

## **The Effect of Context on Noisy-Channel Sentence Comprehension**

Sihan Chen<sup>1,3,\*</sup>, Sarah Nathaniel<sup>1,\*</sup>, Rachel Ryskin<sup>2</sup>, & Edward Gibson<sup>1</sup>

1. Department of Brain and Cognitive Sciences

Massachusetts Institute of Technology, Cambridge, MA 02139

2. Department of Cognitive and Information Sciences,

University of California, Merced, Merced, CA 95343

3. Email: [sihanc@mit.edu](mailto:sihanc@mit.edu)

\* Equal contributions

Word count: 7184

### Abstract

The process of sentence comprehension must allow for the possibility of noise in the input, e.g., from speaker error, listener mishearing, or environmental noise. Consequently, semantically implausible sentences such as *The girl tossed the apple the boy* are often interpreted as a semantically plausible alternative (e.g., *The girl tossed the apple to the boy*). Previous investigations of noisy-channel comprehension have relied exclusively on paradigms with isolated sentences. Because supportive contexts alter the expectations of possible interpretations, the noisy channel framework predicts that context should encourage more inference in interpreting implausible sentences, relative to null contexts (i.e. a lack of context) or unsupportive contexts. In the present work, we tested this prediction in four types of sentence constructions: two where inference is relatively frequent (double object - prepositional object), and two where inference is rare (active-passive). We found evidence that in the two sentence types that commonly elicit inference, supportive contexts encourage noisy-channel inferences about the intended meaning of implausible sentences more than non-supportive contexts or null contexts. These results suggest that noisy-channel inference may be more pervasive in everyday language processing than previously assumed based on work with isolated sentences.

**Keywords:** noisy-channel, sentence comprehension, context, rational inference, error correction

## 1. Introduction

### 1.1. Background

Written and spoken language often includes noise, such as typographical errors in written language, or ambient sound in spoken language. This noise can, at times, make comprehension challenging (e.g., talking to someone on a cell phone with a bad connection). But in many cases the presence of noise is hardly perceptible, despite being an ever-present factor in communication (Fano, 1961). For example, readers often fail to notice inserted function words and their comprehension is rarely affected (Staub, Dodge, & Cohen, 2019). This may be in part because, in spite of noise corruption, the intended meaning can often be inferred from contextual and world-knowledge information. For example, when you are in a video-conferencing meeting (e.g., a Zoom meeting), sometimes you cannot hear all the words another participant is saying (possibly due to internet or other technical issues). Despite this, you are often able to infer what they are saying, based on what they said before. In the present work, we examine the role of discourse context in how comprehenders interpret sentences in noise.

Traditional theories of sentence processing assume a noise-free representation of the input (e.g., Frazier & Fodor, 1978; Gibson, 1998; Levy, 2008a). However, readers often interpret language non-literally when the input is implausible (Ferreira, 2003; Ferreira & Patson, 2007; Bader & Meng, 2018; Meng & Bader, 2021; Cai et al., 2022). Recent proposals argue that comprehenders are, in fact, well-adapted to processing imperfect linguistic input (Gibson, Bergen, & Piantadosi, 2013; Levy, 2008b; Levy, Bicknell, Slattery, & Rayner, 2009; Levy, 2011). On these accounts, based on the framework proposed by Shannon (1948), comprehenders infer the intended sentence,  $s_i$ , from the perceived sentence,  $s_p$ , on the assumption that some noise corruption may have transformed  $s_i$  into  $s_p$  during transmission. Following Bayes' rule (1),

the probability of inferring  $s_i$  given  $s_p$  can be obtained from the probability that the speaker would communicate  $s_i$ , given world and language knowledge,  $p(s_i)$ , and the likelihood of the particular noise operations being applied to the intended sentence,  $p(s_i \rightarrow s_p)$ . Communication is effective when this intended sentence,  $s_i$ , can be understood from the perceived sentence,  $s_p$ .

$$p(s_i | s_p) \propto p(s_i) p(s_i \rightarrow s_p) \quad (1)$$

This framework finds support in empirical evidence that readers maintain uncertainty about previously read words in a sentence and revise prior interpretations according to subsequent input (Levy et al., 2009; Bergen, Levy & Gibson, 2012). Furthermore, Gibson et al. (2013), tested four predictions of the noisy-channel framework: (1) the fewer noise operations (namely insertions and deletions) that are needed to be posited to revert a perceived implausible sentence to a plausible sentence, the more likely a comprehender will be to adopt the plausible interpretation; (2) noise operations should unequally regard insertions and deletions, according to the “Bayesian size principle” – to delete a word, there exists only a limited number of options, whereas to insert a word, there are as many options as one’s vocabulary size, and hence the probability of a word being deleted from a sentence is much higher than a word being added to a sentence; (3) non-literal perception of sentences should increase with perceived noise rate; and (4) non-literal perceptions should decrease as the rate of semantically implausible sentences increases. Participants in Gibson et al. read sentences (e.g., *The mother gave the candle the daughter*) and answered comprehension questions about them (e.g., Did the daughter receive something?) which revealed how they were interpreted. When they encountered sentences that were semantically implausible, they could either answer the comprehension question based on a

literal interpretation of the sentence (i.e., No, the candle received something), or a non-literal interpretation of the sentence (i.e., Yes, the daughter received something). Choosing the non-literal interpretation implies that the comprehender made an inference that the more semantically plausible meaning was intended. The framework is hence tested by comparing the inference rate (or the non-literal interpretation rate) observed in the experiment with the predicted  $p(s_i | s_p)$ . Gibson et al. observed that, consistent with their first hypothesis, participants were more likely to interpret implausible sentences literally if making the sentence plausible required multiple changes. Participants were also more likely to interpret implausible sentences literally if making the implausible sentence plausible required insertion of a word rather than deletion of a word—a result that supports the second hypothesis. Their third hypothesis was supported by evidence showing that an increase in the perceived level of noise led to an increase in the rate of non-literal interpretation of implausible sentences. Finally, their results showed that increased prevalence of implausible sentences led to increased rates of literal interpretation of implausible sentences, thereby lending support to the fourth hypothesis. These results provide further evidence that comprehenders integrate pre-existing expectations with the probability of noise and make rational inferences about the intended meaning of a sentence.

Building on Gibson et al.'s results, Poppels and Levy (2016) found that the noise model is sensitive not only to deletions and possibly insertions but also to function word exchanges. Furthermore, Poppels & Levy (2016), Liu et al. (2020a), and Keshev & Meltzer-Asscher (2021) showed that comprehender-noise models are sensitive to sentence structure. And Ryskin et al. (2018) observed that people are sensitive to the nature of the noise in the local environment that they encounter: participants' assumptions about which edits were most likely were dependent on the types of errors that were present in other sentences within the same experiment.

A key limitation of these existing studies on noisy channel language comprehension is that they all investigate single sentences with no preceding context. In typical language use, there is usually a context for any utterance. For example, the sentences that make up this paragraph come together to construct an overarching meaning that depends on, and affects, the meaning of each individual sentence. Sentences are rarely interpreted on their own: comprehenders constantly extract meanings from a sentence while taking context into consideration. Individual words are recalled more accurately when they form a valid English  $n$ -gram than when they are drawn at random, and the accuracy increases as  $n$  increases (Miller & Selfridge, 1950). Preceding context can exert an important influence on how a sentence is understood online. Sentences are processed more readily when they are coherent with the preceding discourse than when they are not (e.g., Albrecht & O'Brien, 1993; Camblin, Gordon, & Swaab, 2007; Bader & Meng, 2023). Sentence anomalies such as, “when the plane crashes, where should the survivors be buried?”, are less likely to be detected when they are preceded by a coherent context (Barton & Sanford, 1993). A supportive context Furthermore, a strong discourse context can even override language and world-knowledge based expectations. Using event-related potentials as a dependent measure, Kutas and Hillyard (1980) observed an N400 effect when participants were presented with semantically implausible sentences. However, Nieuwland and van Berkum (2006) showed that semantically implausible sentences such as, “*The girl comforted the clock*” were no longer processed as implausible when preceded by a cartoon-like discourse context in which a girl interacts with a clock who is feeling sad. In this supportive context, the N400 effect was reduced, and in some cases reversed (see Kuperberg, 2007, for review).

Indeed, in another ERP experiment, Nieuwland and van Berkum (2005) provide further evidence for the hypothesis that a supportive context can alter the comprehender's interpretation

of an implausible sentence. Nieuwland and van Berkum reported a P600 in response to semantic (e.g., animacy) violations that occur in extended discourse contexts. Participants were presented with short stories (e.g., about a tourist and his suitcase where both entities are mentioned several times). In the critical sentences like *Next, the woman told the tourist / suitcase...*, a P600 (instead of an N400, as is typically observed when there is no preceding context) was observed at the word *suitcase*. Consistent with the proposal that the P600 waveform indexes error correction and ensues whenever it is likely that the received message was corrupted by noise (Ryskin et al., 2021), a P600 ensues in this case because a word substitution error when both lexical entries are highly active is a probable production error and is treated as such by the comprehenders.

## Present Research

In the present work, we replicate and extend Gibson et al. (2013) to investigate the role of discourse context on noisy-channel inferences. In particular, we predict that the discourse context acts on the language prior,  $P(s_i)$ . Semantically implausible but syntactically licensed test sentences (e.g., “*The mother gave the candle the daughter.*”) were preceded by a supportive context (e.g., “*When the power outage happened, the daughter asked the mother for a candle. The mother found a candle in the kitchen cabinet.*”), a non-supportive context (e.g., “*The niece missed the sister when she went to overnight camp. The father offered to buy the truck from the uncle.*”). Test sentences consisted of two types of syntactic alternations from Gibson et al. (2013), each involving a pair of sentence constructions: one with the double-object (DO) construction and the prepositional-phrase-object (PO) construction, and the other with the active construction and the passive construction. Critically, we predicted that participants would be more likely to infer a plausible semantic alternative, rather than interpret the sentences literally,

when they were preceded by a supportive context, compared with when they were preceded by a non-supportive context, or when they were not preceded by any context, because the prior probability of the alternative was increased relative to when the context was not supportive or when there was no context. Such a finding would indicate that the local discourse context plays an important role in shaping the language prior, and, as a result, the rate of noisy-channel inferences during sentence comprehension. In addition, we also expected to replicate key findings in Gibson et al. (2013), namely: 1) DO sentences are less likely to be interpreted literally than PO sentences; 2) DO/PO sentences are less likely to be interpreted literally than active/passive sentences .

From a noisy-channel perspective, a discourse context can affect how a sentence is interpreted by altering the prior probability of the intended meaning  $p(s_i)$ . To illustrate this, we can express  $p(s_i)$  as an integral shown below:

$$p(s_i) = \int_c p(s_i|c) \cdot p(c)dc \quad (2)$$

In Equation (2),  $p(s_i|c)$  indicates the probability of the intended meaning  $s_i$  under a context  $c$ , and  $p(c)$  represents the probability of context  $c$ . A supportive context will yield a relatively high  $p(s_i|c)$  compared with a non-supportive context, in that the context will provide referents for the target sentence, hence increasing the expectation of the intended meaning. For example, consider two types of context preceding an implausible sentence “The girl tossed the apple the boy” in Table 1: the first context is supportive: “The boy and the girl went apple picking together. The girl picked an apple that the boy wanted”. The second context is non-supportive: “The aunt told the nephew she would miss him while he was on vacation. The magician pulled his hat out of the trunk”. Given the first context, a comprehender will expect that the following sentence is about the girl giving the apple to the boy in some form. Therefore, the



conditional probability  $p(s_i|c)$  for the plausible alternative will be high. In contrast, the second context does not establish any expectation that the next sentence would be a girl giving an apple to a boy. Hence, the conditional probability  $p(s_i|c)$  for the plausible alternative will be relatively low.

In addition, Equation (2) can be further broken down into two terms:

$$p(s_i) = \int_{c \in S} p(s_i|c) \cdot p(c)dc + \int_{c \in NS} p(s_i|c) \cdot p(c)dc \quad (3)$$

The first term (henceforth  $C_S$ ) represents the probability of the intended meaning contributed by a supportive context, and similarly, the second term (henceforth  $C_{NS}$ ) represents the probability of the intended meaning contributed by a non-supportive context. In this study, we vary the relative magnitude of  $p(c)$  across  $C_S$  and  $C_{NS}$  in two between-participant experimental conditions: a supportive context condition and a non-supportive context condition (See Section 2.1). In the supportive context condition, participants read target sentences preceded by a supportive context, and therefore, their expectation of a supportive context  $p(c \in \text{supportive})$  should be relatively high, whereas their expectation of a non-supportive context  $p(c \in \text{nonsupportive})$  should be relatively low (see the left panel of Figure 1). As discussed above,  $p(s_i|c)$  is relatively high when  $c$  is supportive, whereas  $p(s_i|c)$  is relatively low when  $c$  is non-supportive. Taken together, the supportive context condition places a higher weight on  $p(s_i|c)$  in  $C_S$  in Equation (3), where it is relatively high, and a lower weight on  $p(s_i|c)$  in  $C_{NS}$ , where it is relatively low. Therefore, from Equation (3), the probability of a plausible interpretation under a supportive context is relatively high. In contrast, in the nonsupportive context condition, participants read target sentences preceded only by non-supportive contexts, and in this case, their expectation of a nonsupportive context  $p(c \in \text{nonsupportive})$  should be higher than that under the supportive context condition (see the right panel of Figure 2). Taken

together, compared with the supportive context condition, the non-supportive context condition places a relatively lower weight on  $p(s_i|c)$  in  $C_s$  in Equation (3), and a relatively higher weight on  $p(s_i|c)$  in  $C_{NS}$ , and therefore, the probability of the plausible interpretation under a non-supportive context is relatively low<sup>1</sup>.

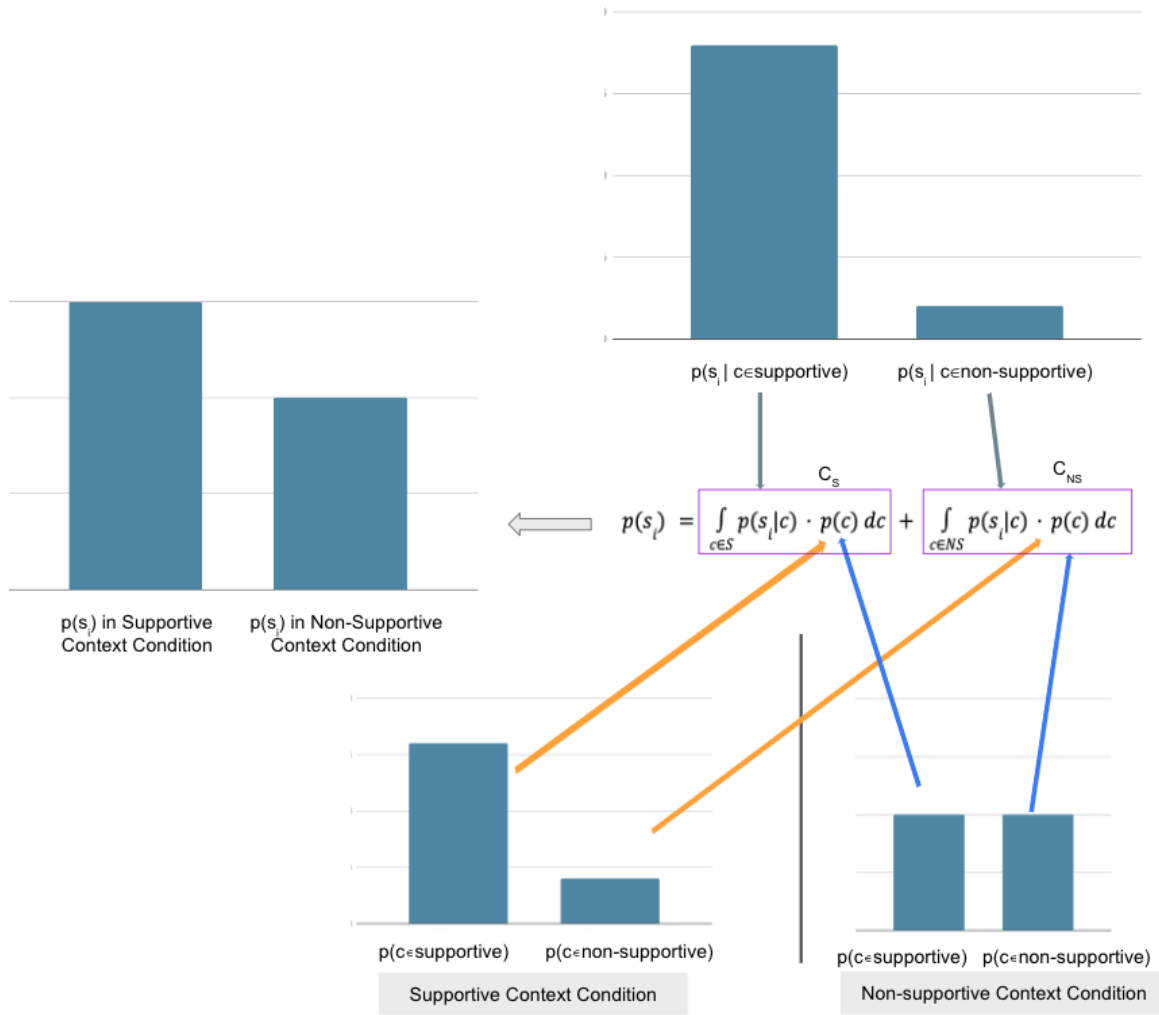


Figure 1. A qualitative illustration of the distribution of the expectation of different kinds of context in the Supportive Context condition and the Non-Supportive Context condition. Top

<sup>1</sup> Under the No Context condition, we expect  $p(s_i)$  to be between the  $p(s_i)$  under the Non-Supportive context condition and that under the Supportive context condition. This is because when target sentences are not preceded by any context, comprehenders may assign their own weights on supportive and non-supportive contexts. Their weights on these conditions should be within the ranges of those under the Supportive and Non-Supportive context conditions, since these two conditions are two of the extreme cases.

panel: a supportive context gives rise to a higher expectation of the target sentence than a non-supportive context. Bottom panel: different context conditions give participants different expectations of supportive and non-supportive context, in that the expectation of a supportive context is higher in the supportive context condition, where participants read target sentences preceded only by a supportive context, than in the non-supportive context condition, where participants read target sentences preceded only by a non-supportive context. Note that both charts are qualitative: the point is simply that there is a higher expectation for a non-supportive context under the Non-Supportive Context condition, compared to under the Supportive Context condition. Left panel: according to Equation (3), the probability of a plausible interpretation  $p(s_i)$  is higher under the supportive context condition than under a non-supportive context condition.

The current study is organized as follows. In Experiment 1, we investigate the effect of a supportive context versus a non-supportive context on sentences under 2 types of syntactic alternations (DO/PO and active/passive). Experiment 2 is a replication of Experiment 1 with twice as many subjects. Experiment 3 is an extension of the previous 2 experiments in that we add a third condition where test sentences are not preceded by any context, as in previous studies, to control for the possibility that results in Exps 1-2 were simply caused by participants paying more attention to the target sentences under a non-supportive context. The no context condition also connects the current study with previous ones by Gibson and colleagues, where sentences were not preceded by any context. Experiment 4 is an exact replication of Experiment 3 due to an oversight error in the pre-registration for Experiment 3.

## 2. Experiment 1

### 2.1. Methods

In this pre-registered experiment, 240 participants were recruited through Amazon's Mechanical Turk platform and compensated 4.00 USD for successful completion. Participation was restricted to users who had U.S.-based IP addresses and 95% approval ratings. Participants were also asked to identify both their native languages and countries of origin. Payment was not linked to participants' responses to demographic questions, but data from participants who were not native English speakers or not from the U.S. was excluded from analysis.

Participants were presented with a questionnaire containing 68 context-sentence-question sets (20 critical, 48 filler), each of which consisted of three sentences and a comprehension question. All sentences and comprehension questions were shown simultaneously, and participants were allowed to read each trial as many times as necessary in order to avoid imposing any memory burden during the task. Participants were not informed about the nature of the experiment, only that the questionnaires contained sentences and comprehension questions for them to read and answer. Critical trials consisted of two context sentences and a test sentence, followed by a comprehension question. For each participant, ten of the test trials were implausible and ten were plausible. Half of the participants were exposed to sentences under DO/PO construction; the other half of the participants were exposed to sentences under the active/passive construction. Example stimuli in each construction and each context condition are presented in Table 1.

| <b>Constructions</b>      | <b>Double Object (DO)</b>                                                                       | <b>Prepositional-Phrase Object (PO)</b> | <b>Active</b>                                                              | <b>Passive</b> |
|---------------------------|-------------------------------------------------------------------------------------------------|-----------------------------------------|----------------------------------------------------------------------------|----------------|
| <i>Supportive Context</i> | The boy and the girl went apple picking together. The girl picked an apple that the boy wanted. |                                         | The boy ordered a pizza at the restaurant. The pizza was cooked by a chef. |                |

|                                             |                                                                                                                     |                                       |                                                                                                         |                                 |
|---------------------------------------------|---------------------------------------------------------------------------------------------------------------------|---------------------------------------|---------------------------------------------------------------------------------------------------------|---------------------------------|
| (Experiments 1-4)                           |                                                                                                                     |                                       |                                                                                                         |                                 |
| Non-Supportive Context<br>(Experiments 1-4) | The aunt told the nephew she would miss him while he was on vacation. The magician pulled his hat out of the trunk. |                                       | When the power outage happened, the husband asked the wife for a candle. Later, the hammer was missing. |                                 |
| No Context<br>(Experiment 3-4)              | N.A.                                                                                                                |                                       | N.A.                                                                                                    |                                 |
| Plausible Target Sentence                   | The girl tossed the boy the apple.                                                                                  | The girl tossed the apple to the boy. | The boy ate the pizza.                                                                                  | The pizza was eaten by the boy. |
| Implausible Target Sentence                 | The girl tossed the apple the boy.                                                                                  | The girl tossed the boy to the apple. | The pizza ate the boy.                                                                                  | The boy was eaten by the pizza. |
| Comprehension Question                      | Did the apple receive something/someone?                                                                            |                                       | Did the pizza eat something/someone?                                                                    |                                 |

*Table 1.* Example stimuli: plausible and implausible target sentences across sentence construction, with supportive contexts, non-supportive contexts, or no contexts.

For each sentence construction, three questionnaires were generated: one containing only supportive contexts, one containing only non-supportive contexts, and one containing no context. For the supportive context condition, two sentences of supporting context were written to precede each of the 68 sentences (20 critical, 48 filler). In order to generate non-supportive contexts, the sentences written to establish a supportive context were randomly paired (so that no two context-establishing sentences related to each other), and randomly assigned to each of the original 68 sentences (20 critical, 48 filler). For each sentence construction and context-type pair, a Latin square was utilized to generate 60 lists from 20 test items and 48 filler items so that there were at least 2 filler items between consecutive test items. If participants were to interpret every sentence literally, the answer would be “Yes” for half of the sentences and “No” for the other half. In the trials involving an implausible sentence, the correct answer is the literal interpretation of that sentence. Hence, a correct response indicates that the participant interpreted the sentence literally, whereas an incorrect response suggests that the participant made an

inference for the non-literal, plausible meaning. Behaviorally, a high literal response rate means the same as a low inference rate.

In Experiment 1, 60 participants received a supportive-context questionnaire with DO/PO sentences; 60 participants received a supportive-context questionnaire with active/passive target sentences; 60 participants received a non-supportive-context questionnaire with DO/PO sentences; 60 participants received a non-supportive-context questionnaire with active/passive target sentences. The full set of materials used is available at [osf.io/s7ck2](https://osf.io/s7ck2). The pre-registration of this experiment is available at <https://osf.io/n3kgr> (Experiment 1).

## 2.2. Results and Discussion

Data were included from a total of 222 participants who correctly answered a minimum of 75% of the plausible questions (plausible target sentences as well as fillers), who reported that English was their native language, and who identified the U.S. as their country of origin. For semantically plausible sentences, across construction type and context conditions, the rate of literal interpretation was above 90% and was not analyzed further. Proportions of literal interpretations for implausible test sentences for Experiment 1 are illustrated in Figure 2.

Regarding the effect of context in each sentence construction, proportions of literally interpreted sentences were analyzed using two mixed-effects logistic regression models, one for each sentence construction (DO/PO, active/passive), using the lme4 package in R (Bates, Maechler, Bolker, & Walker, 2015). Context condition (supportive vs. non-supportive) was entered as a sum coded fixed effect. Items and participants were entered as random intercepts with random by-items slopes for context condition.<sup>2</sup> The fixed effect parameter estimates of

---

<sup>2</sup> In the statistical analysis for the effect of context condition, we use the formula:  $\text{literal interpretation} \sim \text{context} + (1 \mid \text{participant}) + (1 + \text{context} \mid \text{item})$ . The default optimizer is bobyqa. If the model yields singularity, we run the model again with all available optimizers. If the first 3 significant digits of the values given by all the optimizers are

these models are summarized in Table 2 (See Comparison 1).<sup>3</sup> For both sentence construction types, we observe a numerical but not statistical trend that participants are more likely to adopt the non-literal interpretation when they encounter implausible sentences preceded by a supportive context, compared with when preceded by a non-supportive context ( $p > 0.09$ ). The consistent numerical pattern suggests that the supportive context may elicit more noisy-channel inferences than the non-supportive context condition, but the current design may have been under-powered to reliably detect this effect.

Within sentences with the DO/PO construction, in each context condition, we used 2 mixed-effects logistic regression models, one for each context condition (supportive, non-supportive). Sentence type (DO vs. PO) was entered as a sum coded fixed effect. Items and participants were entered as random intercepts with random by-item and by-participant slopes for sentence type<sup>4</sup>. The fixed effect parameter estimates of these models are summarized in Comparison 2 in Table 2. In both context conditions, we found that participants were more likely to interpret PO sentences literally than DO sentences ( $p$ 's  $< 1.84\text{e-}6$ ), which is consistent with the prediction that sentences made implausible with insertion (PO) are more likely to be interpreted literally than those made implausible with deletion (DO), replicating previous studies (e.g. Gibson et al., 2013; Poppels & Levy, 2016). Similarly, within sentences with the active/passive construction, we found that within each context condition, active sentences are approximately as

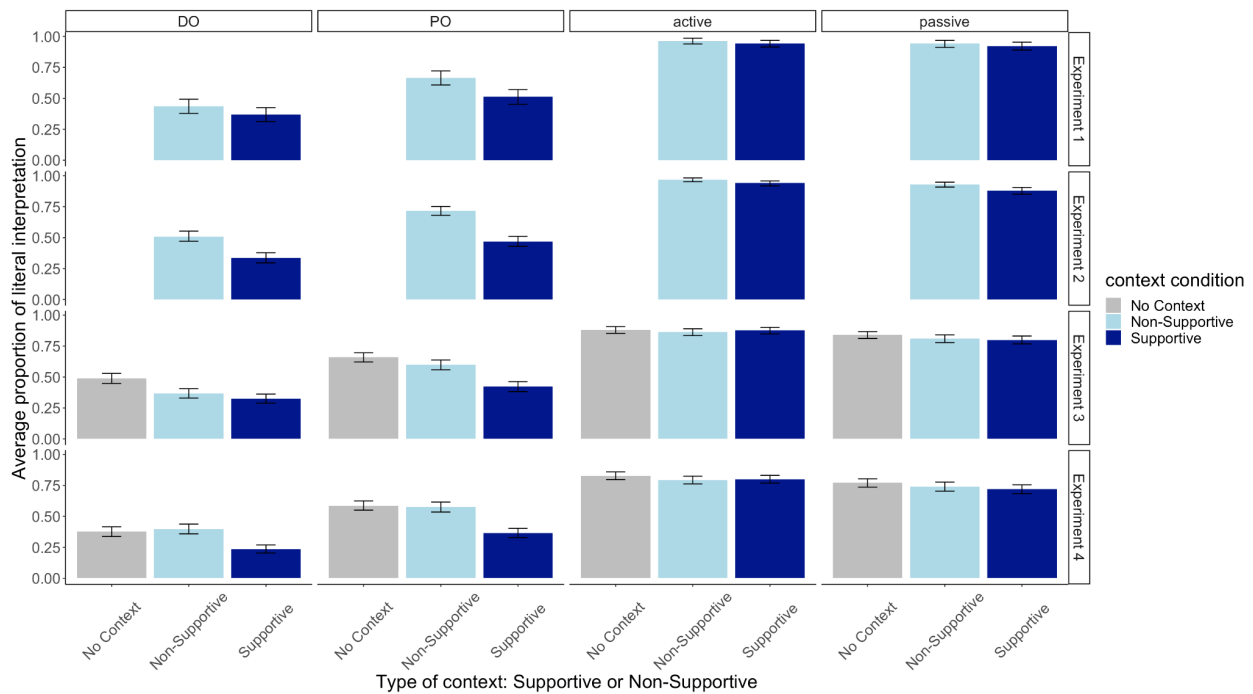
---

identical, we report the values given by the bobyqa optimizer. Otherwise, a simpler model [literal interpretation ~ context + (1 | participant) + (1 | item)] is used.

<sup>3</sup> We also analyzed our data with a Bayesian approach similar to Experiments 3-4, using the MCMCglmm package in R (Hadfield, 2010). We obtained similar results. We report the results with lme4 in the text since it is what was planned in the pre-registration for Experiments 1-2.

<sup>4</sup> Similarly, in the statistical analysis for the effect of sentence type (DO vs. PO) within each context condition, we use the formula: literal interpretation ~ sentence type + (1 + sentence type | participant) + (1 + sentence type | item). The default optimizer is bobyqa. If the model yields singularity, we run the model again with all available optimizers. If the first 3 significant digits of the values given by all the optimizers are identical, we report the values given by the bobyqa optimizer. Otherwise, a simpler model [literal interpretation ~ sentence type + (1 | participant) + (1 | item)] is used.

likely to be interpreted literally as passive sentences, since both types of sentences are made implausible by exchanges. However, although it does not reach significance, passive sentences are consistently less likely to be interpreted literally than active sentences, possibly because of a structural prior effect: passive sentences are used less widely than active sentences, and thus the probability of the passive plausible alternative sentence  $p(s_i)$  will be lower than that of the active plausible alternative sentence.



*Figure 2.* Rates of literal interpretation of DO/PO/active/passive sentences for implausible sentences from 4 experiments, in Non-Supportive contexts (light blue), Supportive contexts (dark blue), and no contexts (gray). The error bars indicate the 95% bootstrapped confidence interval.

### 3. Experiment 2

#### 3.1. Methods

The goal of Experiment 2 was to conduct a pre-registered (link at <https://osf.io/zn579>) replication of Experiment 1 with a larger sample size. The materials and the procedures were



identical to those in Experiment 1 except twice as many participants (120) were recruited for each between-subject condition.

### 3.2. Results and Discussion

Data were included from a total of 456 participants who correctly answered a minimum of 75% of the plausible questions (plausible target sentences as well as fillers), who reported that English was their native language, and who identified the U.S. as their country of origin. For semantically plausible sentences, across construction type and context conditions, the rate of literal interpretation was above 90% and was not analyzed further. Proportions of literal interpretations for implausible test sentences for Experiment 2 are illustrated in Figure 2.

| <b>Experiment 1</b>                                       |                        |                    |               |       |         |             |
|-----------------------------------------------------------|------------------------|--------------------|---------------|-------|---------|-------------|
| Comparison 1: No context vs. Supportive context (n = 222) |                        |                    |               |       |         |             |
| Sentence Alternation                                      | Non-Supportive Context | Supportive Context | $\beta$ value | SE    | z value | p value     |
| DO/PO                                                     | 0.548                  | 0.438              | -0.679        | 0.400 | -1.696  | 0.090       |
| Active/Passive                                            | 0.955                  | 0.932              | -0.572        | 0.423 | -1.349  | 0.177       |
| Comparison 2: Within each alternation, DO vs. PO          |                        |                    |               |       |         |             |
| Context condition                                         | DO                     | PO                 | $\beta$ value | SE    | z value | p value     |
| Non-Supportive Context (n = 111)                          | 0.432                  | 0.664              | 1.620         | 0.250 | 6.474   | 9.52e-11*** |
| Supportive Context (n = 111)                              | 0.364                  | 0.513              | 1.135         | 0.238 | 4.771   | 1.84e-6***  |
| Comparison 3: Within each alternation, active vs. passive |                        |                    |               |       |         |             |
| Context condition                                         | Active                 | Passive            | $\beta$ value | SE    | z value | p value     |
| Non-Supportive Context (n = 111)                          | 0.967                  | 0.942              | -0.736        | 0.479 | -1.538  | 0.124       |
| Supportive Context (n = 111)                              | 0.942                  | 0.921              | -0.373        | 0.359 | -1.046  | 0.296       |
| <b>Experiment 2</b>                                       |                        |                    |               |       |         |             |
| Comparison 1: No context vs. Supportive context (n = 456) |                        |                    |               |       |         |             |
| Sentence Alternation                                      | Non-Supportive Context | Supportive Context | $\beta$ value | SE    | z value | p value     |
| DO/PO                                                     | 0.611                  | 0.402              | -1.543        | 0.306 | -5.042  | 4.6e-7***   |
| Active/Passive                                            | 0.945                  | 0.909              | -0.582        | 0.322 | -1.807  | 0.071       |
| Comparison 2: Within each alternation, DO vs. PO          |                        |                    |               |       |         |             |
| Sentence Construction                                     | DO                     | PO                 | $\beta$ value | SE    | z value | p value     |
| Non-Supportive Context (n = 228)                          | 0.507                  | 0.715              | 1.423         | 0.161 | 8.831   | <2e-16***   |
| Supportive Context (n = 228)                              | 0.334                  | 0.470              | 1.168         | 0.173 | 6.736   | 1.63e-11*** |
| Comparison 3: Within each alternation, active vs. passive |                        |                    |               |       |         |             |
| Context condition                                         | Active                 | Passive            | $\beta$ value | SE    | z value | p value     |

|                                  |       |       |        |       |        |           |
|----------------------------------|-------|-------|--------|-------|--------|-----------|
| Non-Supportive Context (n = 228) | 0.964 | 0.926 | -0.857 | 0.298 | -2.875 | 0.004**   |
| Supportive Context (n = 228)     | 0.938 | 0.880 | -1.004 | 0.260 | -3.859 | <0.001*** |

*Table 2.* Rates of literal interpretation of implausible sentences by construction and context and estimates of the effect of context (Comparison 1) and syntax within each context condition (Comparison 2) on these rates from logistic mixed-effects models in Experiments 1 and 2. The symbols \* and \*\*\* indicate that the p value is smaller than 0.05 and 0.001, respectively.

We adopted the same analysis procedures as in Experiment 1. The fixed-effect parameter estimates of these models are summarized in Table 2. For implausible DO/PO test sentences, participants were less likely to interpret them literally in the supportive context condition (0.402) than in the non-supportive context condition (0.611;  $\beta = -1.543$ ;  $p < 0.001$ ), whereas for implausible active/passive test sentences, there was no significant effect of context ( $p=0.071$ ). In addition, similar to Experiment 1, among DO/PO sentences we observe a significant effect of sentence type ( $p$ 's  $< 1.61e-11$ ) in both supportive and non-supportive context conditions, and on the other hand, the difference in literal interpretation rate among active/passive sentences also reached significance in Experiment 2, possibly because of the difference in structural prior mentioned above. These results show that with twice as many as participants, we were able to observe a significant decrease in literal interpretation rate in DO/PO sentences when test sentences are preceded by a supportive context than by a non-supportive context, consistent with our prediction that a supportive context increases the prior probability of the non-literal but plausible interpretation.

#### 4. Experiment 3

Experiment 3 is a further extension of the previous two experiments and serves two purposes. First, we adopted a Bayesian approach in data analysis to avoid the issue of singularity that we encountered when using linear mixed effects analyses using lme4. In addition, we added

an experimental condition where there is no context as an additional baseline, since results in Experiments 1 and 2 did not preclude the possibility that readers, when seeing a non-supportive context, might scrutinize the target sentence more carefully and therefore might be more likely to interpret it literally. Furthermore, the No Context condition serves as a replication of previous studies in the literature (Gibson et al., 2013; Poppels & Levy, 2016).

#### **4.1. Methods**

The materials and the procedures were identical to those in the previous two experiments except that two more types of questionnaires were added: one containing test sentences in the active/passive construction with no preceding context, and another containing test sentences in the DO/PO construction with no preceding context. We planned to recruit at least 120 participants for each questionnaire type ( $6 * 120 = 720$  in total).

#### **4.2. Results**

1201 participants were recruited on Amazon Mechanical Turk. Data were included from 731 participants who correctly answered a minimum of 75% of the plausible questions (plausible target sentences as well as fillers), who reported that English was their native language, and who identified the U.S. as their country of origin. For semantically plausible sentences, across construction type and context conditions, the rate of literal interpretation was above 90% and was not analyzed further. Proportions of literal interpretations for implausible test sentences for this experiment are illustrated in Figure 1.

For comparison across context types, proportions of literally interpreted sentences were analyzed using four mixed-effects logistic regression models, two for each syntactic alternation ({supportive context vs. no context, supportive context vs. non-supportive context} x {DO/PO, active/passive}), using the MCMCglmm package in R (Hadfield, 2010), with an uninformative

prior for each of the parameters in the model. The number of MCMC iterations for each model is set to be 10000, with thinning set to be every 10 iterations. Context condition (supportive vs. non-supportive and supportive vs. no context) was entered as a sum coded fixed effect. Items and participants were entered as random intercepts with random by-items slopes for context condition, which is the maximally complex random effect structure based on the experimental design (Bates et al., 2013) . The formula is shown below in the typical lme4 syntax (Bates, Maechler, Bolker, & Walker, 2015).

$$\text{literal response} \sim \text{context condition} + (1 + \text{context condition} \mid \text{item}) + (1 \mid \text{participant}) \quad (4)$$

We also analyzed the proportions of literally interpreted sentences within each context condition for each of the construction pairs ({DO vs. PO, Active vs. Passive} x {no context, supportive context, non-supportive context}) using MCMCglmm in R with the same iterations per model and the same thinning interval. In this analysis, construction was coded as a fixed effect. Items and participants were entered as random intercepts, and the slope of construction is allowed to vary by both items and participants. The formula is shown below in a typical lme4 syntax.

$$\text{literal response} \sim \text{construction} + (1 + \text{construction} \mid \text{item}) + (1 + \text{construction} \mid \text{participant}) \quad (5)$$

Then, we analyzed the proportion of literally interpreted sentences within each context condition across construction pairs. In this analysis, the construction pair (active-passive vs. DO-PO) was coded as a fixed effect. Items and participants were entered as random intercepts. The

formula is shown below in a typical lme4 syntax. All these formulas are the maximally complex random effect structure (Bates et al., 2013) based on the experimental design and the analysis.

$$\text{literal response} \sim \text{construction pair} + (1 \mid \text{item}) + (1 \mid \text{participant}) \quad (6)$$

The fixed effect parameter estimates of these models are summarized in Table 3. For implausible DO/PO sentences, participants were less likely to interpret them literally in the supportive context condition (0.373) than in the no context condition (0.574;  $\beta = -1.676$ ;  $p_{\text{MCMC}} < 0.001$ ) and in the non-supportive context condition (0.484;  $\beta = -1.007$ ;  $p_{\text{MCMC}} = 0.009$ ). For implausible active/passive sentences, the literal interpretation rates were approximately the same in all 3 context conditions ( $p > 0.446$ ).

| <b>Experiment 3</b>                                                   |                        |                    |                        |                |                 |           |
|-----------------------------------------------------------------------|------------------------|--------------------|------------------------|----------------|-----------------|-----------|
| Comparison 1: No context vs. Supportive context (n = 731)             |                        |                    |                        |                |                 |           |
| Sentence Construction                                                 | No Context             | Supportive Context | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| DO/PO                                                                 | 0.574                  | 0.373              | -1.676                 | -2.357         | -0.954          | <0.001*** |
| Active/Passive                                                        | 0.859                  | 0.838              | -0.332                 | -1.207         | 0.526           | 0.446     |
| Comparison 2: Non-Supportive Context vs. Supportive Context (n = 731) |                        |                    |                        |                |                 |           |
| Sentence Construction                                                 | Non-Supportive Context | Supportive Context | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| DO/PO                                                                 | 0.484                  | 0.373              | -1.007                 | -1.693         | -0.215          | 0.009**   |
| Active/Passive                                                        | 0.836                  | 0.838              | 0.052                  | -0.778         | 0.893           | 0.871     |
| Comparison 3: Within each construction; PO vs. DO                     |                        |                    |                        |                |                 |           |
| Context condition                                                     | DO                     | PO                 | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| No context (n = 122)                                                  | 0.489                  | 0.659              | 1.629                  | 1.004          | 2.235           | <0.001*** |
| Non-Supportive (n = 121)                                              | 0.368                  | 0.599              | 2.065                  | 1.569          | 2.649           | <0.001*** |
| Supportive (n = 120)                                                  | 0.324                  | 0.422              | 1.002                  | 0.602          | 1.428           | <0.001*** |
| Comparison 4: Within each construction; Active vs. Passive            |                        |                    |                        |                |                 |           |
| Context condition                                                     | Active                 | Passive            | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| No context (n = 125)                                                  | 0.877                  | 0.841              | -0.474                 | -1.186         | 0.200           | 0.189     |
| Non-Supportive (n = 123)                                              | 0.863                  | 0.809              | -0.860                 | -1.404         | -0.262          | 0.002**   |
| Supportive                                                            | 0.876                  | 0.799              | -1.060                 | -1.508         | -0.530          | <0.001*** |

| (n = 120)                          |                |       |                        |                |                 |           |
|------------------------------------|----------------|-------|------------------------|----------------|-----------------|-----------|
| Comparison 5: Across constructions |                |       |                        |                |                 |           |
| Context condition                  | Active/Passive | DO/PO | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| No context (n = 247)               | 0.859          | 0.574 | -3.035                 | -3.906         | -2.289          | <0.001*** |
| Non-Supportive (n = 244)           | 0.836          | 0.484 | -3.402                 | -4.345         | -2.563          | <0.001*** |
| Supportive (n = 240)               | 0.838          | 0.373 | -4.525                 | -5.502         | -3.689          | <0.001*** |

*Table 3.* Rates of literal interpretation of implausible sentences by construction and context (Comparisons 1 and 2), within alternation by context (Comparisons 3 and 4), and across alternations by context (Comparison 5) in Experiment 3. The estimates of the effect of context on these rates are from Bayesian logistic mixed-effects models (\*, \*\*, and \*\*\* indicate that the p\_MCMC value is smaller than 0.05, 0.01, and 0.001, respectively).

Within each alternation, in all context conditions, participants made more inferences when they were presented with implausible DO sentences (0.489, 0.368, 0.324 in no context, non-supportive context, supportive context conditions, respectively) compared with when they were presented with implausible PO sentences (0.659, 0.599, 0.422 respectively;  $\beta$ s > 1.002, p\_MCMCs < 0.001). On the other hand, when the test sentences were preceded with no context, participants were as likely to interpret implausible Active sentences literally (0.877) as they were to interpret implausible Passive sentences literally (0.841;  $p = 0.189$ ). However, when the test sentences were preceded with non-supportive context or supportive context, participants were more likely to interpret implausible Active sentences literally (0.863, 0.876, respectively) than implausible Passive sentences (0.809, 0.799, respectively;  $\beta$ s < -0.860, p\_MCMCs < 0.002).

Finally, across alternations, in all context conditions, participants made more inference when they were presented with sentences in DO/PO construction (0.574 for the No Context Condition, 0.484 for the Non-Supportive Context Condition, and 0.373 for the Supportive Context Condition, respectively) than when they were presented with those in Active/Passive construction (0.859 for the No Context Condition, 0.836 for the Non-Supportive Context

Condition, and 0.838 for the Supportive Context Condition, respectively;  $\beta_s < -3.035$ ,  $p_{\text{MCMCs}} < 0.001$ ).

## **5. Experiment 4**

Upon finishing collecting data for Experiment 3, we spotted an error in one of the key hypotheses stated in the pre-registration. As we found it unjustifiable to simply amend it, we decided to fully replicate Experiment 3 with the correct prediction written in a new pre-registration.

### **5.1. Methods**

The materials and the procedures are identical to those in Experiment 3, where participants were assigned into one of the 6 between-subject conditions, each corresponding to one of the three context conditions (supportive, non-supportive, no context) along with one of the two sentence constructions (active/passive, DO/PO). The pre-registration is available at <https://osf.io/5rtdu>.

### **5.2. Results**

1365 participants were recruited from Amazon Mechanical Turk. Data were included from 729 participants who correctly answered a minimum of 75% of the plausible questions (plausible target sentences as well as fillers), who reported that English was their native language, and who identified the U.S. as their country of origin. Those who have participated in Experiment 3 were excluded from the study. For semantically plausible sentences, across construction type and context conditions, the rate of literal interpretation was above 85% and was not analyzed further. Proportions of literal interpretations for implausible test sentences for this experiment are illustrated in Figure 2.

The data analysis procedures are identical to those in Experiment 3. The results, along with the fitted main effects, are presented in Table 4.

| <b>Experiment 4</b>                                                   |                        |                    |                        |                |                 |           |
|-----------------------------------------------------------------------|------------------------|--------------------|------------------------|----------------|-----------------|-----------|
| Comparison 1: No context vs. Supportive context (n = 729)             |                        |                    |                        |                |                 |           |
| Sentence Construction                                                 | No Context             | Supportive Context | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| DO/PO                                                                 | 0.481                  | 0.301              | -1.376                 | -1.937         | -0.789          | <0.001*** |
| Active/Passive                                                        | 0.799                  | 0.760              | -0.375                 | -1.194         | 0.399           | 0.334     |
| Comparison 2: Non-Supportive Context vs. Supportive Context (n = 729) |                        |                    |                        |                |                 |           |
| Sentence Construction                                                 | Non-Supportive Context | Supportive Context | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| DO/PO                                                                 | 0.486                  | 0.301              | -1.470                 | -2.154         | -0.860          | <0.001*** |
| Active/Passive                                                        | 0.768                  | 0.760              | -0.150                 | -1.018         | 0.688           | 0.774     |
| Comparison 3: Within each construction pair; DO vs. PO                |                        |                    |                        |                |                 |           |
| Context condition                                                     | DO                     | PO                 | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| No context (n = 123)                                                  | 0.376                  | 0.585              | 1.672                  | 1.216          | 2.204           | <0.001*** |
| Non-Supportive (n = 120)                                              | 0.397                  | 0.575              | 1.604                  | 1.054          | 2.150           | <0.001*** |
| Supportive (n = 123)                                                  | 0.237                  | 0.365              | 1.037                  | 0.537          | 1.469           | <0.001*** |
| Comparison 4: Within each construction pair; Active vs. Passive       |                        |                    |                        |                |                 |           |
| Context condition                                                     | Active                 | Passive            | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| No context (n = 121)                                                  | 0.827                  | 0.770              | -0.656                 | -1.137         | -0.163          | 0.02*     |
| Non-Supportive (n = 121)                                              | 0.795                  | 0.740              | -0.625                 | -1.084         | -0.159          | 0.008**   |
| Supportive (n = 121)                                                  | 0.800                  | 0.719              | -0.905                 | -1.355         | -0.462          | <0.001*** |
| Comparison 5: Across construction pairs                               |                        |                    |                        |                |                 |           |
| Context condition                                                     | Active/Passive         | DO/PO              | posterior mean $\beta$ | 2.5 percentile | 97.5 percentile | p_MCMC    |
| No context (n = 243)                                                  | 0.799                  | 0.481              | -2.799                 | -3.742         | -2.092          | <0.001*** |
| Non-Supportive (n = 241)                                              | 0.768                  | 0.486              | -2.758                 | -3.705         | -1.729          | <0.001*** |
| Supportive (n = 244)                                                  | 0.760                  | 0.301              | -3.953                 | -4.855         | -3.140          | <0.001*** |

*Table 4.* Rates of literal interpretation of implausible sentences by construction and context (Comparisons 1 and 2), within construction pair by context (Comparisons 3 and 4), and across construction pair by context (Comparison 5) in Experiment 3. The estimates of the effect of context on these rates are from Bayesian logistic mixed-effects models (\*, \*\*, and \*\*\* indicate that the p\_MCMC value is smaller than 0.05, 0.01, and 0.001, respectively).

The results in Experiment 4 largely replicated those in Experiment 3. Within each construction pair, sentences under DO/PO construction are less likely to be interpreted literally



when they are preceded by supportive context (0.301), than when they are preceded by no context (0.481,  $\beta = -1.376$ ,  $p\_MCMC < 0.001$ ). Similarly, DO/PO sentences preceded by supportive context are also less likely to be interpreted literally than those preceded by non-supportive context (0.486,  $\beta = -1.470$ ,  $p\_MCMC < 0.001$ ). In contrast, as it was also the case in Experiment 3, sentences under Active/Passive construction preceded by supportive context (0.760) are about as likely to be interpreted literally as those preceded by non-supportive context (0.768,  $p\_MCMC = 0.774$ ) and those preceded by no context (0.799,  $p\_MCMC = 0.334$ ).

Compared with PO sentences, DO sentences are less likely to be interpreted literally in all 3 context conditions ( $\beta = 1.672$  in no context condition, 1.604 in non-supportive context condition, and 1.037 in supportive context condition,  $p\_MCMCs < 0.001$ ). In addition, Passive sentences are less likely to be interpreted literally in all 3 context conditions ( $\beta = -0.656$  in no context condition, -0.625 in non-supportive context condition, and -0.905 in supportive context condition,  $p\_MCMCs < 0.02$ ). Notice that the contrast between Active and Passive sentences under the no context condition also reaches significance in this experiment.

Finally, across constructions, participants are less likely to literally interpret implausible DO/PO sentences, than implausible Active/Passive sentences ( $\beta s < -2.758$ ,  $p\_MCMCs < 0.001$ ).

Results in both Experiments 3 and 4 do not support the speculation that the lower literal interpretation rate under the Supportive Context condition than the Non-Supportive Context condition was due to participants placing more scrutiny on the target sentence. If this were to be the case, we would also expect a higher literal interpretation rate under the Non-Supportive Context condition, than under the No Context condition, where nothing incentivizes readers to pay more scrutiny. However, the results show that the literal interpretation rates under the Non-Supportive Context condition and the No Context condition are mostly similar, suggesting that

the differences between the results in the Supportive and the Non-Supportive conditions were indeed driven by the prior  $p(s_i)$ .

## 6. General Discussion

Numerous past studies have investigated how comprehenders draw inferences when they encounter sentences that are implausible to them (e.g. Gibson et al., 2013, Poppels & Levy, 2016, Ryskin et al., 2018; Liu et al., 2020a; Keshev & Meltzer-Asscher, 2021; Washington et al., 2023). These studies suggest that comprehenders adopt a rational approach: they consider both the probability of the intended sentence  $p(s_i)$  and the likelihood for the intended sentence to be corrupted into the perceived sentence due to noise  $p(s_i \rightarrow s_p)$ . However, the test stimuli in these experiments were all isolated sentences. In other words, at each trial, participants only read the test sentence on its own and then were asked to answer a comprehension question. This experimental setting is limited, in that sentences are rarely present on their own in everyday communications: there is always a context. Past studies (e.g. Miller & Selfridge, 1950, Erickson & Mattson, 1981, Nieuwland and van Berkum, 2005, 2006) have shown that discourse context heavily affects how listeners interpret sentences in real time. To see how varying context conditions affects how participants rationally interpret the sentences, in this study, participants read syntactically licit but semantically implausible sentences preceded by a supportive discourse context, non-supportive discourse context, or no discourse context. As in previous studies, readers' interpretations of these implausible sentences were determined by their response to comprehension questions. The addition of two experimental conditions involving discourse context aims to narrow the gap between the standard experimental setting and everyday communication, where sentences are rarely read "out of the blue."

A prediction from the noisy-channel framework is that a supportive context raises the probability of the plausible alternative sentence  $p(s_i)$ . Under the same construction, where the likelihood of noise corruption  $p(s_i \rightarrow s_p)$  stays the same, the inference rate  $p(s_i|s_p)$  increases, according to Bayes Theorem. Behaviorally, this will result in a lower literal interpretation rate in sentences preceded by a supportive context, compared with those preceded by a non-supportive context, or no context. This is what we observed both numerically (Experiments 1-4) and statistically (Experiments 2-4) in this study for sentences under DO/PO construction. On the other hand, in the Active/Passive constructions, the effects of context were hardly visible, consistent with prior findings of ceiling rates of literal interpretations in the active/passive constructions (Gibson et al., 2013). Plausibly, the very low likelihood of an exchange across a main verb (Poppels & Levy, 2016) eclipses the effect of context on the prior (Equation 1). Taken together, our results are consistent with previous work showing that world knowledge and discourse context affect sentence interpretation and can render an implausible sentence plausible (e.g., Hagoort et al., 2004; Nieuwland & van Berkum, 2005; Nieuwland & van Berkum, 2006).

Another prediction from the noisy-channel framework is that sentences made implausible by deletion are more likely to be inferred as their more plausible meaning than those made implausible by insertion. In our materials, this prediction is tested via comparing the literal interpretation rate between DO sentences and PO sentences. We predict that regardless of context conditions, DO sentences will have a lower literal interpretation rate than PO sentences, and this is exactly what we find in Experiments 3 (Comparison 3 in Table 3) and 4 (Comparison 3 in Table 4), as well as in previous studies (e.g. Gibson et al., 2013; Poppels & Levy, 2016).

A third prediction from the framework is that sentences made implausible by one edit are more likely to be inferred as their more plausible meaning than those made implausible by two

edits. This is tested in our study by comparing the literal interpretation rate between sentences under DO/PO construction and those under Active/Passive construction. Since Active/Passive sentences are made implausible by the insertion or deletion of both a copula and a preposition ‘by’ (Gibson et al., 2013), it should be much less likely for them to be corrupted by noise, compared with DO/PO sentences, which are usually made implausible by the insertion or deletion of a preposition ‘to’. This should be reflected in the experiments that DO/PO sentences are less likely to be interpreted literally, in all context conditions, and this is, again, what we observed in Experiments 3 and 4 (See Comparison 5 in Tables 3 and 4), and in past studies.

Both the Active sentences and the Passive sentences can also be made implausible by an exchange of noun phrases across a main verb (Poppels & Levy, 2016), leading to a similar value in the likelihood term  $p(s_i \rightarrow s_p)$ . Interestingly, we observe a numerically (and statistically in Experiments 3 and 4, cf. Comparison 4 in Tables 3 and 4) lower literal interpretation rate for Passive sentences compared to Active sentences, in all context conditions. This is possibly because passive voice is relatively less frequently used than active voice, which results in a lower prior probability  $p(s_i)$  for Passive sentences. This potential effect of structural prior, although not always resulting in statistical significance, has also been pointed out in other related works (Liu et al., 2020a; Poliak et al., in press).

The results indicate that even under a discourse context supporting the plausible alternative interpretation, participants still tend to interpret implausible active/passive sentences literally, possibly because an exchange of NPs across a main verb is perceived to be so unlikely that a supportive context suggesting otherwise cannot exert enough influence on the participants’ interpretations. This indicates that the effects of context is somewhat limited: context can only

influence interpretation of sentences made implausible by noise operations that are likely to happen.

Both the current study and Experiment 3 in Gibson et al. (2013) found that interpretation of implausible sentences could be affected by varying the probability of the plausible alternative interpretation. However, the mechanism in which such a probability is lowered differs in these two studies. In Experiment 3 of Gibson et al. (2013), all the test sentences were not preceded by any discourse context, and the experimenters increased the percentage of implausible target sentences among all sentences read by participants, compared with their first two experiments. In other words, participants in Experiment 3 read more implausible sentences than those from previous experiments. In a later study (Washington et al., 2023), the experimenters increased the probability of implausible sentences by replacing all filler sentences with implausible ones. In both studies, as predicted, participants were more likely to interpret implausible sentences literally. In both Gibson et al. (2013) and Washington et al. (2023), the term  $p(s_i)$  was manipulated directly by varying the proportion of implausible sentences, and as a result, participants possibly considered implausible sentences as more likely to take place. In contrast, in this study, the proportion of implausible sentences is constant in all four experiments, in that under different context conditions, we are varying the proportion of different types of contexts that have disparate probabilities of leading to a specific plausible meaning. For example, in the Supportive context condition, we raise the probability of the context that is very likely to give rise to the intended interpretation, thus increasing  $p(s_i)$ . Both mechanisms are predicted by the noisy-channel framework and seem to be supported by the experimental results in all three studies.

One possible alternative explanation for the lower literal interpretation rate under the supportive context condition is that participants are more likely to “misread” implausible sentences as their plausible alternatives when they are preceded by a supportive context, as previous studies on shallow processing (e.g. Erickson & Mattson, 1981, Barton & Sanford, 1993, Kuperberg, 2007) suggest. In other words, it is possible that when participants saw sentences such as “*The mom gave the candle the daughter*”, they might have read it as “*The mom gave the candle **to** the daughter*”. If this is the case, the participants would adopt the non-literal interpretation simply because they did not notice the implausibility caused by noise in the sentence stimulus (e.g. Huang & Staub, 2021a, 2021b), which did not involve any noisy-channel inference. However, previous studies (Liu et al., 2020b, Bader & Meng, 2018; Meng & Bader, 2021, Cutter et al., 2022a) show that, at least in implausible Active/Passive sentences, participants are fully aware of their implausibility, and that their responses are indeed based on their rational inferences. Similarly, some studies find participants are able to detect the implausibility in implausible DO/PO sentences (e.g. Slevc & Buxó-Lugo, 2020; Cai et al., 2022; Paape, 2023), although others do not or are inconclusive (e.g. Cutter et al., 2022b).

Real-world occurrences of sentences, regardless of the environment being noisy or not, are usually embedded in a discourse context. An important implication of this work is that, because previous reports of noisy-channel sentence comprehension have focused on the understanding of isolated sentences (Gibson et al., 2013; 2017), they likely underestimated the likelihood of noisy-channel inferences in everyday language comprehension. Meanwhile, the word “context” is merely an umbrella term covering a wide range of information sources, acting on diverse timescales (Christiansen & Chater, 2016; Ryskin & Fang, 2021). The discourse context tested in this work is only one of these potential information sources. An interesting

future direction would be to look into the effect of context at a timescale other than discourse, and how different timescales interact with each other (Hess et al., 1995, Nieuwland and van Berkum, 2006). Moreover, future work may take advantage of more implicit measures (e.g., ERPs) to examine noisy-channel inference in more naturalistic settings (e.g., while reading a story). In addition, the present study is mainly focused on the effect of a discourse context on the semantic prior, while on the other hand, it might also have an effect on a comprehender's noise model. Specifically, both the supportive context and non-supportive context in this study contain syntactically licit, semantically plausible sentences, which could possibly lower a comprehender's perceived noise rate and therefore raise the literal interpretation rate of implausible sentences. We leave the question of how a discourse context would affect the noise model to future research.

## **7. Conclusion**

The present findings provide further support for a noisy-channel approach to language processing and highlight the importance of considering how this rational inference process is affected by prior context. Furthermore, these results suggest that noisy-channel processing is even more prevalent outside of the research setting, where there is often greater contextual support for the intended meaning of a sentence. Indeed, everyday language typically involves comprehending a sentence using relevant information from various sources, including the context derived from preceding sentences. Thus, our findings shed light on the possibility that in daily life, noisy-channel inferences are more commonplace than previously suggested by investigations of isolated sentence comprehension.

**8. Conflict of interest statement**

The authors declare no conflict of interest.



### References

- Albrecht, J. E., & O'Brien, E. J. (1993). Updating a mental model: Maintaining both local and global coherence. *Journal of Experimental Psychology: Learning, memory, and cognition*, 19(5), 1061.
- Bader, M., Meng, M., Bader, M., & Meng, M. (2023). Processing noncanonical sentences: effects of context on online processing and (mis) interpretation. *Glossa Psycholinguistics*, 2(1).
- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. J. (2013). Random effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of memory and language*, 68(3), 255-278.
- Bader, M., & Meng, M. (2018). The misinterpretation of noncanonical sentences revisited. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44(8), 1286.
- Barton, S. B., & Sanford, A. J. (1993). A case study of anomaly detection: Shallow semantic processing and cohesion establishment. *Memory & cognition*, 21(4), 477-487.
- Bates, D., Maechler, M., Bolker, B., & Walker, S. (2015). Fitting Linear Mixed-Effects Models Using lme4. *Journal of Statistical Software*, 67(1), 1-48.<doi:10.18637/jss.v067.i01>.
- Bergen, L., Levy, R., & Gibson, E. (2012). Verb omission errors: Evidence of rational processing of noisy language inputs. *Proceedings of the 34th Annual Conference of the Cognitive Science Society*
- Cai, Z. G., Zhao, N., & Pickering, M. J. (2022). How do people interpret implausible sentences?. *Cognition*, 225, 105101.
- Camblin, C. C., Gordon, P. C., & Swaab, T. Y. (2007). The interplay of discourse congruence and lexical association during sentence processing: Evidence from ERPs and eye tracking. *Journal of Memory and Language*, 56(1), 103-128.

- Christianson, K., Hollingworth, A., Halliwell, J. F., & Ferreira, F. (2001). Thematic Roles Assigned along the Garden Path Linger. *Cognitive Psychology*, 42(4), 368–407.  
<https://doi.org/10.1006/cogp.2001.0752>
- Fano, R. M. (1961). Introductory remarks in *Structure of language and its mathematical aspects* (pp. 261-263). Providence, RI: AMS.
- Christiansen, M. H., & Chater, N. (2016). The now-or-never bottleneck: A fundamental constraint on language. *Behavioral and brain sciences*, 39.
- Cutter, M. G., Paterson, K. B., & Filik, R. (2022a). Online representations of non-canonical sentences are more than good-enough. *Quarterly Journal of Experimental Psychology*, 75(1), 30-42.
- Cutter, M. G., Paterson, K. B., & Filik, R. (2022b). The effect of age and eye-movements on noisy-channel inference. Presented at the 35rd Annual Human Sentence Processing Conference, the University of California - Santa Barbara. Santa Barbara, CA.
- Erickson, T. D., & Mattson, M. E. (1981). From words to meaning: A semantic illusion. *Journal of Verbal Learning and Verbal Behavior*, 20(5), 540-551.
- Fano, R. M. (1961). Introductory remarks in *Structure of language and its mathematical aspects* (pp. 261-263). Providence, RI: AMS.
- Ferreira, F., & Patson, N. D. (2007). The ‘good enough’ approach to language comprehension. *Language and Linguistics Compass*, 1(1-2), 71-83.
- Frazier, L., & Fodor, J. D. (1978). The sausage machine: A new two-stage parsing model. *Cognition*, 6(4), 291–325.
- Gibson, E. (1998). Linguistic complexity: Locality of syntactic dependencies. *Cognition*, 68(1), 1-76.

Gibson, E., Bergen, L., & Piantadosi, S. T. (2013). Rational integration of noisy evidence and prior semantic expectations in sentence completion. *Proceedings of the National Academy of Sciences of the United States of America*, 110(20), 8051-8056.

Gibson, E., Tan, C., Futrell, R., Mahowald, K., Konieczny, L., Hemforth, B. & Fedorenko, E. (2017). Don't underestimate the benefits of being misunderstood. *Psychological Science*.  
<https://doi.org/10.1177/0956797617690277>

Hadfield, J. D. (2010). MCMC methods for multi-response generalized linear mixed models: the MCMCglmm R package. *Journal of statistical software*, 33, 1-22.

Hagoort, P., Hald, L., Bastiaansen, M., & Petersson, K. M. (2004). Integration of word meaning and world knowledge in language comprehension. *Science*, 304(5669), 438-441.

Hess, D. J., Foss, D. J., & Carroll, P. (1995). Effects of global and local context on lexical processing during language comprehension. *Journal of Experimental Psychology: General*, 124(1), 62.

Huang, K. J., & Staub, A. (2021a). Why do readers fail to notice word transpositions, omissions, and repetitions? A review of recent evidence and theory. *Language and Linguistics Compass*, 15(7), e12434.

Huang, K. J., & Staub, A. (2021b). Using eye tracking to investigate failure to notice word transpositions in reading. *Cognition*, 216, 104846.

Keshev, M., & Meltzer-Asscher, A. (2021). Noisy is better than rare: Comprehenders compromise subject-verb agreement to form more probable linguistic structures. *Cognitive Psychology*, 124, 101359.

Kuperberg, G. R. (2007). Neural mechanisms of language comprehension: Challenges to syntax. *Brain research*, 1146, 23-49.

- Kutas, M., & Hillyard, S. A. (1980). Reading senseless sentences: Brain potentials reflect semantic incongruity. *Science*, 207(4427), 203-205.
- Levy, R. (2008a). Expectation-based syntactic comprehension. *Cognition*, 106(3), 1126-1177.
- Levy, R. (2008b). A noisy-channel model of rational human sentence comprehension under uncertain input. *Proceedings of the 13<sup>th</sup> Conference on Empirical Methods in Natural Language Processing* (Association for Computational Linguistics, Stroudsburg, PA), 234-243.
- Levy, R., Bicknell, K., Slattery, T., Rayner, K. (2009). Eye movement evidence that readers maintain and act on uncertainty about past linguistic input. *Proc. Natl. Acad. Sci. USA* 106(50): 21086-21090.
- Levy, R. (2011). Integrating surprisal and uncertain-input models in online sentence comprehension: Formal techniques and empirical results. *Proceedings of the 49<sup>th</sup> Annual Meeting of the Association for Computational Linguistics* (Association for Computational Linguistics, Portland, OR), 1055-1065.
- Liu, Y., Ryskin, R., Futrell, R., Gibson, E. (2020a). Structural frequency effects in comprehenders' noisy-channel inferences. Presented at the 33rd Annual CUNY Human Sentence Processing Conference, the University of Massachusetts, Amherst. <https://osf.io/h9tdf/> Liu, Y., Ryskin, R., Futrell, R. & Gibson, E. (2020b). Structural Frequency Effects in Noisy-channel Comprehension. *Presentation at the Penn Linguistics Conference*.
- Meng, M., & Bader, M. (2021). Does comprehension (sometimes) go wrong for noncanonical sentences?. *Quarterly Journal of Experimental Psychology*, 74(1), 1-28.
- Miller, G. A., & Selfridge, J. A. (1950). Verbal context and the recall of meaningful material. *The American journal of psychology*, 63(2), 176-185.

Nieuwland, M. S., & van Berkum, J. J. (2005). Testing the limits of the semantic illusion phenomenon: ERPs reveal temporary semantic change deafness in discourse comprehension.

*Cognitive Brain Research*, 24 (3), 691-701.

Nieuwland, M. S. & van Berkum, J. J. A. (2006). When Peanuts Fall in Love: N400 Evidence for the Power of Discourse. *Journal of Cognitive Neuroscience*, 18:7, 1098-1111.

Paape, D. (2023). Good-enough processing versus rational inference in linguistic illusions: Modeling judgments of formal correctness and meaning recoverability with a race model.

*Presented at the 36th Annual Human Sentence Processing Conference.*

Poliak, M., Ryskin, R., Braginsky, M., & Gibson, E. (in press). It's not What You Say but How You Say It: Evidence from Russian Shows Robust Effects of the Structural Prior on Noisy Channel Inferences.

Poppels, T. & Levy, R. P. (2016). Structure-sensitive Noise Inference: Comprehenders Expect Exchange Errors. In *Proceedings of the 38th annual meeting of the Cognitive Science Society* (pp. 378–383).

Ryskin, R., & Fang, X. (2021). The many timescales of context in language processing. In *Psychology of Learning and Motivation* (Vol. 75, pp. 201-243). Academic Press.

Ryskin, R., Futrell, R., Kiran, S. & Gibson, E. (2018) Comprehenders model the nature of noise in the environment. *Cognition*, 181, 141-150.

Ryskin, R., Stearns, L., Bergen, L., Eddy, M., Fedorenko, E. & Gibson, E. (2021). An ERP index of real-time error correction within a noisy-channel framework of human communication.

*Neuropsychologia*.

Shannon, C. (1948). A Mathematical Theory of Communication. *Bell System Technical Journal*, 27(4), 623-656.

Slevc, L.R., & Buxó-Lugo, A. (2020). Noisy evidence and plausibility influence structural priming. *Poster presented at 33rd Annual Human Sentence Processing Conference*, Amherst, Massachusetts. <https://osf.io/huw86>.

Staub, A., Dodge, S., & Cohen, A. L. (2018). Failure to detect function word repetitions and omissions in reading: Are eye movements to blame? *Psychonomic Bulletin & Review*.  
<https://doi.org/10.3758/s13423-018-1492-z>

Washington, L., Chen, S., & Gibson, E. (2023). The Influence of Prior Semantic Knowledge Noisy Channel Interpretation. *Poster presented at 36th Annual Human Sentence Processing Conference*, Pittsburgh, Pennsylvania.