction between the two. For subject re-mention there was a significant effect of sample ($\beta = 0.73$ (SE = 0.24), $t = 3.02$, $p = 0.006$), no effect of verb type ($p = 0.82$), and an interaction between sample and verbytype ($\beta = 2.7$ (SE = 0.50), $t = 5.52$, $p < .0001$). For object re-mention there was a significant effect of sample ($\beta = 1.4$ (SE = 0.20), $t = 6.9$, $p < .0001$), a marginal effect of verbytype ($\beta = -0.59$ (SE = 0.32), $t = -1.83$, $p = 0.08$) and a significant interaction ($\beta = -2.22$ (SE = 0.41), $t = -5.39$, $p < .0001$). This confirms that verbytype had significantly different effects on remention patterns for our two samples.

Question 2: Do coherence relations modulate re-mention patterns?

Several theories suggest that semantically-based referential predictability is conditioned on the coherence relations between utterances (Arnold, 2001; Kehler, et al., 2008; Kehler & Rohde, 2013). In particular, emotion verbs are claimed to make the implicit cause predictable when the next sentence describes the cause of the event (i.e., Explanation continuations). This predicts that implicit cause referents might only be preferentially continued in supportive coherence contexts, or perhaps that the pattern would be stronger in these contexts.

However, testing this question with the Fisher sample is difficult, because the frequency of explanation relations was very low in our sample, representing only 10% of the object-biased verbs ($n = 12$) and only 10.2% of the subject-biased verbs ($n=8$). In and of itself this is notable, because it shows that emotion verbs frequently occur without a following explanation relation. This pattern might help explain the contrast between the Fisher and Google analyses, since all

Why are implicit causes predictable?21

the target items in the Google sample were followed by an explanation, whereas almost none were in the Fisher analysis. Nonetheless, the referents in these two samples also differed in animacy and discourse status, making it difficult to pinpoint the critical source of difference.

Question 3: Do implicit causes have a more topical discourse status than non-causes?

Our next question is whether cause and non-cause roles tend to be correlated with topical discourse status characteristics. We tested this question by examining each of four discourse properties: animacy, 1st and 2nd person vs. 3rd; givenness, and pronominalization. We asked whether verb type (subject-biased vs. object-biased) predicted the rate of each property occurring, running separate models for subjects and objects. Thus, for each model the predictor was verb type (centered), and the dependent measure was the discourse property. For example, the subject models tested whether verbytype significantly modulated the likelihood that the referent in subject position was animate; 1st/2nd person, pronominal/zero; or given/inferable. Table 2 reports the effect of verb type for each discourse property, with all significant effects bolded.

Table 2. The discourse properties of subjects and objects in emotion verbs, depending on their thematic role, plus the inferential statistics from separate logistic regression models for each property, tested separately for subjects and objects.

Dependent measure	implicit causes (Subj-biased verbs)	non-causes (Obj-biased verbs)	Subject model, verb type effect
Animate (vs. Inanimate) Subject	27%	100%	$\beta = -17.56$ (361.15), $t = -0.05$, $p = 0.96$, odds ratio = .000000235) Pearson's chi- square: 123.14; p < .001
1st or 2nd person (vs. 3rd) Subject	9%	68%	$\beta = -3.09$ (0.55), $t = -5.62$, $p < .0001$, odds ratio = .05
Pronominal or zero (vs. nominal) Subject	85%	92%	$\beta = -0.69$ (0.46), $t = -1.52$, odds ratio = 0.5, $p = 0.13$
Given or inferrable (vs. new) Subject	81%	93%	$\beta = -1.1$ (0.5), $t = -2.17$, odds ratio = .33, $p = 0.05$

- NOTE: models included a random slope for verb, except in the pronominal model where it was estimated to be zero and in the animacy model where it prevented convergence.
- NOTE: For the animacy model, the logistic regression estimated the standard error to be extremely high, presumably due to the fact that there was no variation for the object-biased verbs, which rendered the difference between conditions insignificant. We therefore tested the effect of animacy with Pearson's chi-square, which showed that the difference between 100% and 27% subject-continuation was indeed significant.

Dependent measure	implicit causes (Obj-biased verbs)	non-causes (Subj-biased verbs)	Object model; verb type effect
Animate (vs. Inanimate) Objects	33%	99%	$\beta = 5.54$ (1.5), $t = 3.7$, odds ratio = 254, $p = 0.0013$
1st or 2nd person (vs. 3rd) Objects	8%	68%	$\beta = 3.2$ (0.66), $t = 4.86$, odds ratio = 24, $p = 0.0004$
Pronominal or zero (vs. nominal) Objects	58%	87%	$\beta = 1.57$ (0.43), $t = 3.65$, odds ratio = 4.8, $p = 0.003$
Given or inferrable (vs. new) Objects	75%	85%	$\beta = 0.46$ (0.55), $t = 0.83$, odds ratio = 1.5, $p = 0.43$

- NOTE: all models included a random slope for verb

Notably, our findings were all in the opposite direction that was predicted. The “predictable” referent, the implicit cause role, was actually less likely to occur with prominent discourse features than the non-cause role. The non-cause was numerically more prominent in all eight comparisons, and this difference reached significance in all but two. These analyses are possible because of the painstakingly detailed hand coding performed for this analysis, but by the same token this coding limited the sample size. Thus, these findings must be interpreted with caution.

Nevertheless, these analyses are clearly inconsistent with the hypothesis that predictability patterns with discourse prominence, which might be predicted by the topicality account for re-mention patterns. But it makes sense if we think about the meaning of emotion verbs, which describes the mental state of the non-cause entity, who often has the thematic role of experiencer. This requires the speaker to take the perspective of the non-cause entity, which seems more likely to occur when that character is central to the discourse. That is, speakers are more likely to use emotion verbs to talk about the emotions of topical characters than less-topical characters. In addition, typically only humans can be experiencers, and humanness is correlated with two of our discourse prominence metrics (animacy and 1st/2nd person; see also Corrigan, 1992). This question will be further explored in Experiment 2.

Text analyses discussion

The primary question in the text analyses was whether there was any support for a frequency account of implicit cause expectation. The Google analysis targeted the kind of sentences that are frequently used for psycholinguistic experiments, namely sentences with two animate arguments and followed by the connector *because*. In this case, we found that implicit causes were significantly more likely to be rementioned. This provides support for the

Why are implicit causes predictable?²⁴

semantically-based frequency account, in that we observed that implicit causes, which are perceived to be predictable, are also likely to be re-mentioned in natural language. However, by contrast, the Fisher corpus analysis examined a random selection of emotion verbs in natural conversation, and verb type had no effect on remention patterns in this context. Our combined analysis confirmed that verb type had different effects across experiments. While there were several differences in the tokens across corpora, a major difference was that the Google corpus always included a following explanation, whereas the Fisher corpus almost never did. This is consistent with the **coherence-contingent** version of the frequency account, and suggests that learning about implicit cause expectancy might require tracking statistics of re-mention contingent on both verb semantics and coherence relations.

These findings are consistent with a corpus analysis of newspaper text (Long and de Ley, 2000). Long and de Ley did not test referential re-mention explicitly, but instead analyzed the reverse question, “Given an anaphor in subject position that refers to a referent from a preceding implicit causality verb with two arguments, which referent is it more likely to refer to?” While their selection criteria were somewhat different, they found that for their object-biased verbs, anaphors were more likely to re-mention the implicit cause overall (70% of all reported samples in active voice), but not for their subject-biased verbs (47% of all reported samples in active voice referred to the implicit cause). But this rate of implicit-cause mention was much higher for the subset of items with the connective “because” (96% for object-biased verbs and 74% for subject-biased verbs). This further supports the conclusion that the frequency-based account of implicit causality expectation is possible, but only in the coherence-contingent version. In addition, people may also track frequency as a function of other semantic features, for example

Why are implicit causes predictable?25

only when the two referents are equal in discourse prominence (as they were in the Google analysis, where both arguments were pronominalized).

The Fisher corpus also addressed our question about whether implicit causes would tend to occur in topical or discourse prominent positions. In short, they do not. Instead, most implicit causes in natural language are inanimate, and inanimate roles tend to not be topical. By contrast, the non-cause referent (the experiencer role) tends to be human and topical.

These text analyses provide evidence of how implicit causality events occur in natural speech, both the types of referents that tend to occur as implicit cause and non-cause arguments, and which ones tend to be re-mentioned. Our next question is whether people use implicit causality to guide judgments of both predictability and topicality in controlled discourses.

Experiments 1 and 2

Our experiments sought to directly test how people view implicit causes and non-causes in terms of predictability and topicality. Some studies have used pronoun production to argue that implicit causes are not topical, based on evidence that speakers tend to use pronouns to refer to the subject irrespective of implicit causality (e.g., Fukumura & van Gompel, 2010; Rohde & Kehler, 2014). If topicality guides pronoun production, this might imply that implicit causality is not related to topicality. Yet this story is complicated by the fact that other studies found that implicit causality can indeed affect pronoun production (Weatherford & Arnold, 2021; Ye, Weatherford, & Arnold, 2021).

Here we instead use a metacognitive judgment task to test predictability and topicality in parallel. Our task differed from many previous tests of implicit causality in that we provide a

context sentence to facilitate representations of the critical emotion event.⁴ Participants read stories like *Liz and Ana were volunteering at the library. Liz offended Ana (because)....* and then were asked to judge either 1) who is likely to be mentioned next (predictability) or 2) who is the main character (topicality). To assess the importance of the coherence relation to both judgments, we manipulated the presence of the word *because*. Experiment 1 tested the predictability question, and Experiment 2 tested the topicality question. The two experiments are discussed together because the stimuli were otherwise identical.

Methods

Participants

There were 32 participants in Exp. 1 and 32 in Exp. 2. All were native speakers of English residing in the United States or the United Kingdom and were recruited on Amazon Mechanical Turk in exchange for \$1.00. Another restriction was that they had not previously completed another experiment from the Arnold Lab. Ages ranged from 24 to 76 years old.

Materials and Design

We used four characters in these stories who had simple and common English names: Ana and Liz, who were explicitly described as female, and Will and Matt, who were explicitly described as male. For each target item, the story consisted of one context sentence in the past progressive tense followed by an incomplete sentence containing an implicit causality verb in the simple past tense. The context sentence served to introduce two characters and the setting. For example, *Liz and Ana were volunteering at the library. {Liz offended Ana... / Liz disliked Ana...}*

⁴ But see for example Koorneef & Sanders, 2013, Experiment 2; Majid, Sanford, & Pickering, 2006, and van den Hoven & Ferstl, 2018 for a similar stimulus feature.

Why are implicit causes predictable?27

Each story was manipulated to appear in one of two conditions such that the second clause included either a subject-biased (e.g., *offend*) or object-biased (e.g., *dislike*) verb. We also manipulated the presence of *because* at the end of the second clause. Verb type (subject-biased versus object-biased) was manipulated within subjects. Connective type (*because* versus null) was manipulated between subjects because reading “because” in one stimulus would be likely to activate the causal relation for other stimuli as well.

We tested a total of 12 critical stories. Each story used two verbs (one in each verbtype condition), so in total we examined 24 different emotion verbs. To keep the task short, each participant only saw six of the critical items, together with six filler items (i.e. 12 items total), plus two practice items. The inclusion of *because* was manipulated between-subjects. Thus, each experiment used a 2 (verbtype: subject-biased vs. object-biased) x 2 (with *because* vs. without *because*) design; this design was applied to two sets of verbs, resulting in 8 lists per experiment. For each experiment, there were 16 participants in the with-*because* condition and 16 participants in the without-*because* condition, yielding a total of 192 datapoints per experiment.

Subject-biased verbs were taken from Levin class 31.1, and object-biased verbs were taken from Levin class 31.2, again with the exception of “tease” which is class 31.1 but is an object-biased verb (Hartshorne & Snedeker, 2013). Across all verbs, implicit cause bias scores ranged from 0.73 to 0.89, based on Hartshorne and Snedeker (2013).

Filler and practice items followed a similar format. For all filler items, the context sentence was in the past progressive, and the following fragment had a non-emotion verb in the simple past tense. For variety, two filler items and one practice introduced two characters and continued talking about both, two filler items and one practice introduced two characters and continued talking about only one of them (represented in Table 3), and 2 filler items introduced

Why are implicit causes predictable?28

only one character and continued talking about him/her. Example items are given in Table 3. The complete list of items is given in the Supplement.

Table 3: Example experimental stimuli.

Item Type	Item
Practice	Will and Matt were doing the laundry. Will folded the clothes with Matt...
Subject-biased	Matt and Will were working out at the gym. Matt aggravated Will...
Object-biased	Liz and Ana were volunteering at the library. Liz disliked Ana...
Filler	Matt and Will were eating breakfast. Will took out the cereal...

For measuring predictability and topicality, respectively, we had two critical post-story questions: *Think about the **rest of the second sentence** in this story. Who do you think will be mentioned?* (Exp. 1) and *Who do you think is the **main character** of this story?* (Exp. 2). Both were given as 2-alternative forced choice questions. Each story also had a content question, which asked about either who was involved in the story (e.g. *Who was working on a project for class?*) or what happened in the story (e.g. *What were they doing?*). These extremely simple questions served as a check that participants were engaged and actually processing the stories. Participants were informed that they would not be paid if they had too many errors; if a participant answered more than 25% of the content questions incorrectly, the survey automatically ended, and we did not use their data.

In sum, from the participants' perspective, in one survey they saw 2 practice items, and 6 target items (3 subject-biased verbs, 3 object-biased verbs) intermixed with 6 filler items. For each item, they answered the content question on one page followed by the content question and one of the two critical questions on the next page.

Data Analysis

We performed logistic regressions using SAS proc glimmix. Fixed effects included verb type and connective type, which were coded with zero-sum contrasts. The dependent measure was whether the participant chose the subject character or not. We constructed mixed-effects logistic regression models with maximal random effects except when estimated to be zero by the model (random intercepts for both participants and items, random slopes for verb type by both participants and items, and a random slope for connective and the interaction between connective and verbytype by items).

Results and Discussion

For our question about predictability (Exp. 1), participants selected the subject more when it was the implicit cause than when it was not (Figure 1). The same pattern emerged for both the with-because and without-because conditions, but the difference was bigger in the with-because condition. This was supported by a model with subject chosen as the dependent measure, given in Table 4a below, in which there was a main effect of verb type and an interaction between verbytype and causal coherence. We further probed the interaction by estimating the effect of verb type in each condition (with-because and without-because). There was a significant effect of verbytype with because ($\beta = 2.99$ (SE = .64), $t = 4.65$, $p < .0001$) but only a marginal effect of verbytype without because ($\beta = 0.98$ (SE = 0.54), $t = 1.83$, $p = .0812$).

Table 4a. Exp. 1: Model of predictability with subject chosen as dependent measure.

Table 4b. Exp. 2: Model of topicality with subject chosen as dependent measure.

Effect	Estimate (SE)	t value	p value
Intercept	0.25 (0.36)	0.70	0.49
Verb type	1.99 (0.42)	4.73	<.0001
Causal coherence	-0.35 (0.65)	-0.53	0.60
Verb type x Causal coherence	1.99 (0.84)	2.36	0.02

* odds ratio for the verbytype effect in the with-because condition: 19.8, and in the without-because condition: 2.7

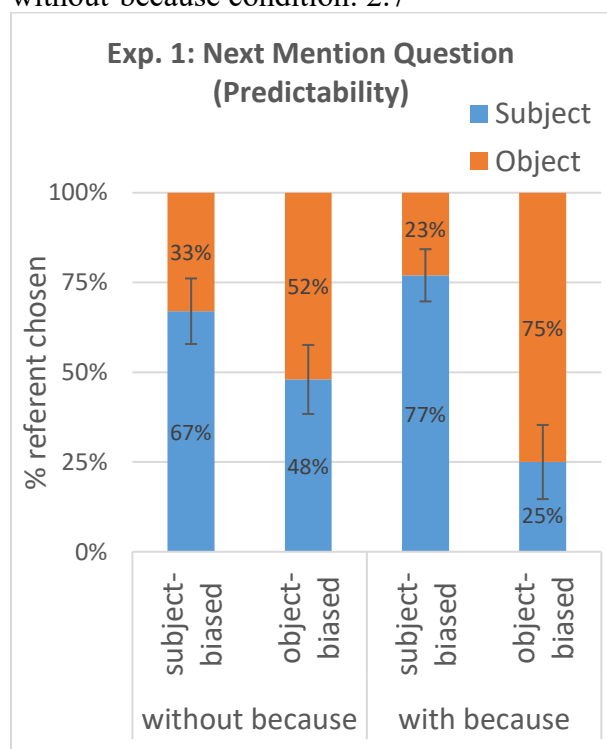


Figure 2: Experiment 1 results, Next Mention question. The standard error of the subject means in each condition are plotted.

Effect	Estimate (SE)	t value	p value
Intercept	1.73 (0.35)	4.90	0.0002
Verb type	-0.20 (0.54)	-0.37	0.72
Causal coherence	-0.57 (0.57)	-0.99	0.33
Verb type x Causal coherence	-0.20 (1.12)	-0.18	0.86

* odds ratio for the verbytype effect in the with-because condition: 0.7, and in the without-because condition: 0.9.

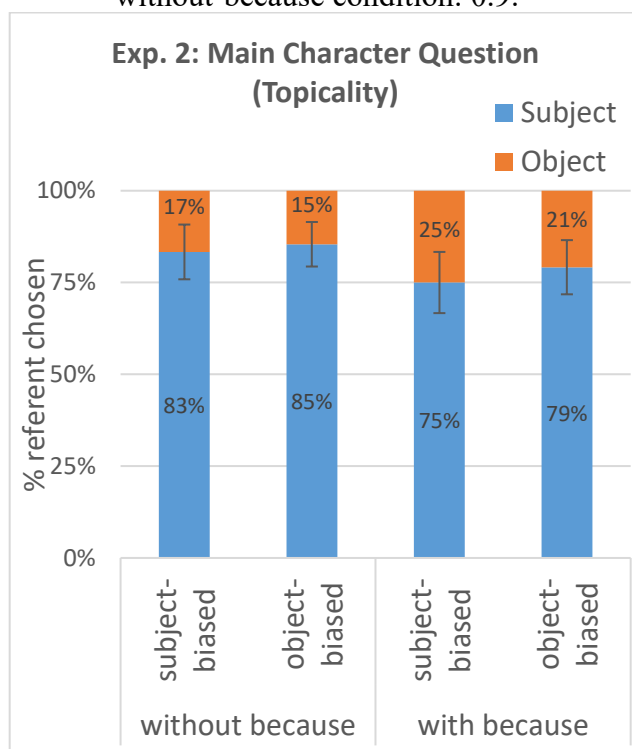


Figure 3: Experiment 2 results, Main Character question. The standard error of the subject means in each condition are plotted.

For our question about topicality (Exp. 2), people consistently chose the subject whether it was the implicit cause or not in both connective conditions (Figure 3). This observation was supported by a model with subject chosen as the dependent measure, given in Table 4b above.

Why are implicit causes predictable?31

The only significant factor was the intercept ($p < 0.001$), indicating that subject bias was the only predictor for judging topicality.

The robustness of our findings are supported from three sources. First, to confirm the contrast between predictability and topicality questions, we performed a post-hoc analysis to compare the experiments (see Appendix C). This analysis revealed a significant interaction between verb type and question type, and estimates confirmed that the verb type effect was significant for only the predictability and not the topicality questions. Second, we used a power simulation to test whether a sample of this size would have adequate power to detect the verb type effect for experiment 1 and the intercept effect for experiment 2. In an analysis of about 1000 simulated datasets, we found that it did (see Appendix D). Third, these findings are supported by conceptually similar patterns detected in other studies. A similar effect of verb type for the predictability question was reported in two other studies using implicit causality verbs in a similar task but with different stories, and only in the condition with “because” present (Weatherford & Arnold, 2021; Johnson & Arnold, in press), providing a direct replication for one condition from Exp. 1. In addition, a similar pattern for the topicality question was observed for the same task with a different verb type (Zerkle & Arnold, 2019).

The results from our experiments suggest that implicit causality affects predictability but not topicality judgments. These findings replicate previous findings with the sentence-completion method, which have shown that speakers tend to frequently continue talking about the implicit cause, supporting the hypothesis that the implicit cause is predictable (Fukumura & van Gompel, 2010; Kehler et al., 2008; Kehler & Rohde, 2013; Holler & Suckow, 2016). We also found that subjecthood affects topicality but not predictability judgments for implicit

Why are implicit causes predictable?³²

causality verbs. Our findings are consistent with the suggestion that predictability and topicality are represented separately (e.g., Kehler & Rohde, 2013).

General Discussion

This project examined implicit causality scenarios for the purpose of better understanding where referential predictability comes from, and the relationship between referential predictability and topicality. There were two major findings from this project: 1) implicit causes are frequently mentioned but only in some conditions, and 2) predictability and topicality are not correlated for this verb type. We take up each of these points in turn, and then discuss their implications for how predictability is calculated.

Does referential predictability stem from frequency?

Our first question was whether implicit causes have a higher re-mention frequency in natural language compared to non-causes. We found that implicit causes were frequently re-mentioned in the Google sample, which was restricted to contexts with two animate arguments and a clear explanation in the following clause. By contrast, when we considered all emotion verbs with two arguments (the Fisher corpus), there was no preference to re-mention the implicit cause. This difference could stem from three differences between our analyses. First, the Google analysis was restricted to cases including an explanation in the second clause, while the Fischer analysis had very few tokens of this type. Second, in the Fisher corpus, the implicit cause argument was frequently an inanimate referent (73% in object-biased verbs, and 67% in subject-biased verbs), whereas it was always animate in the Google analysis. We know that people like

Why are implicit causes predictable?33

to talk about people, so it may not be surprising that implicit causes are more often mentioned when they are animate.

Third, the discourse contexts in our two corpora are very different, and the context can have severe consequences for the kinds of explanations that would be appropriate. Explanations are expected to provide newsworthy information, but the implicit cause is not newsworthy if the cause has already been explained (see Kehler & Rohde, 2018, Solstad & Bott, 2013 for a similar idea). For example, consider this excerpt from the Fisher corpus: *her breath smells bad her hair her hands smell so bad and I noticed it a lot because I'm a nonsmoker*. Here, the context already provides an explanation for why the smell would be noticeable, making the implicit cause (*it* = the smell) less likely to be mentioned. Instead, the knowledge of the speaker's smoking status is unknown, so this information is more newsworthy, so the non-cause is re-mentioned. While we did not code for the status of the explanation in either corpus, this highlights the fact that the “right” sort of explanation is a subset of all explanations.

Together, our findings provide partial support for the frequency account, which suggests that people could learn through experience that implicit causes tend to be rementioned by observing that this pattern frequently holds in discourse. However, such learning would require learners to track frequencies on a detailed level, paying attention not just to the overall frequency of re-mention for all implicit causes, but rather to the frequency of re-mention for specific contexts. This suggests that comprehenders would have to filter natural language, raising questions about what critical constraints are used for this filtering, and whether it is computationally plausible. Many researchers refer to this as the “grain problem” (Desmet & Gibson, 2003; Desmet, De Baecke, Drieghe, Brysbaert, & Vonk, 2006; Mitchell, Cuetos, Corley,

Why are implicit causes predictable?³⁴

& Brysbaert, 1995; Saffran, 2002; Townsend & Bever, 2001). A language learner may have to categorize tokens along several dimensions.

Evidence that listeners might do this comes from the domain of syntactic parsing, where some researchers believe that comprehenders are capable of the kind of fine-grained computation necessary for reconciling corpus patterns with online comprehension results. Desmet and colleagues conducted a series of experiments involving relative clause attachment in Dutch. They argue that if natural language tokens are filtered for animacy and concreteness of referents, frequencies in corpus analyses do match the attachment biases found in online comprehension and sentence completion experiments. They take this as evidence that experience-based accounts of processing are plausible (Desmet, Brysbaert, & De Baecke, 2002, Desmet & Gibson, 2003; Desmet et al., 2006; see also Jurafsky, 1996; MacDonald, 2013; MacDonald & Thornton, 2009; Saffran, 2002, 2003; Tabor, Juliano, & Tanenhaus, 1997; Wells, Christiansen, Race, Acheson, & MacDonald, 2009). Their results suggest that by extension, people may also track referential frequencies for subsets of contexts. However, more work is needed to identify which features of the discourse context correlate with next-mention probability for different verb types, and test whether these features also guide predictability judgments during language comprehension.

A vast amount of research has shown that for reference frequency, coherence relations are one dimension that is likely to be relevant. Our text analyses suggest that animacy of the referents may be another dimension (see also Corrigan, 1992). Finally, the comprehender may have to track whether the explanation for the event has not been previously provided. In sum, the corpora results suggest that natural re-mention frequency might help listeners learn that implicit causes are likely to be re-mentioned, but only if comprehenders are able to constrain their learning along several dimensions.

Does referential predictability correlate with topicality?

Our second question asked how referential predictability relates to topicality. This question is theoretically relevant because some models consider predictability and topicality as overlapping properties (Givón, 1983), while others consider them as distinct properties (Kehler & Rohde, 2013). In the Fisher corpus analysis, we examined discourse properties that are associated with discourse topicality and accessibility, namely animacy, 1st/2nd person (vs. 3rd), givenness, and pronominalization. The implicit cause did not pattern with topical discourse properties, and instead, the non-cause role did. We then used a metalinguistic task in Experiments 1 and 2 to collect parallel judgments of predictability and topicality and found similar results. Participants overwhelmingly selected the grammatical subject as the most topical referent, but chose the implicit cause as the most predictable referent. The Fisher findings could result from the fact that most implicit causes were non-human in that corpus, but this explanation does not account for the experimental results, where both characters were human and given.

Our findings suggest that for implicit causality verbs, people tend to judge the implicit cause as predictable (likely to be mentioned), and implicit cause status has no effect on topicality judgments. This conclusion differs from findings for other verbs where the same entity is viewed as both topical and predictable (Zerkle & Arnold, 2019). However, other scholars have come to the same conclusion from sentence completion studies (e.g., Kehler and Rohde (2013; Kehler et al., 2008; Rohde & Kehler, 2014). Although sentence completion studies only provides indirect information about how topicality relates to implicit causality, the fact that both methods provide converging evidence suggests that implicit causality affects predictability but not topicality.

Conclusion

In sum, our study contributes to the study of referential predictability by testing two hypotheses about how people might calculate that implicit causes are predictable. The frequency hypothesis suggested that people might learn that implicit causes are highly likely to be mentioned again through experience, if implicit causes are frequently re-mentioned overall in natural language. We found that this idea is only plausible for emotion verbs if people track re-mention frequencies for very specific contexts.

The topicality hypothesis suggested that people might assign predictability on the basis of topicality, or conversely that they might perceive predictable referents as topical. Again, while this idea is plausible for some discourse contexts, such as the predictability of subjects (Zerkle & Arnold, 2019), it does not account for the predictability of implicit causes. In our spoken-language corpus analysis, implicit causes did not tend to fall in discourse roles associated with discourse topicality. In our experiments, topicality judgments were not influenced by implicit causality, whereas predictability judgments were. In addition, the presence of the word *because* strengthened predictability judgments, but had no effect on topicality judgments. One reason that topicality is aligned with predictability in the literature is that previous topics often continue as the topic of upcoming speech input. This means that the previous topic is likely to be re-mentioned. We saw here that in implicit causality contexts, this pattern does not hold.

The current study did not directly test the semantic inference account (e.g., Hartshorne et al., 2015), and this remains a plausible mechanism for calculating predictability. The intuition behind implicit causality is that the assumed implicit cause is dependent on the meaning of the event. Together with expectations about an upcoming explanation, this may be enough to generate a prediction that the implicit cause will be mentioned.

At the same time, our results raise the possibility that referential frequencies may contribute to expectations about upcoming language input. We observed that for emotion verbs, the frequency account only can explain referential expectations if we assume a fine-grained learning mechanism, where people integrate knowledge about the semantic roles of discourse participants and the expectation for an upcoming explanation, perhaps based on whether the prior discourse already provided an explanation or not. Animacy may matter as well. Yet if people must remember such a rich representation of contingent frequencies, it raises questions about the difference between a frequency-driven mechanism and a semantic inference mechanism. Are semantic inferences merely complex representations of referential frequencies?

For example, consider the sentence like *Ana impressed Liz because....* Under the frequency account, it might evoke a memory that in scenarios with emotion verbs, two animate characters and an upcoming explanation, mention of the subject (Ana) is frequent. Under the semantic inference account, this sentence might lead to a semantic representation that generates an inference about the likely cause (Ana) which combined with an expectation for an upcoming event leads to an expectation for re-mention of that person. Both accounts require representation of the same sets of conditioning features, but the frequency account makes predictions based on what has happened in the past, whereas the inference account makes predictions based on what is plausible. These calculations may be very similar.

In conclusion, our study has shown that implicit causality tends to be unrelated to judgments of topicality and its relation to frequency of reference in natural language is modulated by constraints on the context. These findings highlight the importance of manipulating the context in experimental studies, and of examining real-world discourse contexts in natural speech. A full theory of referential predictability will require an understanding of how

Why are implicit causes predictable?38

multiple different verb types relate to both frequency of mention and other discourse indicators of topicality.

Author note

Thank you to Irene Tang for her work on the Google analysis, and to Michaela Neely, Grant Huffman, Elise Rosa, Simon Wolf, Leela Rao, and Ana Medina Fetterman for their work on the Fisher corpus analysis. This work was partially supported by NSF grants 1651000 and 1348549 to J. Arnold, and by the 2018 Swarthmore Career Services Summer Experiential Fellowship to S. Guan. We are grateful to UNC's statistical consulting services for assistance with our models.

References

- Altmann, G., & Kamide, Y. (1999). Incremental interpretation at verbs: Restricting the domain of subsequent reference. *Cognition*, 73, 247–264. doi:10.1016/S0010-0277(99)00059-1
- Ariel, M. (1990). *Accessing Noun-Phrase Antecedents*. London: Routledge.
- Arnold, J. E. (1998). Reference form and discourse patterns. Doctoral dissertation, Stanford University.
- Arnold, J. E. (2001). The effect of thematic roles on pronoun use and frequency of reference continuation. *Discourse Processes*, 3, 137-162. doi:10.1207/S15326950DP3102_02
- Arnold, J. E. (2010). How speakers refer: The role of accessibility. *Language and Linguistics Compass*, 4, 187-203. doi:10.1111/j.1749-818X.2010.00193.x
- Arnold, J. E. (2011). Ordering choices in production: For the speaker or for the listener. In E. M. Bender & J. E. Arnold (Eds.), *Language from a cognitive perspective: Grammar, usage, and processing*, 199-222. CSLI Publishers.

Why are implicit causes predictable?39

Arnold J. E. (2017). Corpus analysis of transfer verbs. Unpublished analysis, UNC Chapel Hill.

Arnold, J. E., Hudson Kam, C., & Tanenhaus, M. K. (2007). If you say thee uh you are describing something hard: The on-line attribution of disfluency during reference comprehension. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 33, 914-930. doi:10.1037/0278-7393.33.5.914

Arnold, J. E., Kaiser, E., Kahn, J. M., & Kim, L. K. (2013). Information structure: linguistic, cognitive, and processing approaches. *Wiley Interdisciplinary Reviews: Cognitive Science*, 4, 403-413. doi:10.1002/wcs.1234

Arnold, J. E., & Lao, S. C. (2008). Put in last position something previously unmentioned: Word order effects on referential expectancy and reference comprehension. *Language and Cognitive Processes*, 23, 282-295. doi:10.1080/01690960701536805

Arnold, J. E., Strangmann, I., Hwang, H., & Zerkle, S. (2018). Reference frequency: what do speakers tend to talk about? Technical Report #2. UNC Language Processing Lab, Department of Psychology & Neuroscience, University of North Carolina – Chapel Hill, Chapel Hill, North Carolina.

Arnold, J. E., Wasow, T., Losongco, T., & Ginstrom, R. (2000). Heaviness vs. newness: The effects of structural complexity and discourse status on constituent ordering. *Language*, 76, 28-55. doi:10.1353/lan.2000.0045

Au, T. K. F. (1986). A verb is worth a thousand words: The causes and consequences of interpersonal events implicit in language. *Journal of Memory and Language*, 25, 104-122. doi:10.1016/0749-596X(86)90024-0

Why are implicit causes predictable?40

Bock, J. K., & Irwin, D. E. (1980). Syntactic effects of information availability in sentence production. *Journal of Verbal Learning & Verbal Behavior*, 19, 467-484.

doi:10.1016/S0022-5371(80)90321-7

Bower, Gordon & Morrow, Daniel. (1990). Mental Models in Narrative Comprehension. *Science*, 247, 44-8. doi:10.1126/science.2403694.

Bransford, J. D., Barclay, J. R., & Franks, J. J. (1972). Sentence memory: A constructive versus interpretive approach. *Cognitive Psychology*, 3, 193-209. [doi.org/10.1016/0010-0285\(72\)90003-5](https://doi.org/10.1016/0010-0285(72)90003-5)

Brennan, S. E., Friedman, M. W., & Pollard, C. J. (1987). A centering approach to pronouns. In *Proceedings from the 25th Annual Meeting of the Association for Computational Linguistics*, pp. 155-162, Stanford, CA.

Brocher, A., Chiriacescu, S. I., & von Heusinger, K. (2018). Effects of information status and uniqueness status on referent management in discourse comprehension and planning. *Discourse Processes*, 55(4), 346–370. <https://doi-org.libproxy.lib.unc.edu/10.1080/0163853X.2016.1254990>

Brown, R., & Fish, D. (1983). The psychological causality implicit in language. *Cognition*, 14, 237–273. doi:10.1016/0010-0277(83)90006-9

Caramazza, A., Grober, E., Garvey, C., & Yates, J. (1977). Comprehension of anaphoric pronouns. *Journal of Verbal Learning & Verbal Behavior*, 16, 601-609. doi:10.1016/S0022-5371(77)80022-4

Chafe, W. (1976). Givenness, contrastiveness, definiteness, subjects, topics, and point of view. In Charles N. Li (Ed.), *Subject and Topic*, 25-56. New York: Academic Press Inc.

Chafe, W. (1994). *Discourse, Consciousness, and Time*. Chicago: Chicago University Press.

Why are implicit causes predictable?41

Cieri, C., Graff, D., Kimball, O., Miller, D., & Walker, K. (2004). *Fisher English Training*

Speech Part 1 Transcripts LDC2004T19. Web Download. Philadelphia: Linguistic Data Consortium, 2004.

Cieri, C., Graff, D., Kimball, O., Miller, D., & Walker, K. (2005). *Fisher English Training Part*

2, Transcripts LDC2005T19. Web Download. Philadelphia: Linguistic Data Consortium, 2005.

Corrigan, R. (1992). The relationship between causal attributions and judgments of the typicality of events described by sentences. *British Journal of Social Psychology*, *31*, 351-368.

Coulson, S., Federmeier, K. D., Van Petten, C., & Kutas, M. (2005). Right hemisphere

sensitivity to word and sentence level context: Evidence from event-related brain potentials. *Journal of Experimental Psychology: Learning, Memory and Cognition*, *31*, 129-147. doi:10.1037/0278-7393.31.1.129

Cozijn, R., Commandeur, E., Vonk, W., & Noordman, L. G. (2011). The time course of the use of implicit causality information in the processing of pronouns: A visual world paradigm study. *Journal of Memory and Language*, *64*, 381-403. doi:10.1016/j.jml.2011.01.001

Crinean, M., & Garnham, A. (2006). Implicit causality, implicit consequentiality and semantic roles. *Language and Cognitive Processes*, *21*, 636-648. doi:10.1080/01690960500199763

Delogu, F., Jachmann, T., Staudte, M., Vespignani, F., & Molinaro, N. (2020). Discourse expectations are sensitive to the question under discussion: Evidence from ERPs.

Discourse Processes, *57*(2), 122–140. <https://doi-org.libproxy.lib.unc.edu/10.1080/0163853X.2019.1575140>

Why are implicit causes predictable?42

Desmet, T., Brysbaert, M., & De Baecke, C. (2002). The correspondence between sentence production and corpus frequencies in modifier attachment. *Quarterly Journal of Experimental Psychology*, 55A, 879–896. doi:10.1080/02724980143000604

Desmet, T., De Baecke, C., Drieghe, D., Brysbaert, M., & Vonk, W. (2006). Relative clause attachment in Dutch: On-line comprehension corresponds to corpus frequencies when lexical variables are taken into account. *Language and Cognitive Processes*, 21, 453-485. doi:10.1080/01690960400023485

Desmet, T., & Gibson, E. (2003). Disambiguation preferences and corpus frequencies in noun phrase conjunction. *Journal of Memory and Language*, 49, 353-374. doi:10.1016/S0749-596X(03)00025-1

Ehrlich, K. (1980). Comprehension of pronouns. *Journal of Experimental Psychology*, 32, 247-255. doi:10.1080/14640748008401161

Falkauskas, K., & Kuperman, V. (2015). When experience meets language statistics: Individual variability in processing English compound words. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 41, 1607-1627. doi:10.1037/xlm0000132

Federmeier, K. D., & Kutas, M. (2001). Meaning and modality: Influences of context, semantic memory organization, and perceptual predictability on picture processing. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 27, 202-224. doi:10.1037/0278-7393.27.1.202

Fukumura, K., van Gompel, R.P.G. (2010). Choosing anaphoric expression: Do people take into account likelihood of reference? *Journal of Memory and Language*, 62, 52-66. doi:10.1016/j.jml.2009.09.001

Why are implicit causes predictable?43

Garnham, A., Traxler, M., Oakhill, J., & Gernsbacher, M A. (1996). The Locus of Implicit Causality Effects in Comprehension. *Journal of Memory and Language*, 35, 517-543.

Garnham, A., Oakhill, J., & Cruttenden, H. (1992). The role of implicit causality and gender cue in the interpretation of pronouns. *Language and Cognitive Processes*, 7, 231-255.
doi:10.1080/01690969208409386

Garvey, C. & Caramazza, A. (1974). Implicit causality in verbs. *Linguistic Inquiry*, 5, 459-64.

Garvey, C., Caramazza, A., & Yates, J. (1974). Factors influencing assignment of pronoun antecedents. *Cognition*, 3, 227-243.

Givón, T. (1983). *Topic Continuity in Discourse: A Quantitative Cross-Language Study* (Vol. 3). John Benjamins Publishing.

Grosz, B. J., Joshi, A. K., & Weinstein, S. (1995). Centering: A framework for modeling the local discourse. *Computational Linguistics*, 21, 203–225.

Gundel, J. K. (1988). Universals of topic-comment structure. *Studies in Syntactic Typology*, 17, 209-239.

Hartshorne, J. K. (2014). What is implicit causality? *Language, Cognition and Neuroscience*, 29, 804-824. doi:10.1080/01690965.2013.796396

Hartshorne, J.K., O'Donnell, T.J., & Tenenbaum, J.B. (2015). The causes and consequences explicit in verbs. *Language, Cognition and Neuroscience*, 30, 716-734.
doi:10.1080/23273798.2015.1008524

Hartshorne, J. K. & Snedeker, J. (2013). Verb argument structure predicts implicit causality: The advantages of finer-grained semantics. *Language and Cognitive Processes*, 28, 1474–1508. doi:10.1080/01690965.2012.689305

Why are implicit causes predictable?44

Hendriks, P. (2016). Cognitive modeling of individual variation in reference production and comprehension. *Frontiers in Psychology*, 7, 506. doi: [10.3389/fpsyg.2016.00506](https://doi.org/10.3389/fpsyg.2016.00506)

Holler, A., & Suckow, K. (2016). How clausal linking affects noun phrase salience in pronoun resolution. *Empirical Perspectives on Anaphora Resolution*, 61-85.

Järvikivi, J., van Gompel, R. P., & Hyönä, J. (2017). The interplay of implicit causality, structural heuristics, and anaphor type in ambiguous pronoun resolution. *Journal of Psycholinguistic Research*, 46, 525-550. doi:10.1007/s10936-016-9451-1

Johnson-Laird, P. N. (1983). Mental models: toward a cognitive science of language, inference and consciousness. Cambridge, MA: Harvard University Press. doi:[10.2307/414498](https://doi.org/10.2307/414498)

Jurafsky, D. (1996). A probabilistic model of lexical and syntactic access and disambiguation. *Cognitive Science*, 20, 137–194. doi:10.1207/s15516709cog2002_1

Kaiser, E. (2019). Order of mention in causal sequences: Talking about cause and effect in narratives and warning signs. *Discourse Processes*, 56(8), 599–618. <https://doi-org.libproxy.lib.unc.edu/10.1080/0163853X.2018.1522913>

Kehler, A. (2002). Coherence, Reference, and the Theory of Grammar, CSLI Publications.

Kehler, A., Kertz, L., Rohde, H., Elman, J.L. (2008). Coherence and coreference revisited. *Journal of Semantics*, 25, 1-44. doi:10.1093/jos/ffm018

Kehler, A., & Rohde, H. (2013). A probabilistic reconciliation of coherence-driven and centering-driven theories of pronoun interpretation. *Theoretical Linguistics*, 39, 1-37. doi:10.1515/tl-2013-0001.

Kehler, A., & Rohde, H. (2017). Evaluating an expectation-driven question-under-discussion model of discourse interpretation. *Discourse Processes*, 54(3), 219–238. <https://doi-org.libproxy.lib.unc.edu/10.1080/0163853X.2016.1169069>

Why are implicit causes predictable?45

Kehler, A., & Rohde, H. (2019). Prominence and coherence in a Bayesian theory of pronoun interpretation. *Journal of Pragmatics*. doi:10.1016/j.pragma.2018.04.006

Kochari, A. R., & Flecken, M. (2019). Lexical prediction in language comprehension: a replication study of grammatical gender effects in Dutch. *Language, Cognition and Neuroscience*, 34, 239-253. doi:10.1080/23273798.2018.1524500

Koornneef, A. W., & Sanders, T. J. (2013). Establishing coherence relations in discourse: The influence of implicit causality and connectives on pronoun resolution. *Language and Cognitive Processes*, 28, 1169-1206. doi:10.1080/01690965.2012.699076

Koornneef, A. W., & Van Berkum, J. J. (2006). On the use of verb-based implicit causality in sentence comprehension: Evidence from self-paced reading and eye tracking. *Journal of Memory and Language*, 54, 445-465. doi:10.1016/j.jml.2005.12.003

Koornneef, A., Dotlačil, J., van den Broek, P., & Sanders, T. (2016). The influence of linguistic and cognitive factors on the time course of verb-based implicit causality. *The Quarterly Journal of Experimental Psychology*, 69(3), 455–481. <https://doi-org.libproxy.lib.unc.edu/10.1080/17470218.2015.1055282>

Kowalski, A., & Huang, Y. T. (2017). Predicting and priming thematic roles: flexible use of verbal and nonverbal cues during relative clause comprehension. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 43, 1341-1351. doi:10.1037/xlm0000389

Kutas, M., & Hillyard, S. A. (1984). Brain potentials during reading reflect word expectancy and semantic association. *Nature*, 307, 161-163. doi:10.1038/307161a0

Langlois, V. J., Zerkle, S., & Arnold, J. E. (Under review). Referential Expectation Explains Linguistic and Social Constraints on Pronoun Comprehension. Ms., University of North Carolina. Submitted March 2021.

Why are implicit causes predictable?46

Levin, B. (1993). *English Verb Classes and Alternations: A Preliminary Investigation*. Chicago: University of Chicago Press.

Levy, R. (2008). Expectation-based syntactic comprehension. *Cognition*, 106, 1126-1177.
doi:10.1016/j.cognition.2007.05.006

Long, D. L., & De Ley, L. (2000). Implicit causality and discourse focus: The interaction of text and reader characteristics in pronoun resolution. *Journal of Memory and Language*, 42, 545-570. doi: [10.1006/jmla.1999.2695](https://doi.org/10.1006/jmla.1999.2695)

Lowder, M. W., & Ferreira, F. (2016). Prediction in the processing of repair disfluencies: Evidence from the visual-world paradigm. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 42, 1400-1416. doi:10.1037/xlm0000256

MacDonald, M.C. (2013). How language production shapes language form and comprehension. *Frontiers in Psychology*, 4. doi:10.3389/fpsyg.2013.00226

MacDonald, M. C., & Thornton, R. (2009). When language comprehension reflects production constraints: Resolving ambiguities with the help of past experience. *Memory & Cognition*, 37, 1177-1186. doi:10.3758/MC.37.8.1177

Magliano, J. P., Baggett, W. B., Johnson, B. K., & Graesser, A. C. (1993). The time course of generating casual antecedent and causal consequence inferences. *Discourse Processes*, 16(1-2), 35-53. <https://doi-org.libproxy.lib.unc.edu/10.1080/01638539309544828>

Mak, W. M., Tribushinina, E., & Andreiushina, E. (2013). Semantics of connectives guides referential expectations in discourse: An eye-tracking study of Dutch and Russian. *Discourse Processes*, 50(8), 557-576. <https://doi-org.libproxy.lib.unc.edu/10.1080/0163853X.2013.841075>

Why are implicit causes predictable?47

Majid, A., Sanford, A. J., & Pickering, M. J. (2006). Covariation and quantifier polarity: What determines causal attribution in vignettes?. *Cognition*, 99, 35-51.

doi:10.1016/j.cognition.2004.12.004

Majid A. Sanford A. J., Pickering M. J. (2007). The linguistic description of minimal social scenarios affects the extent of causal inference making *Journal of Experimental Social Psychology*. 43: 918-932. DOI: 10.1016/j.jesp.2006.10.016

McDonald, J. L., & MacWhinney, B. (1995). The time course of anaphor resolution: Effects of implicit verb causality and gender. *Journal of Memory and Language*, 34, 543-566.

doi:10.1006/jmla.1995.1025

McKoon, G., Greene, S. B., & Ratcliff, R. (1993). Discourse models, pronoun resolution, and the implicit causality of verbs. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 19, 1040–1052. [doi:10.1037/0278-7393.19.5.1040](https://doi.org/10.1037/0278-7393.19.5.1040)

Mitchell, D. C., Cuetos, F., Corley, M. M. B., & Brysbaert, M. (1995). Exposure-based models of human parsing: Evidence for the use of coarse-grained (nonlexical) statistical records. *Journal of Psycholinguistic Research*, 24, 469–488. doi:10.1007/BF02143162

Niemi, L., Hartshorne, J., Gerstenberg, T., & Young, L. (2016). Implicit measurement of motivated causal attribution. In *Proceedings of the 38th Annual Conference of the Cognitive Science Society*, Austin, TX, 2016 (pp. 1745-1750). Cognitive Science Society.

Pickering, M. J., & Garrod, S. (2007). Do people use language production to make predictions during comprehension?. *Trends in Cognitive Sciences*, 11, 105-110.

doi:10.1016/j.tics.2006.12.002

Pickering, M. J., & Majid, A. (2007). What are implicit causality and consequentiality?

Language and Cognitive Processes, 22, 780-788. doi:10.1080/01690960601119876

Why are implicit causes predictable?48

- Pyykkönen, P., & Järvikivi, J. (2010). Activation and persistence of implicit causality information in spoken language comprehension. *Experimental Psychology*, 57, 5-16.
doi:10.1027/1618-3169/a000002
- Prince, E. (1981). *Toward a Taxonomy of Given-New Information*. In P.Cole, (Ed.). *Radical Pragmatics* (pp. 223-256). NY: Academic Press.
- Prince, E. (1992). The ZPG letter: Subjects, definiteness, and information-status. In W. Mann & S. Thompson (Eds.), *Discourse description: Diverse linguistic analyses of a fund-raising text* (pp. 295-326). Amsterdam: Benjamins.
- Reinhart, T. (1982). Pragmatics and linguistics: An analysis of sentence topics. *Philosophica*, 27, 53-94.
- Rohde, H., & Kehler, A. (2014). Grammatical and information-structural influences on pronoun production. *Language, Cognition and Neuroscience*, 29, 912-927.
doi:10.1080/01690965.2013.854918
- Rosa, E. C., & Arnold, J. E. (2017). Predictability affects production: Thematic roles can affect reference form selection. *Journal of Memory and Language*, 94, 43-60.
doi:10.1016/j.jml.2016.07.007
- Ryskin, R., Ng, S., Mimnaugh, K., Brown-Schmidt, S., & Federmeier, K. D. (2019). Talker-specific predictions during language processing. *Language, Cognition and Neuroscience*, 1-16. doi:10.1080/23273798.2019.1630654
- Solstad, T., & Bott, O. (2017). Causality and causal reasoning in natural language. In M. R. Waldmann (Ed.), *The Oxford handbook of causal reasoning*. (pp. 619–644). Oxford University Press.

Why are implicit causes predictable?49

Saffran, J. R. (2002). Constraints on statistical language learning. *Journal of Memory and Language*, 47, 172-196. doi:10.1006/jmla.2001.2839

Saffran, J. R. (2003). Statistical language learning: Mechanisms and constraints. *Current Directions in Psychological Science*, 12, 110-114. doi:10.1111/1467-8721.01243

Smith, N. J., & Levy, R. (2013). The effect of word predictability on reading time is logarithmic. *Cognition*, 128, 302–319. doi:10.1016/j.cognition.2013.02.013

Stevenson, R.J., Crawley, R.A., & Kleinman, D. (1994). Thematic roles, focus, and the representation of events. *Language and Cognitive Processes*, 9, 519-548.
doi:10.1080/01690969408402130

Stevenson, R., Knott, A., Oberlander, J., & McDonald, S. (2000). Interpreting pronouns and connectives: Interactions among focusing, thematic roles and coherence relations. *Language and Cognitive Processes*, 15, 225-262.
doi:10.1080/016909600386048

Stewart, A. J., Pickering, M. J., & Sanford, A. J. (2000). The time course of the influence of implicit causality information: Focusing versus integration accounts. *Journal of Memory and Language*, 42, 423-443. doi:10.1006/jmla.1999.2691

Tabor, W., Juliano, C., & Tanenhaus, M. K. (1997). Parsing in a dynamical system: An attractor-based account of the interaction of lexical and structural constraints in sentence processing. *Language and Cognitive Processes*, 12, 211–271.
doi:10.1080/016909697386853

Tily, H., & Piantadosi, S. (2009). Refer efficiently: Use less informative expressions for more predictable meanings. In *Proceedings of the workshop on the production of referring*

Why are implicit causes predictable?50

expressions: Bridging the gap between computational and empirical approaches to reference.

Townsend, D. J., & Bever, T. G. (2001). *Sentence Comprehension: The Integration of Habits and Rules*. MIT Press.

van den Hoven, Emiel & Ferstl, Evelyn. (2018). Discourse context modulates the effect of implicit causality on rementions. *Language and Cognition*, 10. 1-34.
doi:10.1017/langcog.2018.17.

van Dijk, T. A. & Kintsch, W. (1983). *Strategies of discourse comprehension*. New York: Academic Press. doi:[10.2307/415483](https://doi.org/10.2307/415483)

van Rij, J., van Rijn, H., & Hendriks, P. (2013). How WM load influences linguistic processing in adults: A computational model of pronoun interpretation in discourse. *Topics in Cognitive Science*, 5, 564–580. doi:10.1111/tops.12029

Viebahn, M. C., Ernestus, M., & McQueen, J. M. (2015). Syntactic predictability in the recognition of carefully and casually produced speech. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 41, 1684-1702. doi:10.1037/a0039326

Weatherford, K. & Arnold, J. E. (2019) Semantic predictability of implicit causes affects referential form choice. Poster, CUNY conference on human sentence processing, University of Colorado (March 2019).

Weatherford, K., & Arnold, J. E. (2021). Semantic predictability of implicit causality can affect referential form choice. *Cognition*, 214, 104759. DOI:
<https://doi.org/10.1016/j.cognition.2021.104759>

Why are implicit causes predictable?51

Wells, J. B., Christiansen, M. H., Race, D. S., Acheson, D. J., & MacDonald, M. C. (2009).

Experience and sentence processing: Statistical learning and relative clause

comprehension. *Cognitive Psychology*, 58, 250-271. doi:10.1016/j.cogpsych.2008.08.002

Williams, E. & Arnold, J. E. (Under review). Individual differences in print exposure predict use of implicit causality in pronoun comprehension and referential prediction. Ms., University of North Carolina.

Zerkle, S., & Arnold, J. E. (2019, March). *Similar properties guide both topicality & referential predictability judgments*. Poster session presented at the CUNY Conference on Human Sentence Processing, Boulder, CO.

Appendix A: Google Corpus Analysis Methods

We assembled a dataset by searching Google for tokens with 16 emotion verbs (8 subject-biased and 8 object-biased), listed in Table A1. For each verb, we put it in different strings, hereafter referred to as the “target clauses,” according to the 14 pairwise combinations of the pronouns, *I/me, you, he/him, she/her*, followed by *because*. For example, *I appreciated him because*. We searched for up to three examples of each string, including fewer if we could not find three. This meant each verb had a maximum of 42 possible tokens. Verbs had to have at least 20 tokens to be included for analysis; all 16 verbs met this requirement.

We searched for exact matches of the target clauses on Google in Incognito Mode (to prevent search history biases). For each target clause, we found the first three search results that matched our token selection criteria, given in Table A2, and recorded them in the database, including information on the source, type of text, and sentence in which the target clause appeared. The top five most common types of text were book, blog, comment, article, and fanfiction. Since we were looking for language appearing naturally in a rich discourse context, we excluded memes, pictures or short videos with words on top, and quote translations.

Table A1. Google analysis emotion verbs

Verb type	Verbs included in analysis (n for each verb)
Subject-biased	Amazed (25), disappointed (33), impressed (24), surprised (30), upset (28), annoyed (36), offended (30), scared (41)
Object-biased	Appreciated (26), hated (42), liked (42), noticed (39), respected (38), trusted (42), admired (41), adored (31)

The criteria used for inclusion were that the verb must occur in an active transitive frame with critical implicit cause and non-cause roles in subject and object position, and that “because” is used in a causal sense. In order to guarantee that the following *because* clause was only explaining the event related to the implicit causality verb, we excluded instances where the target clause containing the implicit causality verb was embedded, like *I just figured **she hated me because** last week I changed my profile picture to this one she took of me in a suit...* since it is ambiguous whether the *because* explains *I just figured* or *she hated me*. Examples of included and excluded tokens are shown in Table A2.

The first half of the text analysis was conducted in collaboration with Irene Tang; a subset of tokens was double coded and disagreements were resolved by discussion. For the remainder of the analysis, the first author (SG) coded all tokens and discussed unclear cases with the second author (JA).

Table A2. Token selection criteria

Criterion	Example
The target clause must be a tensed main or subordinate clause, i.e. embedded clauses and infinitival clauses were excluded. Results that matched the string but were clearly not simple past tense were not accepted. One common example of this was questions with a fronted modal. Structurally ambiguous cases were not	<u>Not usable</u> : “I hope I offended you. Because that's what I do.” <u>Not usable</u> : “I know I'm leaving out tons of details – not sure if they would be relevant or not – but the core issue was she felt I was too nice all the time and I annoyed her because I was nice to people she felt I should blow off.”

Why are implicit causes predictable?53

accepted.	<u>Not usable:</u> Had I annoyed Him because I hadn't been able to get it sooner? <u>Usable:</u> I say, if I annoyed you because I lost a few pounds, what would you say if I were to lose hundreds—thousands.
The target clause must be followed by a tensed clause.	<u>Not usable:</u> ...he admired him because of the intelligence, alertness and understanding that he saw in him
Disfluency, unfinished fragments and discourse markers were excluded for the purpose of clause inclusion; excluded material shown here in brackets.	You hated him because {...well. Because} he was a man.
For results that matched the string, but the first pronoun did not fill the subject, the token was not accepted.	<u>Not usable:</u> Clearly from reading the reviews above, many of you appreciated him because you need people like him to give you a break and help you achieve things in life.
The target clause and the following clause did not have to occur in the same sentence, given that it was clear the second clause was a continuation of the first. Cases where the string was separated between two speakers were accepted when “PRO VERB PRO” was spoken by one speaker and “Because” was spoken by the other speaker, with PRO’s being categorized under the appropriate referents from the first speaker’s perspective.	<u>Not usable:</u> “I figured I'd probably offended you.” “Offended me? Because you weren't interested in hiring me as your station attendant?” <u>Usable:</u> “You scared me.” “Because you wandered too far from camp. 'Tis dark and not safe for you alone.”
Ungrammatical results, or results in which it was clear that the target clause was being used in a nonstandard way were excluded.	<u>Not usable:</u> I amazed her because she continually taking care for us even he is so tired for us.
Items where the referent mention was too ambiguous to be determined from context were excluded.	<u>Not usable:</u> Moreover, Goku never fought Super Janemba. He annoyed him, because Janemba is a mindless, savage beast with no combat ability whatsoever.

After finding appropriate tokens, we then applied the criteria shown in Table A3 below to select the clause immediately following the target clause, the window in which we looked for re-mention. The general rule was to select the next tensed clause, subordinate or not, and all material embedded under it.

Table A3. Selection criteria for the following clause, which is underlined.

Criterion	Example
The following clause must be a tensed clause, and can be a subordinate clause.	You amazed him, because <u>when there were other people around you trying to comfort you, you rejected them.</u>
Hanging fragments were included in the following clause.	She amazed him, because <u>in his final year she joined his class, becoming the youngest graduate in micro-surgery, a mere twenty-two year old, younger than many university entrants.</u>
In cases of ellipsis, if only the pronoun was elided, the second section after a conjunction was considered a separate clause and excluded. If more than the pronoun was elided, the section after a conjunction was included in the following clause.	You liked me because <u>I was young</u> and made you feel young. He liked me, because <u>I was good to the old folks, and to Emily,</u> - and had a sort of respect for me, because I was the oldest, and because I could talk, and because of the great thick books in my room.

Coding of dependent measure: reference continuation

For each token in the database, we coded whether the speaker or writer referred back to the critical entities (implicit cause or non-cause). Our question was centered on reference – that is, is the subject entity re-mentioned? Is the object entity re-mentioned? Usually re-mentions also used a pronoun, but sometimes a name or description was used instead (e.g., *Then he surprised her, because Eli was rarely demonstrative.*)

Our analysis focuses on the broadest definition of continued reference, including both **direct** and **indirect** reference. In direct reference, the following clause refers directly and entirely to an entity in the target clause. Indirect reference includes cases where the following clause refers to either a superset or subset of the critical referent. For example, consider *you surprised him because your teleportation technique is silent*. The cause in the target clause is *you*, and the following clauses mentioned *your teleportation technique*, which includes the possessive *your*. Other examples of indirect reference include *I...we*, or *me...someone I once loved*. We found that out of a total 548 tokens, there were very few instances of indirect reference (8.6% indirect subject re-mentions and 7.3 indirect object re-mentions), so these were collapsed with direct mentions.

Appendix B: Fisher Corpus Analysis Methods

Our goal was to generate a sample of utterances using emotion verbs that matched our criteria (see below), with at least 10 (and no more than 25) tokens per verb. This sample was smaller than the Google sample because a greater degree of hand coding was needed for both item selection and analysis. Our final sample included 78 tokens of subject-biased verbs, and 120 tokens of object-biased verbs. While other contextual features may change causality judgments (e.g., negation, Garvey, Caramazza, & Yates, 1974), we did not have a large enough sample to test the impact of these features. Our method of token selection was to search for the emotion verbs listed in Table B1, which represent commonly used verbs from psycholinguistic studies of implicit causality. We began our search with the files in part 2 (files 058-116), and if that search did not produce at least 10 tokens, we extended our search to part 1 (files 001-057). We limited the search to verbs in the past tense, and no other restrictions, using a sequence of grep searches. These tokens were examined for the criteria listed in Table B2. Based on our first (rough) examination, verbs with fewer than 10 usable tokens in the sample were excluded from the sample. However, some tokens were later discarded due to ambiguity or extreme disfluency, leaving fewer than 10 tokens for *offended*. In addition, we found that our initial pass only identified 5 subject-cause verbs with more than 10 usable tokens, so we relaxed our criterion for this sample and included two more verbs that had 5-7 usable tokens (*disappointed*, *impressed*).

Table B1. Fisher analysis emotion verbs

Verb type	Verbs searched for	Included in analysis (n for each verb)
Subject-biased	amused, annoyed, astonished, bored, deceived, delighted, disappointed, disgusted distressed, frightened, impressed, inspired, offended, pleased, scared, surprised, startled, terrified, troubled, upset, worried	Annoyed (10), offended (7), surprised (20), upset (13), scared (16), disappointed (7), impressed (5)
Object-biased	adored, appreciated, blamed, criticized, despised, detested, disliked, distrusted, dreaded, envied, feared, forgave, hated, idolized, loathed, liked, noticed, pitied, respected, trusted	Appreciated (20), hated (23), liked (21), noticed (20), respected (20), trusted (16)

The criteria used for inclusion of tokens are shown in Table B2. We selected transitive clauses in the active voice, including only clauses that contained both critical entities for each verb type, controlling for the syntactic structure of the clause, and only in cases where there was a following utterance in which a reference could be made.

We did not restrict our sample based on whether the following utterance specified an explanation for the previous event. While we know that implicit causality judgments are most strongly tied to re-mention predictions under explanation coherence relations, we did not know a priori how people might learn about re-mention frequency. It is possible that they count all instances of a verb, and the proportion of cases on which the implicit cause is re-mentioned, especially since the coherence relation may not be apparent until later in the utterance. By including all coherence relations, this also allowed us to observe how frequently each verb is followed by an explanation.

Table B2. Criteria for inclusion in analysis

Criterion	Example
The verb must occur in an active transitive frame with critical implicit cause and non-cause roles in subject and object position	<u>Not usable</u>: I'm really annoyed at the darn television <u>Not usable</u>: you shouldn't really be getting offended <u>Usable</u>: she appreciated the things that I would do for her <u>Usable</u>: I trusted them
The clause must be a tensed main or subordinate clause. Relative clauses and infinitival clauses were not accepted.	<u>Not usable</u>: like the big brother show which I never really liked <u>Usable</u>: if something uh really really insulted me or offended me or something
Disfluency, unfinished fragments, editorials and discourse markers are excluded for the purpose of clause inclusion; examples of excluded material shown here in brackets. In most cases, disfluent segments were excluded from the analysis, but the token was included. A few tokens that were extremely disfluent were excluded altogether.	we {um} noticed all the new barriers my friend noticed {the ah} the date on the pepper I always respected Russia and the people {and from the writings you know from and things that they went through} that upset me {you know however you say it}

Coding of dependent measure: reference continuation

For each token in the database, we coded whether the speaker referred back to the critical entities (implicit cause or non-cause). Our question was centered on reference – that is, is the subject entity re-mentioned? Is the object entity re-mentioned? Sometimes re-mentions involved the same word as the first mention (*he impressed these children / ... he came with a suit and everything*), but sometimes they didn't (*Robin Williams surprised me / when I saw him do stand up*). If the next utterance was spoken by a different speaker, the words *you* and *me* could refer to the same person.

Our analysis focuses on the broadest definition of continued reference, including both **direct** and **indirect** reference. In direct reference, the following clause refers directly and entirely to an entity in the target clause. For example, in (B1), both the non-cause *we* and the implicit cause *our freedom* are mentioned directly in the following clause. Note that direct reference does not have to use the same words, as pronouns are frequently used to refer to recently mentioned entities, like *it* in (B1).

(B1) **Target clause:** we appreciated our freedom

Following clause: and we liked it

Indirect reference includes cases where the following clause refers to either a superset or subset of the critical referent. For example, consider *I hated Melissa / and Mo Jo we hated Mo Jo too*. The non-cause in the target clause is *I*, and the following clause mentions *we*, which is a superset of *I*. Other examples of indirect reference include *her kids... she; I... the girl I got*

matched with last time; him...his father. There were relatively few items with indirect reference (out of the 101 cases of subject reference, 11 were indirect, and of the 93 cases of object reference, 15 were indirect) so we selected the broader category for analysis, including all cases of either direct or indirect reference.

Coding of predictors: discourse status

For each token, we also coded the discourse status of both the subject and object referents on four dimensions. First, was the referent first or second person (I, we, or you) or a third-person referent? Second, was the subject or object pronominalized? Third, did the subject or object refer to a referent that was already given in the context? First and second person pronouns were automatically considered given, because they referred to the discourse participants. For third person referents, coders examined the discourse context and coded whether the referent had previously been mentioned in the conversation, in which case they were coded as given. Referents that had not been mentioned were coded as new. If a referent could be inferred from the prior context, it was coded as inferable (Prince, 1992); this category was combined with the given category for analysis. Fourth, was the subject or object entity animate? Coders looked at the context to identify the meaning of each referent, and coded it as animate (people or animals) or inanimate (things or ideas).

Coding of predictors: coherence relations

We coded the coherence relation between the critical emotion-verb clause and the following clause, using the system devised by Hannah Rohde and Andy Kehler⁵. We also added the category “other” to capture cases that did not fit into any of the existing categories. The critical question was whether implicit cause continuations are more likely in explanation vs. other contexts. We therefore collapsed the coherence ratings into binary categories: explanation vs. other. The final coherence ratings were done by a single RA (SW), and compared with coherence ratings for subsets of the data by two other RAs (LR, AMF). Coders averaged 91% agreement, and cases of disagreement were solved through discussion between the second author (JA) and SW. Examples are given in Table B3 below.

Table B3. Coherence relation coding system

Coherence Relation and Definition	Example
<u>Explanation</u> : Explanation about the previous event or general information about the cause of an event	I hated that because I hate to be that kind of person.
<u>Elaboration</u> : Elaborates on the same event, e.g. how it is carried out or where/when	My dad hated New Mexico. My whole family did.
<u>Occasion</u> : Temporal relation between two sentences where second sentence describes an event that follows the first sentence, but	A: I never liked this sort of {type} patriotism. B: {exactly} and it has increased since September eleventh

⁵ We are very grateful to Hannah Rohde and Andrew Kehler for sharing their coding schema, which is based on the inventory of relations in Kehler (2002).

there is no casual relation (if there is, the explanation coding takes precedence)	
<u>Parallel</u> : Similar event with different referents or similar referents in parallel event	I liked the fear factor show and the aspect of the daring things they have to do like getting on a airplane and walking down the wing and getting the flag and coming back. I don't like the bug things.
<u>Result</u> : Causal result of previous event	He really impressed me so we switched our whole family over to the same doctor.
<u>Violated Expectation</u> : An unexpected outcome given general real-world knowledge about likely events and their typical consequences/reactions	The parents certainly appreciated it. Although {I} even then it was funny
<u>Background</u> : Background information that elaborates on some aspect of the event	I certainly appreciated the things that she did for me. There are a bunch of other things that I look for too.
<u>Other</u> : No specific relationship with prior event, or both are related to some higher-level event, or the following clause refers to the entire previous statement.	A: she appreciated the things that I would do for her B: {mhm mhm oh} that's great (here <i>that's great</i> refers to the entire previous segment and really serves to introduce B's taking the floor).

Appendix C. Combined analysis for Experiments 1 and 2.

Experiment 1 tested the predictability question, and Experiment 2 tested the topicality question. We found that Verbttype affected predictability questions but not topicality questions. To confirm the contrast between these experiments, we analyzed the two together. In both cases the dependent measure was whether the subject was chosen or not. We included both predictability and topicality questions in the same analysis, adding question type (predictability vs. topicality) as a predictor, plus interactions between it and all other predictors. Otherwise the modeling procedure was identical to the two single experiments.

As shown in Table D1, there was a main effect of Verb type, a main effect of Question type, and an interaction between the two. There was also a marginal three-way interaction between Verbttype, Question type, and Causal coherence. We probed the significant verbttype x question interaction by estimating the verbttype effect for the Prediction and Topicality questions, and found that it was significant only for the Prediction question (Table C2).

Table C1. Analysis of Experiments 1 and 2 together with subject chosen as dependent measure.

Effect	Estimate (SE)	t value	p value
Intercept	1 (0.23)	4.35	0.0002
Verb type	0.88 (0.32)	2.8	0.0067
Causal coherence	-0.46 (0.44)	-1.06	0.2964
Verb type x Causal coherence	0.83 (0.63)	1.32	0.1904
Question type	-1.51 (0.55)	-2.73	0.0135
Verbttype x Question	2.13 (0.63)	3.38	0.0012
Causal x Question	0.24 (0.87)	0.28	0.7815
Verbttype x Causal x Question	2.15 (1.26)	1.71	0.0928

Table C2. Estimates for the Verbttype effect in different conditions.

Effect	Estimate (SE)	t value	p value
Verbttype effect for Prediction question	1.95 (0.43)	4.58	<.0001
Verbttype effect for Topicality question	-0.18 (0.47)	-0.39	0.6986

Appendix D. Power analysis for Experiments 1 and 2.

We used the results from the analyses for experiments 1 and 2 as an estimate of the population of responses with this paradigm. We simulated 1000 random datasets for exp. 1 and 999 for exp. 2* by sampling from a multivariate normal distribution for all random effects, and ran each one through the original model. Table D1 shows the average estimates and standard error for the intercept and each predictor. We calculated the power as the percentage of items where that parameter had a p value of .05 or lower.

Table D1. Power simulations for Experiments 1 and 2

Experiment 1 (Predictability question)			Experiment 2 (Topicality question)		
Effect	Estimate (SE)	power	Effect	Estimate (SE)	power
Intercept	0.20 (0.30)	0.10	Intercept	1.44 (0.30)	0.99
Verb type	1.65 (0.37)	1.0	Verb type	-0.15 (0.42)	0.04
Causal coherence	-0.27 (0.55)	0.08	Causal coherence	-0.47 (0.49)	0.15
Verb type x Causal coherence	1.66 (0.74)	0.59	Verb type x Causal coherence	-0.16 (0.84)	0.06

There are two results of interest. First, in the main models we found an effect of Verbytype for the Predictability question (Experiment 1) but not for the Topicality question (Experiment 2). Our simulations suggest that with a sample of this size we have adequate power for detecting these patterns. For the Predictability question, the verbytype effect was estimated to be positive, reflecting a greater chance of choosing the subject for the subject-biased verbs than for the object-biased verbs. The likelihood of detecting a verbytype effect was 1.0, reflecting the fact that it had a p value at .05 or lower in every single one of the 1000 simulated datasets. For experiment 2, the verb type effect was highly unlikely to be significant (resulting in only 4% of the cases), and in addition was estimated to run in the opposite direction.

Second, in the main models the intercept was significant for the Topicality question but the not Predictability question. Our sample size has adequate power for detecting this pattern as well. For the topicality question, the intercept was significant in 99% of the simulated datasets, but for the predictability question, only in 10%.

* 1000 were requested but one failed and was not reported

Supplement

A. Verbs used in all studies

We analyzed subject-biased (e.g., *impress*, *anger*) and object-biased (e.g., *admire*, *like*) interpersonal verbs. We focused on main clause, active tensed forms of these verbs, in structures where one person was the subject and another was the object, since this has been the focus of the majority of discourse processing work.

Verbs were chosen based on their object bias scores from Hartshorne and Snedeker (2013), which are the frequency with which participants choose the object in response to the following task involving the nonce word “dax”, e.g. “Sally frightens Mary because she is a dax”; Who do you think is a dax? (Sally/Marry). Scores thus have a possible range from 0 to 1 and the complement of a verb’s object bias score can be considered its subject bias score. Scores can be considered both a measure of bias for who is the implicit cause and who the pronoun is resolved to.

Overall, 19 different subject-biased verbs ranged from 0.13 to 0.38 in object bias (thus 0.62 to 0.87 in subject bias), and 16 different object-biased verbs ranged from 0.74 to 0.89 in object bias. Most of these verbs come from Levin’s (1993) verb classes 31.1 (*amuse*-type verbs) and 31.2 (*admire*-type verbs), and assign the thematic roles of stimulus and experiencer to their arguments. However, the object-biased verbs include *notice* from verb class 2.13.1, and *tease* from class 31.1 (even though most 31.1 verbs are subject-biased verbs). We therefore discuss roles with the labels “implicit cause” and “non-cause.” We used a slightly different selection of verbs in the experiments and text analyses, because the experiments were designed to be compared with other pronoun comprehension studies in our lab (Williams, 2020; Johnson & Arnold, 2021). The full list of verbs used in this paper is listed in Table A1.

Table A1. Verbs used and their object biases according to Hartshorne & Snedeker (2013) Experiment 2, except for *notice* from Experiment 1.

Subject-biased Verbs	Fisher Text Analysis	Google Text Analysis	Exp. 1 and 2	Levin verb class	H & S (2013) Exp. 2 Object Bias
annoyed	x	x		31.1	0.27
disappointed	x	x		31.1	0.16
impressed	x	x		31.1	0.15
offended	x	x	x	31.1	0.23
scared	x	x		31.1	0.13
surprised	x	x		31.1	0.28
upset	x	x		31.1	0.38
amaze		x		31.1	0.23
amused			x	31.1	0.18
inspired			x	31.1	0.19
pleased			x	31.1	0.22
bored			x	31.1	0.24
dazzled			x	31.1	0.22

Why are implicit causes predictable?62

enraged			x	31.1	0.17
fascinated			x	31.1	0.26
aggravated			x	31.1	0.27
irritated			x	31.1	0.22
distracted			x	31.1	0.24
frightened			x	31.1	0.19
Object-biased Verbs					
appreciated	x	x		31.2	0.79
hated	x	x	x	31.2	0.85
liked	x	x		31.2	0.87
noticed	x	x		2.13.1	0.74 (Exp. 1)
respected	x	x		31.2	0.84
trusted	x	x	x	31.2	0.85
admired		x	x	31.2	0.89
adored		x	x	31.2	0.88
envied			x	31.2	0.86
resented			x	31.2	0.85
despised			x	31.2	0.89
disliked			x	31.2	0.89
idolized			x	31.2	0.87
loathed			x	31.2	0.86
teased			x	31.1	0.80
worshipped			x	31.2	0.84

B. Corpus analysis results by verb

Table B1. Fisher Corpus analysis: Rate of continued mention of implicit cause and non-cause entities, by verb.

	Implicit cause continuation	Non-cause continuation
<u>Subject-biased verbs</u>		
Annoyed (n = 10)	30%	50%
Offended (n = 7)	43%	71%
Surprised (n = 20)	70%	55%
Upset (n=13)	38%	46%
Scare (n=16)	13%	56%
Disappoint (n=7)	43%	14%
Impress (n=5)	40%	40%
<u>Object-biased verbs</u>		
Appreciated (n=20)	40%	50%
Hated (n=23)	35%	70%
Liked (n=21)	48%	62%
Noticed (n=20)	35%	40%
Respected (n=20)	55%	65%
Trusted (n=16)	63%	56%

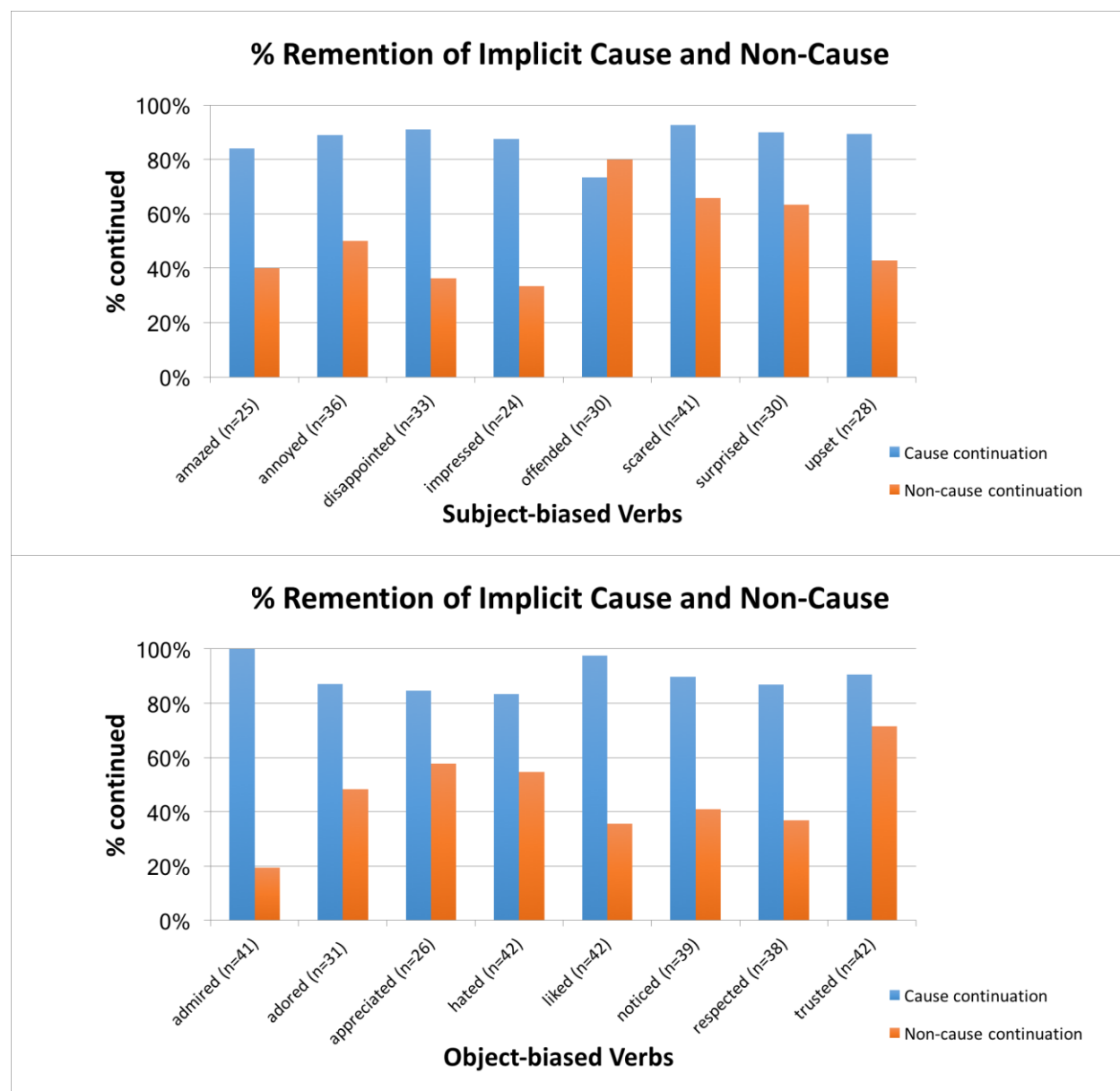


Figure B1 Implicit causality verbs. Rate of continued mention of implicit cause and non-cause entities. Top panel shows subject-biased verbs; bottom panel shows object-biased verbs. The percentage continued reflects the percentage of tokens in which each role is mentioned in the immediately next clause.

C. Materials used in Experiments 1 and 2

Item ID [†]	Context Sentence	Target Fragment
P1	Will and Matt were doing the laundry.	Will folded the clothes with Matt (because)...
P2	Liz and Ana were fixing the bathroom sink.	Liz took out the wrench (because)...
T1	Will and Matt were fishing with their kids.	Will admired Matt (because)...
T1	Will and Matt were fishing with their kids.	Will inspired Matt (because)...
T2	Will and Matt were working at the same company.	Will idolized Matt (because)...
T2	Will and Matt were working at the same company.	Will fascinated Matt (because)...
T3	Matt and Will were rooming together in college.	Matt adored Will (because)...
T3	Matt and Will were rooming together in college.	Matt amused Will (because)...
T4	Matt and Will were working out at the gym.	Matt loathed Will (because)...
T4	Matt and Will were working out at the gym.	Matt aggravated Will (because)...
T5	Liz and Ana were working on a project for class.	Liz despised Ana (because)...
T5	Liz and Ana were working on a project for class.	Liz bored Ana (because)...
T6	Liz and Ana were driving to a family reunion.	Liz resented Ana (because)...
T6	Liz and Ana were driving to a family reunion.	Liz irritated Ana (because)...
T7	Liz and Ana were volunteering at the library	Liz offended Ana (because)...
T7	Liz and Ana were volunteering at the library	Liz disliked Ana (because)...
T8	Liz and Ana were putting up Christmas decorations.	Liz distracted Ana (because)...
T8	Liz and Ana were putting up Christmas decorations.	Liz teased Ana (because)...
T9	Will and Matt were attending an office party.	Will dazzled Matt (because)...
T9	Will and Matt were attending an office party.	Will envied Matt (because)...
T10	Will and Matt were camping in the woods.	Will frightened Matt (because)...
T10	Will and Matt were camping in the woods.	Will trusted Matt (because)...
T11	Ana and Liz were competing in a marathon.	Ana enraged Liz (because)...
T11	Ana and Liz were competing in a marathon.	Ana hated Liz (because)...
T12	Ana and Liz were practicing for a ballet performance.	Ana pleased Liz (because)...
T12	Ana and Liz were practicing for a ballet performance.	Ana worshipped Liz (because)...
F1	Ana and Liz were watching TV.	Ana changed the channel (because)...
F2	Ana and Will were changing their furniture.	Ana assembled the bed with Will (because)...
F3	Matt was playing the piano.	Matt enjoyed himself (because)...
F4	Liz was preparing a presentation.	Liz created an outline (because)...
F5	Matt and Liz were going to decorate the house.	Matt went to the store with Liz (because)...
F6	Matt and Will were eating breakfast.	Will took out the cereal (because)...

[†]P stands for Practice item, T stands for Target item, and F stands for Filler item.

Supplement References

- Hartshorne, J. K. & Snedeker, J. (2013). Verb argument structure predicts implicit causality: The advantages of finer-grained semantics. *Language and Cognitive Processes*, 28, 1474–1508. doi:10.1080/01690965.2012.689305
- Johnson, E. & Arnold, J. E. (2021). Individual differences in print exposure predict use of implicit causality in pronoun comprehension and referential prediction. *Frontiers*.
- Williams, E. (2020). Language experience affects pronoun comprehension in implicit causality sentences. Masters thesis, University of North Carolina at Chapel Hill.