

How abstract are logical representations? The role of verb semantics in representing quantifier scope

Mieke Sarah Slim, Language Development Department, Max Planck Institute for Psycholinguistics, Nijmegen, The Netherlands; Department of Experimental Psychology, Ghent University, Ghent, Belgium, mieke.slim@mpi.nl

Peter Lauwers, Department of Linguistics, Ghent University, Ghent, Belgium, peter.lauwers@ugent.be

Robert J. Hartsuiker, Department of Experimental Psychology, Ghent University, Ghent, Belgium, Robert.Hartsuiker@ugent.be

Language comprehension involves the derivation of the meaning of sentences by combining the meanings of their parts. In some cases, this can lead to ambiguity. A sentence like *Every hiker climbed a hill* allows two logical representations: One that specifies that every hiker climbed a different hill and one that specifies that every hiker climbed the same hill. The interpretations of such sentences can be primed: Exposure to a particular reading increases the likelihood that the same reading will be assigned to a subsequent similar sentence. Feiman and Snedeker (2016) observed that such priming is not modulated by overlap of the verb between prime and target. This indicates that mental logical representations specify the compositional structure of the sentence meaning without conceptual meaning content. We conducted a close replication of Feiman and Snedeker's experiment in Dutch and found no verb-independent priming. Moreover, a comparison with a previous, within-verb priming experiment showed an interaction, suggesting stronger verb-specific than abstract priming. A power analysis revealed that both Feiman and Snedeker's experiment and our Experiment 1 were underpowered. Therefore, we replicated our Experiment 1, using the sample size guidelines provided by our power analysis. This experiment again showed that priming was stronger if a prime-target pair contained the same verb. Together, our experiments show that logical representation priming is enhanced if the prime and target sentence contain the same verb. This suggests that logical representations specify compositional structure and meaning features in an integrated manner.



1. Introduction

The computation of sentence meaning involves combining the meanings of single lexical items into a complex sentence meaning. This process is assumed to involve two cognitive systems: the *lexical semantic* system, which stores the conceptual meanings of content words, and the *compositional semantic* system, which specifies the abstract combinatorial rules needed to piece together the meaning of the content words (e.g., Feiman & Snedeker, 2016; Heim & Kratzer, 1998; see also Piantadosi et al., 2008).¹ This latter system is involved in the interpretation of function words, like quantifiers (e.g., *every*, *a*). Such quantifiers do not refer to anything out there in the world but instantiate abstract semantic operations that are required to compute the compositional structure of the sentence meaning. Consider the sentence in (1a).

- (1) a. Every dog is breakdancing.

In this sentence, the quantifier *every* specifies that the predicate ‘is breakdancing’ applies to all things that are ‘dogs’ (anything that is a dog is also a breakdancing thing). This compositional structure of the sentence meaning is typically represented as a *logical representation* (May, 1985; Raffray & Pickering, 2010; see also Ruys & Winter, 2011, for overview).² The logical representation corresponding to (1a) is given in (1b).

- (1) b. $\forall x[\text{DOG}(x) \rightarrow \text{IS-BREAKDANCING}(x)]$
For all x , if x is a dog, then x is a breakdancing thing.

In language use, we compute such logical representations to derive complex meanings from sentences. But what information exactly is specified in logical representations? Do they only specify the compositional structure of sentence meaning (as computed by the compositional semantic system) or do they also store the (aspects of) meaning content of the content words (as specified in the lexical semantic system)? The present study centers on this latter question, by conducting two replication studies of a previous experiment by Feiman and Snedeker (2016, Experiment 3), which tested whether the conceptual meaning content of verbs is specified in mental logical representations.

¹ These two systems are sometimes referred to as the combinatorial system and the conceptual system (Feiman & Snedeker, 2016). However, these terms are somewhat misleading, because it is typically assumed that the “combinatorial” system is responsible for combining concepts in thought (making it part of our wider “conceptual system”; e.g., Fodor, 1982; Frankland & Greene, 2020; Jackendoff, 2002)

² We use the term *logical representations* in a theory-neutral way to refer to meaning representations at a sentence level, without any commitments to *logical form*-based accounts of sentence interpretation (e.g., Chomsky, 1995; Heim & Kratzer, 1998; May, 1985)

1.1 Are logical representation abstract?

The nature of mental logical representations is especially well-studied in doubly-quantified sentences, which are notorious for their ambiguity in the compositional semantic structure (see Brasoveanu & Dotlacil, 2019; Ruys & Winter, 2011, for overview). Consider (2a), containing the universal quantifier *every* and the existential quantifier *a*.

- (2) a. Every man bit a dog.

Does this sentence mean that every man bit a different dog, or does it mean that every man bit the same dog? This ambiguity emerges because the compositional rules specified by *every* and *a* can interact in two ways (e.g., Ruys & Winter, 2011; Szabolcsi, 1997). This, in turn, results in two possible logical structures, presented in (2b) and (2c).

- (2) b. *Universal-wide*: $\forall x[\text{MAN}(x) \rightarrow \exists y[\text{DOG}(y) \wedge \text{BIT}(x,y)]]$
 For all x , if x is a man, there exists a y such that y is a dog, and x bit y
- c. *Existential-wide*: $\exists y[\text{DOG}(y) \wedge \forall x[\text{MAN}(x) \rightarrow \text{BIT}(x,y)]]$
 There exists a y , such that y is a dog, and for all x , if x is a man, then x bit y

The logical representation in (2b) corresponds to the interpretation in which every man bit a (potentially) different dog, and the one in (2c) corresponds to the interpretation in which the same dog was bitten by every man. These formulae highlight that the ambiguity can be captured by the relative ordering in which *every* and *a* combine in the compositional structure of the sentence's interpretation. In (2b) *every* takes wide scope over the other elements in the sentence (this reading will therefore be called the *universal-wide interpretation*), whereas in (2c) *a* takes wide scope over the other elements in the sentence (henceforth the *existential-wide interpretation*).

The central question in this paper is whether the representation of a sentence's compositional meaning structure is independent from the representations of the conceptual meaning of the content words in a sentence. Theorists do not agree on the answer to this question. Some theorists argue that the compositional structure of sentence meaning is processed independently from the processing of the concepts denoted by the content words (e.g., Feiman & Snedeker, 2016; Heim & Kratzer, 1998). Under this view, the logical representations in (2b) and (2c) specify the verb arguments as abstract variables (x and y), and the verb as a relation between these variables (i.e., $R(x,y)$). Importantly, the logical representations specify how these variables are pieced together, but not the meanings of these variables themselves (that is, the conceptual meaning of *DOG*, *MAN* and *BIT*). In other words, the logical representations specify how things that are 'men', a 'dog', and involved in 'biting' must be combined, but they do not specify what it means to be a 'man', a 'dog', or 'biting'. The compositional rules, which specify how concepts are combined,

are stored in the compositional system, and the content of these concepts themselves, in turn, is stored separately in the lexical semantic system (Feiman & Snedeker, 2016; Heim & Kratzer, 1998; Piantadosi et al., 2008). So, in this view, logical representations are abstract, in the sense that they specify content-free structure. Here, the term *abstract* is borrowed from literature on syntactic representations, where *abstract* representations are understood as representations that only specify bare syntactic phrase structure (and no additional lexical or semantic features related to the structure of the sentence; e.g., Bock & Loebell, 1990; Ziegler et al., 2019).

This abstract view of logical representations is contested by theorists who stipulate that logical representations integrate both compositional structure and meaning content of the concept words. Within these theories, more emphasis is put on the role of pragmatic (real-world) knowledge in the computation of sentence meaning (e.g., Altmann & Steedman, 1988; Saba & Corriveau, 2001). This is based on the observation that the conceptual meaning of content words greatly affects the computation of compositional structure. Consider the sentences in (3) and (4), taken from Dwivedi (2013).

(3) Every customer returned an item.

(4) Every jeweller praised a diamond.

The abstract compositional structure of (3) and (4) is identical. Nevertheless, Dwivedi (2013) observed that the sentence in (3) is interpreted much more often as universal-wide than the sentence in (4). This difference is due to pragmatic biases evoked by the events denoted in the sentence. More specifically, the “returning an item” and “praising a diamond” events differ in their prototypical number of individuals and objects involved: Whereas it is common that different jewellers praise the same diamond, customers typically return different items (see also Saba & Corriveau, 2001). These examples show that the lexical semantics of a sentence plays a pivotal role in the processing of the compositional semantics, suggesting that conceptual meaning content and compositional structure are processed in an integrated manner, resulting in logical representations that specify both structure and content.

This integrated view of logical representations fits with theoretical descriptions of logical representations that assume that logical representations are conceptual representations that capture the event situation that is described in the sentence (e.g., Fodor, 1982; Johnson-Laird et al., 1989). Fodor (1982), for example, postulated that logical representations can be thought of as situation models that capture the number of participants, objects, and actions that are part of the situation described in the sentence. This closely resembles accounts of event representations that posit that event situations are represented as ‘mental models’ (e.g., Goodwin & Johnson-Laird, 2011; Johnson-Laird, 1983; Johnson-Laird et al., 1989; Khemlani et al., 2012; but cf. Altmann & Ekves, 2019, for an alternative account on event structure). We, however, remain

agnostic to the exact conceptual format of event representations. More relevant to our research purposes is that a theoretical description of logical representations as event representations does not propose distinct levels for encoding structure and meaning content.

1.2 Priming logical representations

The central question in this paper is whether mental logical representations are abstract representations of compositional structure alone, or integrated event representations that specify both structure and content. An efficient tool to examine the nature of mental logical representations is *structural priming*, which refers to the tendency to re-use the structure of a prior sentence (e.g., Bock, 1986; Bock & Loebell, 1990; Branigan et al., 2005; Ziegler & Snedeker, 2018; see Branigan & Pickering, 2017; Pickering & Ferreira, 2008, for reviews; and Mahowald et al., 2016, for meta-analysis). Raffray and Pickering (2010) observed that logical representations are also susceptible to priming using a sentence-picture matching task. On each trial, the participants matched a sentence with one out of two pictures. The prime trials contained a doubly-quantified sentence (e.g., *Every kid climbed a tree*), a matching picture, and a mismatching foil picture. The matching picture only matched one of the two readings of the sentence, forcing the participants to assign that particular reading to the sentence. Each prime was directly followed by a target trial. The target trials again contained a doubly-quantified sentence (e.g., *Every hiker climbed a hill*). Unlike the primes, however, the two response pictures in the target trials display both possible readings of the sentence. Therefore, the participants now have a real choice between the two possible readings of the sentence. The results showed an effect of priming: Participants were more likely to assign the universal-wide reading to the target sentence if the previous prime trial forced a universal-wide reading than if the previous prime trial forced an existential-wide reading (and vice versa for the existential-wide reading).

Effects of structural priming indicate that exposure to a certain structure facilitates the future processing of that structure. This effect is assumed to emerge because it is easier to re-use a representation if it was used in the processing of a previous, related, sentence (e.g., Chang et al., 2006; Pickering & Branigan, 1998). This makes structural priming a suitable window into (linguistic) representation: Priming effects indicate that two sentences share underlying representations, with larger effects denoting more representational overlap between the two sentences (e.g., Bock, 1986; Branigan et al., 2005; Pickering & Branigan, 1998; for review, see Branigan & Pickering, 2017). Crucially, structural priming can be used to distinguish between abstract structure representations and integrated structure/content representations (see also Ziegler et al., 2019; Ziegler & Snedeker, 2018). Specifically, if priming emerges on the basis of an abstract logical representation, the priming effect should not be modulated by overlap of the content words in the prime and target. If, on the other hand, priming emerges on the basis of integrated structure and content representations, priming effects are enhanced by overlap of the content words.

Important evidence in favour of abstract logical representation priming comes from Feiman and Snedeker (2016, Experiment 3). They adapted Raffray and Pickering's (2010) above-described sentence-picture matching task. In Raffray and Pickering's original task, target sentences contained the same verb as the preceding prime sentence (see also Slim et al., 2021). Therefore, the effects of priming observed in this study can emerge on the basis of abstract structure or on the basis of integrated structure and content representations. To tease apart these two interpretations, Feiman and Snedeker altered Raffray and Pickering's task in such a way that the prime and target sentences contained different verbs (e.g., *Every beggar hit a reporter* to *Every hiker climbed a hill*). This task showed a priming effect, and crucially, the size of this effect was statistically comparable to priming within the same verb (although there was a small numerical difference: Priming within the same verb showed a descriptive effect of 6.47%, whereas priming between verbs showed a descriptive effect of 4.98%).

Feiman and Snedeker's (2016) findings thus indicate that logical representations are abstract representations of compositional structure. This suggests that sentence comprehension involves a level of processing at which abstract content-free logical representations are computed, indicating a division of labour between the compositional and the lexical system (see also Heim & Kratzer, 1998). In the present paper, we re-evaluate Feiman and Snedeker's findings in a replication study. However, it is important to first note that the role of lexical repetition in structural priming has been widely studied, especially outside research on logical representations.

1.3 The role of lexical repetition in priming

Feiman and Snedeker (2016) argued that logical representations only specify abstract compositional structure based on the finding that repetition of the verb does not modulate logical representation priming. The role of lexical repetition, especially of the verb, has been the subject of much investigation in research on structural priming. Most of these studies focused on the nature of *syntactic* representations (e.g., Bock, 1986; see Mahowald et al., 2016, for review and meta-analysis). In these studies, priming is characterised by the re-use of a particular syntactic structure (e.g., a double-object dative like *The man gave the dog a bone*) after exposure to an earlier sentence in that structure (e.g., *The girl gave the boy a balloon*) compared to exposure to an earlier sentence in an alternative structure (e.g., a prepositional-object dative like *The girl gave a balloon to the boy*). This effect was originally observed in production studies (Bock, 1986), but has also been found in comprehension studies (e.g., Arai et al., 2007; Branigan et al., 2005; Thothathiri & Snedeker, 2008; Ziegler & Snedeker, 2019; for review, see Tooley, 2022).

Many studies on syntactic priming have shown that the effect of priming is enhanced if the prime and target sentences contain the same verb (Pickering & Branigan, 1998; Schoonbaert et al., 2009). This effect, known as the *lexical boost* effect, is robust in studies that examined syntactic

priming in production (Mahowald et al., 2016). In comprehension, however, the evidence for a lexical boost effect in syntactic priming is piecemeal (Tooley, 2022). Some studies show that priming only emerges if the verb is repeated in prime and target (e.g., Arai et al., 2007), whereas other studies show priming regardless of whether the verb is repeated (with no modulating effect of verb repetition, e.g., Thothathiri & Snedeker, 2008; Ziegler & Snedeker, 2019).

Feiman and Snedeker's (2016) results suggest that there is no lexical boost effect in logical representation priming in comprehension. However, care is required in drawing this conclusion, since the lexical boost effect is less robust in comprehension than in production. A relevant difference between studies on production and comprehension is that production studies typically measure which structure is produced (i.e., the *outcome* is measured, e.g., Bock, 1986). In comprehension, on the other hand, the *online* processing of the target sentences is often measured (e.g., by studying looking or reading times, e.g., Arai et al., 2007; Ziegler & Snedeker, 2019; see Tooley, 2022, for review), which may explain – at least partially – why the lexical boost effect is not consistently observed in comprehension whereas it is in production.

The sentence-picture matching task used in Feiman and Snedeker (2016) tested outcome measures and not online measures. Moreover, the general sentence-picture matching paradigm that Feiman and Snedeker used is also used to test syntactic priming in comprehension. Branigan et al. (2005) used a similar sentence-picture matching task to test effects of syntactic priming in comprehension. In their task, the prime and target structures were syntactically ambiguous descriptions like *The waitress prodding the clown with the umbrella* (did the waitress use the umbrella to prod the clown, or did the waitress prod a clown who is holding an umbrella? – this ambiguity emerges because the sentence allows multiple syntactic analyses). Branigan et al.'s results showed priming, but only if the prime and target contained the same verb.

Branigan et al.'s (2005) findings suggest that it is possible to find a lexical boost of syntactic priming in comprehension with the sentence-picture matching task, which contrasts with Feiman and Snedeker's (2016) findings on logical representation priming. The present study further assessed abstract priming of logical representations in such a paradigm and compared it to priming with verb overlap. We note though that any effect reminiscent of a 'lexical boost' in logical representation priming is presumably driven by different mechanisms from the lexical boost in syntactic priming; we return to that point in Section 6 (the General Discussion).

1.4 The present study

Below, we first report a close replication experiment of Feiman and Snedeker's (2016) Experiment 3 in Dutch. Our initial motivation for this experiment was not to re-evaluate the strength of logical representation priming between verbs. However, to foreshadow the results, Experiment 1 did not find logical representation priming in the absence of verb overlap. This finding required

us to replace our initial goals with the goal of testing whether Feiman and Snedeker's finding of abstract logical representation priming is generalisable and whether such priming is enhanced by (or even conditional on) verb overlap. To gain more insight into the reliability of the between-verb priming effect, we next performed a power analysis on Feiman and Snedeker's data (a procedure inspired by Harrington Stack et al., 2018). This power analysis revealed that both Feiman and Snedeker's original experiment and our replication were underpowered. A new close replication of Feiman and Snedeker's experiment, which followed the sample size guidelines provided by the power analyses (reported below as Experiment 2), again showed that logical representation priming is stronger with verb overlap than priming between different verbs. Combined analyses of Experiment 1 and Experiment 2 further corroborated that logical representation priming is enhanced if the prime-target pair contains the same verb. These findings suggest that logical representation priming between verbs is not robust, and that the conceptual meaning of verbs plays a role in processing and representing the compositional structure of a sentence's interpretation.

2. Experiment 1

Experiment 1 was our first replication of Feiman and Snedeker's (2016) Experiment 3. We used the same task as Feiman and Snedeker with similar materials, but conducted our experiment in Dutch instead of English (the language of Feiman and Snedeker's experiment). Nevertheless, Experiment 1 can be considered a close replication following the taxonomy of replications defined by LeBel et al. (2017), because the dependent and independent variables were implemented in the same way.

Like Feiman and Snedeker (2016), we gained further insight into the role of verb repetition in logical representation priming by comparing the data of Experiment 1 to the data of a similar experiment that tested logical representation priming within verbs (in Dutch) in the analyses. This latter experiment is not reported here but in Slim et al. (2023, Experiment 1).

2.1 Method

2.1.1 Participants

We recruited 69 Dutch-speaking participants. They were first-year psychology students at Ghent University and received course credit for their participation. We excluded two participants from the analyses because they guessed the aim of the experiment in a post-experimental questionnaire, and three further participants because they responded incorrectly on more than 10% of the filler trials (see below). Thus, 64 participants were included in the analysis.

2.1.2 Materials

Experiment 1 was a sentence-picture matching task, like the one used in Feiman and Snedeker (2016). The stimuli were adapted from Slim et al. (2021), who in turn partly borrowed the

materials from Raffray and Pickering (2010). The materials from Raffray and Pickering were also used in Feiman and Snedeker's original experiment, making our materials very similar to those of Feiman and Snedeker.

Each trial consisted of a sentence and two response pictures. In the prime and target trials, we presented a doubly-quantified sentence with the universal quantifier *elke* ('every') in the subject position and the existential quantifier *een* ('a') in the direct object position (e.g., *Elke beer naderde een tent* 'Every bear approached a tent'). In the prime trials, one of the response pictures depicted a possible reading of the sentence, whereas the other response picture was a foil picture that mismatched either reading of the sentence (because the picture depicted the wrong agent or the wrong theme). Prime trials were presented in two prime conditions: the *universal-wide* and the *existential-wide* condition. In the universal-wide condition, the matching picture displayed the universal-wide reading of the prime sentence and in the existential-wide condition, the matching picture depicted the existential-wide reading of the prime sentence (Figure 1).

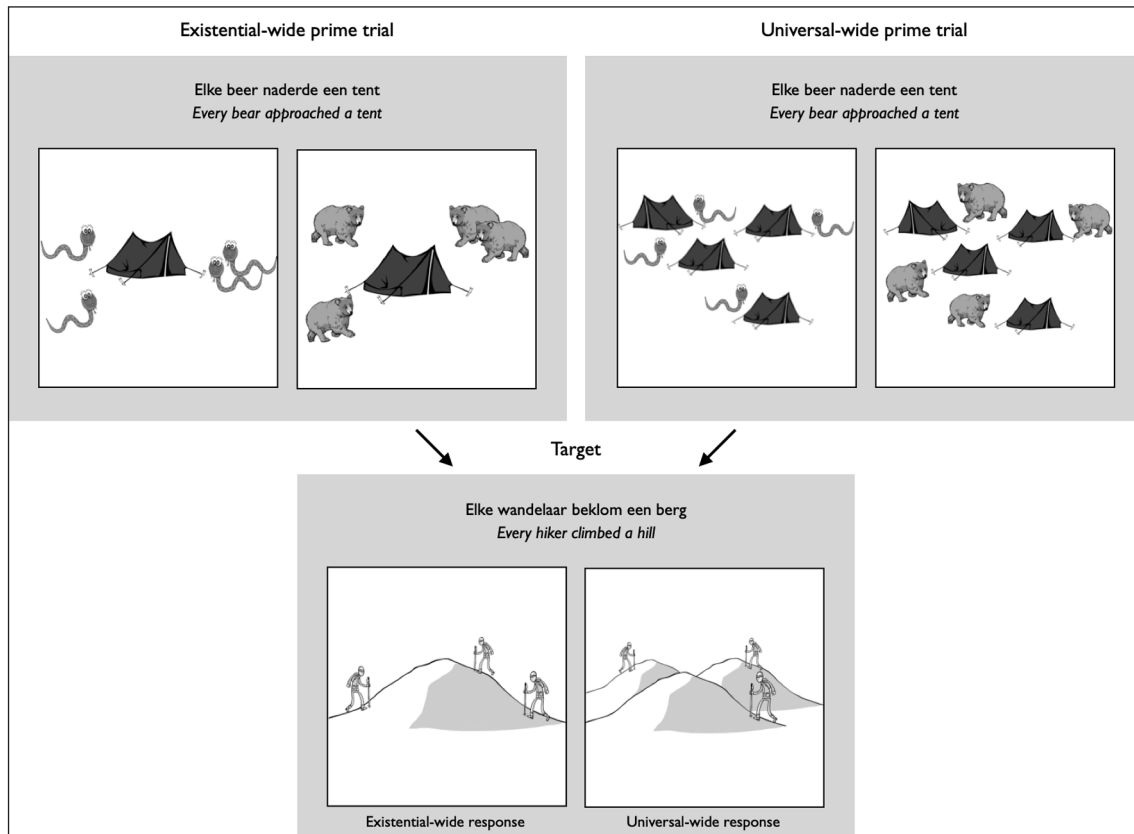


Figure 1: Example of a prime-target pair in Experiment 1. The English translation of the sentences and the labels 'Universal/Existential-wide prime condition/response' are added to this figure for the sake of illustration.

Each prime trial directly preceded a target trial. In the target trials, the participants read another *elke...een* sentence, which always contained a different verb than the preceding prime sentence (see <https://osf.io/697wg/> for a list of the prime-target pairs). The two response pictures in the target trials displayed both possible readings of the target sentence (**Figure 1**). Thus, the participants were forced to assign one of the two possible readings to the prime sentence but had a free choice between either interpretation in the target trials.

Finally, the filler trials contained an unambiguous sentence, a matching picture, and a non-matching picture. Half of the filler trials contained unambiguous transitive sentences (e.g., *De cowboy sloeg de boef* ‘The cowboy punched the burglar’) and the other half involved intransitive sentences with *elke* in subject position (e.g., *Elke heks sliep* ‘Every witch slept’). The aim of adding these latter fillers was to further mask the prime-target pattern in the trials, because a previous study from Slim et al. (2021) showed that a considerable number of participants guessed the aim of the experiment.

The experiment contained 54 prime trials, 54 target trials, and 162 filler trials. Trial order was organised pseudo-randomly: Primes and targets appeared in pairs, and prime-target sets were intervened by two to five filler trials (following Raffray & Pickering, 2010). We created two lists in which the prime condition was counterbalanced across trials. Our materials differed in three aspects from Feiman and Snedeker’s (2016) original experiment: (i) We presented our stimuli in Dutch, (ii) Feiman and Snedeker’s experiment only contained transitive fillers, whereas we constructed additional intransitive fillers, and (iii) our experiment contained 54 prime-target sets, whereas Feiman and Snedeker’s experiment contained 24 prime-target sets. We also added additional fillers to our task, so that the ratio of fillers to primes/targets was the same. Note further that Feiman and Snedeker recruited more participants than us (107 compared to 64) but fewer critical trials; we therefore collected a larger number of observations (3456 observations vs. 2568 in Feiman and Snedeker’s experiment).

2.1.3 Procedure

The experiment was implemented using the *PennController for Ibex (PCIbex)* and carried out online via the *PCIbex Farm* (Zehr & Schwarz, 2018). The experiment started with an informed consent form. In all trials, the sentence and two pictures were shown simultaneously on a screen. The participants were instructed to click on the picture that best fitted the sentence. The instructions also said that participants needed to select their spontaneous preference if they thought that both pictures matched the sentence. After completing the sentence-picture matching task, the participants filled in a short questionnaire about their language background. A final question asked whether the participant had any ideas about the research aims of the experiment. Those who guessed the pattern of the trials (prime-target), and/or guessed that the experiment examined the possible influence of the preceding trial on the target trials were excluded from the analyses (following Slim et al., 2021).

2.2 Analyses and results

2.2.1 Data treatment and analysis procedure

A target response was discarded if the non-matching picture in the preceding prime was selected (following Raffray & Pickering, 2010). Moreover, all responses of a participant were removed if they responded incorrectly on more than 10% of the filler trials. The remaining responses were coded as `true` if the universal-wide response was selected, and `false` if the existential-wide response was selected.

The data were analysed by modelling the response-type likelihood with logit mixed-effect models (Jaeger, 2008). The full model contained the target response as a binomial dependent variable and Prime Condition as a (sum-coded) fixed effect. The random effects structure was maximal (both random slopes and intercepts by Item and Subject; Barr et al., 2013), unless this maximal random-effect structure did not converge. In that case, we omitted random effects until we reached convergence (following the guidelines from Bates et al., 2015). We calculated p -values by running χ^2 -tests on the log-likelihood values of the full model compared to models in which the predictor values were removed one-by-one.

We then compared the data of Experiment 1 to the data from Slim et al. (2023, Experiment 1). These are data from a sentence-picture matching task in which logical representations were primed in Dutch *elke...een* sentences. The materials and procedure of Slim et al.'s Experiment 1 were identical to the current Experiment 1, with the only difference that each prime-target pair in our previous experiment contained the same verb. Moreover, the sample of participants was comparable to those of Experiment 1: 63 (unique) participants that were first-year psychology students at Ghent University who took part for course credit.

Since the only critical difference between Experiment 1 reported here and our previous experiment reported in Slim et al. (2023, Experiment 1) is the repetition of verbs in the prime-target pairs, these combined analyses will reveal whether repetition of the verb modulates priming of logical representations. These analyses were again carried out with logit mixed-effect modelling. The model contained Prime Condition, Verb Overlap, and the interaction term as (sum-coded) fixed effects. In addition, we carried out pairwise comparison to test the effect of Prime Condition in the within-verb and between-verb datasets in isolation. The rest of the analysis procedure was similar to the one described above. All analyses were carried out in R (R Core Team, 2019) using the `lme4` (Bates et al., 2014) and `phia` (De Rosario-Martinez, 2015) packages.

2.2.2 Results and discussion

We discarded 171 target responses because the participants responded incorrectly on the preceding prime (4.95% of the total numbers of prime-target pairs, 88 of these responses were from the

universal-wide and the other 83 were from the existential-wide condition). The remaining responses firstly reveal that the participants mostly selected the universal-wide response. This is not surprising: The Dutch quantifier *elke* is a distributive quantifier. These quantifiers have a strong tendency to take wide scope, which causes the preference for the universal-wide reading of our test sentences (Ioup, 1975). Secondly, and more importantly, the responses reveal a difference of only 1.7% in the percentage of universal-wide responses in both prime conditions (Figure 2; upper panel).

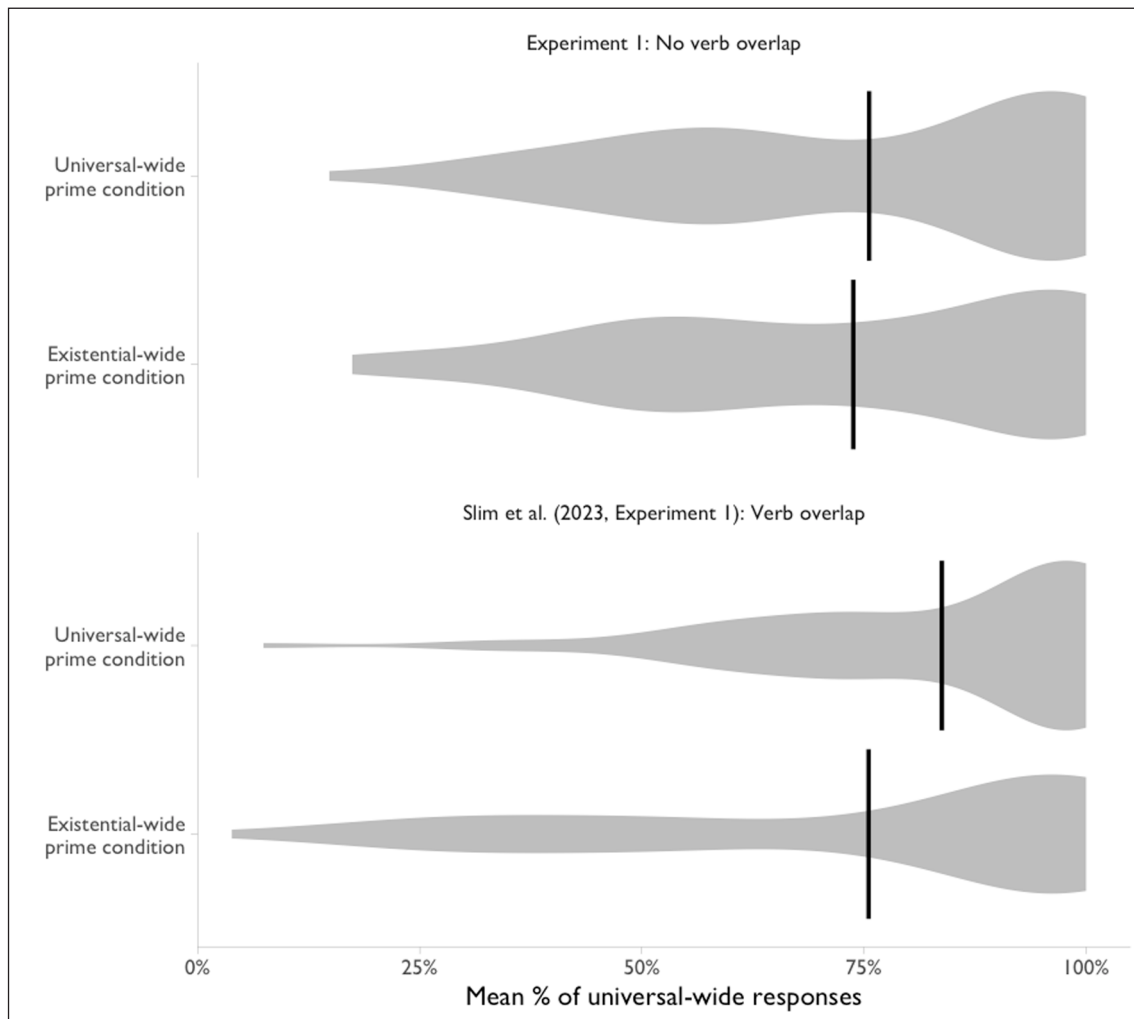


Figure 2: The participants' mean percentage of universal-wide responses on the target trials in Experiment 1 (upper panel) as a function of prime condition, compared to the participants' mean percentage of universal-wide responses on the target trials in the within-verb dataset (lower panel). The black vertical lines represent the overall mean response rate, and the width of the outlined area represents the proportion of the data located at that point.

The mixed-effect model that contained the maximal random-effects structure converged, so we used this model in our analyses. This is a small deviation from the analysis procedure of Feiman and Snedeker (2016), who removed the correlation between the random slopes and random intercepts from their model. We, however, decided to retain this parameter, since a log-likelihood comparison test between the full model and a model that did not include the correlation parameter indicated that the maximal effect structure significantly improved model fit ($\chi^2(2) = 6.99, p = 0.03$).

Log-likelihood tests on the full model revealed no main effect of Prime Condition ($\chi^2(1) = 0.96, p = 0.328$). The analysis on the combined dataset of Experiment 1 and the within-verbs data of Slim et al. (2023, Experiment 1) (**Figure 2**; lower panel) revealed a significant interaction between Prime Condition and Experiment ($\chi^2(1) = 5.74, p = 0.017$), no main effect of Experiment ($\chi^2(1) = 3.26, p = 0.071$), and no main effect of Prime Condition ($\chi^2(1) = 2.47, p = 0.116$). Pairwise comparisons on the effect of Prime Condition in the within-verb and between-verb experiments, carried out using the `testInteractions()` function from R's `phia` package (De Rosario-Martinez, 2015), revealed that the effect of Prime Condition is significant in the data from Slim et al.'s (2023) within-verb experiment ($\chi^2(1) = 8.720, p = 0.006$), but again not in the data of the current between-verb experiment carried out here in Experiment 1 ($\chi^2(1) = 0.004, p = 0.949$).

These results do not replicate Feiman and Snedeker's (2016) finding that logical representation priming is not modulated by verb repetition: Unlike Feiman and Snedeker, we find no strong evidence for priming of logical representations between different verbs, and our analyses show stronger logical representation priming if prime and target involve the same verb. Logical representation priming between verbs is therefore, at the least, not a robust effect. It is unclear, however, whether Feiman and Snedeker or the present experiment were sufficiently powered to detect effects of between-verb logical representation priming. Therefore, we ran a power analysis on the original data of Feiman and Snedeker's Experiment 3.

3. Power analysis

To gain insight into the reliability of logical representation priming between verbs, we conducted power simulations on the original data of Feiman and Snedeker's (2016) Experiment 3. First, we estimated the effect size observed in Feiman and Snedeker's data, by retracting the coefficient of interest from a logit mixed-effect model constructed to analyse the effect of priming in Feiman and Snedeker's Experiment 3 (inspired by Harrington Stack et al., 2018).

As described in Section 2, the materials, design, and procedure of Feiman and Snedeker's (2016) Experiment 3 were similar to those of the present Experiment 1. Moreover, both experiments were carried out remotely over the internet (albeit with different populations: We recruited Dutch-speaking first-year psychology students, whereas Feiman and Snedeker recruited a more diverse group of English-speaking participants on Amazon MTurk). Feiman and Snedeker's

test sentences were English doubly-quantified sentences that contained *every* and *a* (e.g., *Every hiker climbed a hill*) and the verb differed between prime and target. Feiman and Snedeker tested 107 English-speaking participants. Each participant was presented with 12 universal-wide prime-target pairs and 12 existential-wide prime-target pairs.

Based on a re-analysis of Feiman and Snedeker’s data, we extracted the effect size of the prime effect. Then, we ran power simulations to estimate the number of observations required to obtain 80% and 95% power. We aimed for 95% power, rather than commit to the usual 80% power threshold, to further reduce the risk of a false result.

3.1 Re-analysis of Feiman and Snedeker’s (2016) Experiment 3

Feiman and Snedeker’s (2016) Experiment 3 revealed a numerical priming effect of 4.98% of logical representation priming between verbs (**Figure 3**). Feiman and Snedeker analysed the data using logit mixed-effect models (similar to our analysis procedure in Experiment 1). We re-analysed their original data by modelling the response-type likelihood using logit mixed-effect models as well. The model constructed in this analysis was subsequently used in our power simulations. The model was similar to the one constructed in the analyses of Experiment 1 above: It contained the binomial target selection (`true` if universal-wide and `false` if existential-wide) as the dependent variable and Prime Condition as a (sum-coded) predictor variable. The random-effects structure contained random slopes and intercepts for Item and Subject. See **Table 1** for the output of the full model.

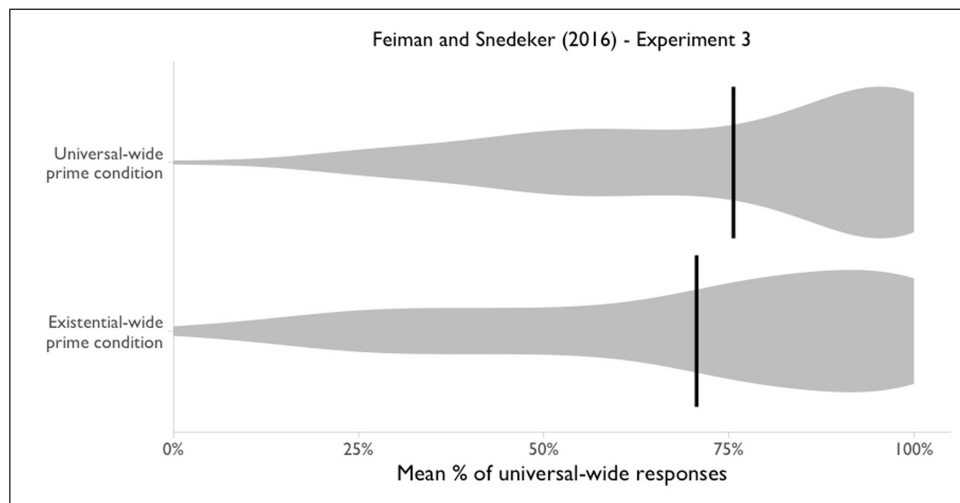


Figure 3: The participants’ mean percentage of universal-wide responses on the target trials in Feiman and Snedeker’s (2016) Experiment 3. The responses are broken down per prime condition. The black vertical lines represent the overall mean response rate, and the width of the outlined area represents the proportion of the data located at that point. The results show a descriptive priming effect of 4.98%.

Table 1: Model output of the logit mixed-effect model constructed to re-analyse Feiman and Snedeker’s (2016) Experiment 3.

	β	se	z-value	p-value
Prime Condition	−0.208	0.089	−2.345	0.019

Note that Feiman and Snedeker (2016) used a slightly reduced random effects structure in their original analyses, since they removed the parameter of correlation between the random slope and intercept. We followed the analysis procedure of Experiment 1 and Barr et al.’s (2013) recommendations, especially because the results of Experiment 1 revealed that the maximal random-effects structure significantly improved model fit (see Section 2.2).

The significant effect of Prime Condition indicates that significantly more universal-wide responses were given in the universal-wide condition compared to the existential-wide condition. In our power analysis, we estimated the power to find an effect of the size observed in this analysis (as expressed by the β -coefficient of −0.208).

3.2 Power simulation and analysis

At the core of simulation-based power analysis is the simulation of new datasets based on the existing dataset. All these simulated datasets are analysed for significance using the same model and procedure as described in the previous subsection. Power is defined as the proportion of significant outcomes in the complete set of simulations.

We conducted this power analysis using the `mixedpower` package in R (Kumle et al., 2021). First, we estimated the required sample size to obtain 95% power if we take the design of Feiman and Snedeker’s (2016) original experiment (with 12 items in each prime condition). We estimated the power based on 1000 simulations for a sample size of 80, 100, 120, 140, 160, 180, 200, and 220 participants. To reach 80% power with 12 items in each prime condition, 140 participants would be needed. To reach 95% power with 12 items in each prime condition, 220 participants would be required (**Figure 4**, left panel). This is almost double the number of participants that Feiman and Snedeker originally tested.

Second, we repeated this power analysis, but now simulated the data as if there were 27 items in each prime condition (like in Experiment 1). This power analysis revealed that 120 participants are needed to obtain 80% power and 160 participants are needed to obtain 95% power (**Figure 4**, right panel). This is roughly 100 participants more than we tested in Experiment 1. Thus, the power analyses revealed that both Feiman and Snedeker’s (2016) experiment and our Dutch replication were underpowered.

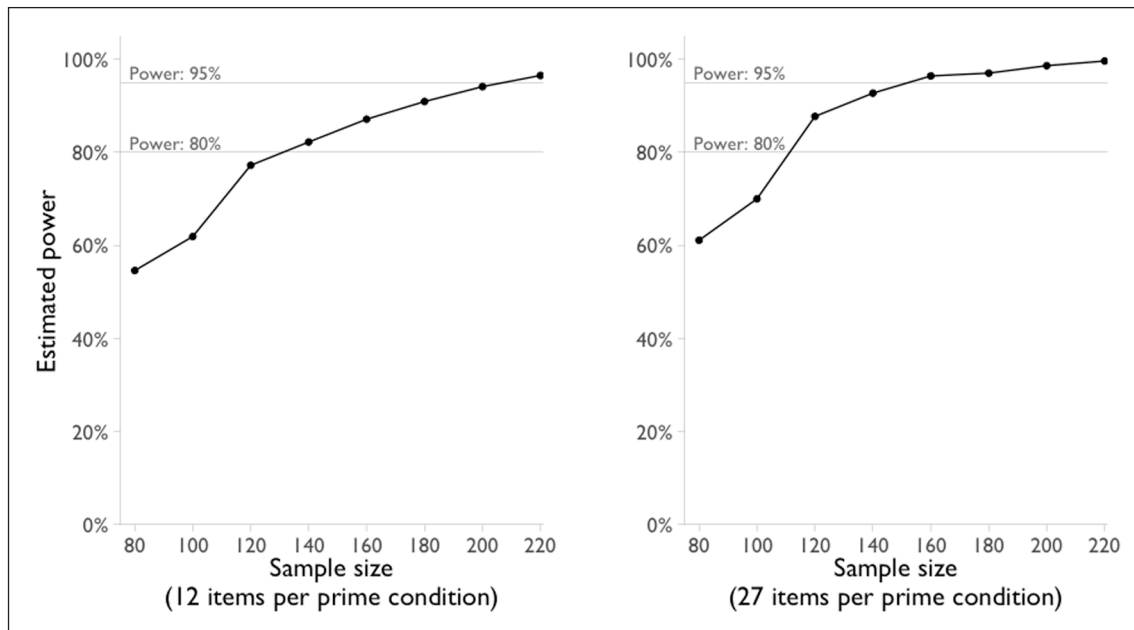


Figure 4: The required sample size to obtain power levels. The left panel shows that 140 participants are needed to obtain 80% power and 220 participants would be needed to obtain 95% if the experiment contained 12 items per prime condition. In the right panel, you can see that 120 participants are needed to reach the 80% power threshold and 160 participants are needed to reach the 95% power threshold if the experiment involved 27 items per prime condition.

4. Experiment 2

In Experiment 2, we tested logical representation priming between verbs again by replicating Experiment 1, but we now determined the sample size based on the outcomes of our power analyses described in Section 2.

4.1 Method

4.1.1 Participants

We recruited 172 native speakers of Dutch. These participants were all recruited via Prolific and were paid £3.75 for their participation. Of the 172 participants, 9 were removed because they guessed the aim of the experiment and 3 participants were removed because they responded incorrectly on more than 10% of the filler items. Thus, 160 participants were included in the analyses (which follows the guidelines to obtain 95% power, as calculated in our power analysis).

4.1.2 Materials and procedure

The materials and procedure were identical to Experiment 1.

4.2 Analyses and results

4.2.1 Data treatment and analyses procedure

The data treatment and analyses procedure were similar to those in Experiment 1.

4.2.2 Results and discussion

Due to an error in the counterbalancing, list A contained 28 universal-wide items and 26 existential-wide items, while list B contained 27 items in both prime conditions. We removed 186 target responses in the universal-wide prime condition (4.28%) and 132 target responses in the existential-wide prime condition (3.11%) due to incorrect responses to the preceding prime trial. The remaining responses indicate that there was a small difference of 1.5% between the prime conditions in the percentage of universal-wide target responses (**Figure 5**, top panel).

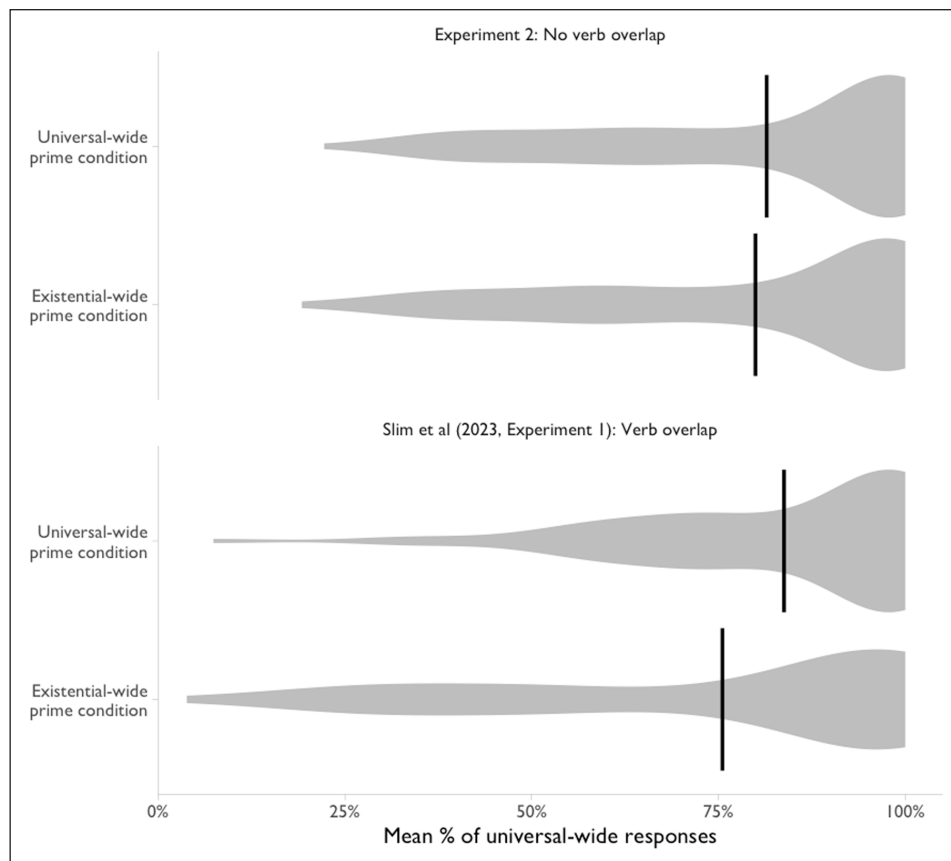


Figure 5: The participants' mean percentage of universal-wide responses on the target trials in Experiment 2 (upper panel) as a function of prime condition, compared to the participants' mean percentage of universal-wide responses on the target trials in the within-verb dataset (lower panel). The black vertical lines represent the overall mean response rate, and the width of the outlined area represents the proportion of the data located at that point.

Like the model constructed in the analysis of Experiment 1, the random-effects structure of the model constructed in the analysis of Experiment 2 was maximal. The log-likelihood test showed no main effect of Prime Condition ($\chi^2(1) = 1.98, p = 0.159$). This suggests that Experiment 2 also failed to detect a reliable effect of logical representation priming between verbs, thereby replicating the null effect observed in Experiment 1.

Additionally, a comparison of Experiment 2 and the within-verb data from Slim et al. (2023, Experiment 1) (**Figure 5**; bottom panel) revealed a significant interaction between Prime Condition and Experiment ($\chi^2(1) = 8.11, p = 0.004$), a significant main effect of Prime Condition ($\chi^2(1) = 4.46, p = 0.035$), and no main effect of Experiment ($\chi^2(1) = 0.13, p = 0.715$). This demonstrates stronger logical representation priming within than between verbs. Pairwise comparisons revealed that the effect of Prime Condition is significant in the data of Slim et al.'s (2023) within-verb experiment ($\chi^2(1) = 15.08, p < 0.001$), but again not in the data of the between-verb experiment carried out here in Experiment 2 ($\chi^2(1) = 0.472, p = 0.492$). Thus, like Experiment 1, Experiment 2 fails to detect robust priming of logical representations if the prime and target sentences contain different verbs.

5. Combined analyses

Both Experiments 1 and 2 failed to replicate Feiman and Snedeker's (2016) finding that priming of logical representations is not modulated by verb overlap in the prime-target pair: We did not observe robust logical representation priming between verbs and observed that priming is stronger if the verb is repeated. We further examined the robustness of logical representation priming between verbs by conducting combined analyses of the data of Experiments 1–2.

We conducted two analyses. First, we combined the data of Experiments 1–2 (thereby increasing the number of observations even more compared to Experiment 2) and tested whether there was a robust effect of logical representations. Second, we examined the role of verb overlap on the strength of logical representation priming in Experiments 1–2 by comparing the data to those of Slim et al. (2023, Experiment 1) in which the priming was tested within verbs. However, in Experiment 2, the number of participants ($n = 160$) is higher than in Slim et al.'s within-verb experiment ($n = 63$). Therefore, we conducted an additional analysis in which we randomly took 63 participants from the combined dataset of Experiments 1–2 and tested whether this dataset (i) showed an influence of logical representation priming between verbs and (ii) revealed a modulating effect of verb overlap on logical representation priming if the data are compared to the within-verb dataset. We repeated these analyses 1000 times (selecting 63 different participants on each iteration).

5.1 Combined analysis 1

In our first combined analysis, we combined the data of Experiments 1 and 2 (**Figure 6**). This dataset included the data of 224 participants. Note that the prime-targets in Experiments 1–2

never involved verb overlap, and therefore this dataset only provides insight into priming logical representations between different verbs. The combined data showed that the participants selected the universal-wide response in 79.83% of the target trials following a universal-wide prime and in 78.42% of the target trials following an existential-wide prime (a descriptive priming effect of 1.41%).

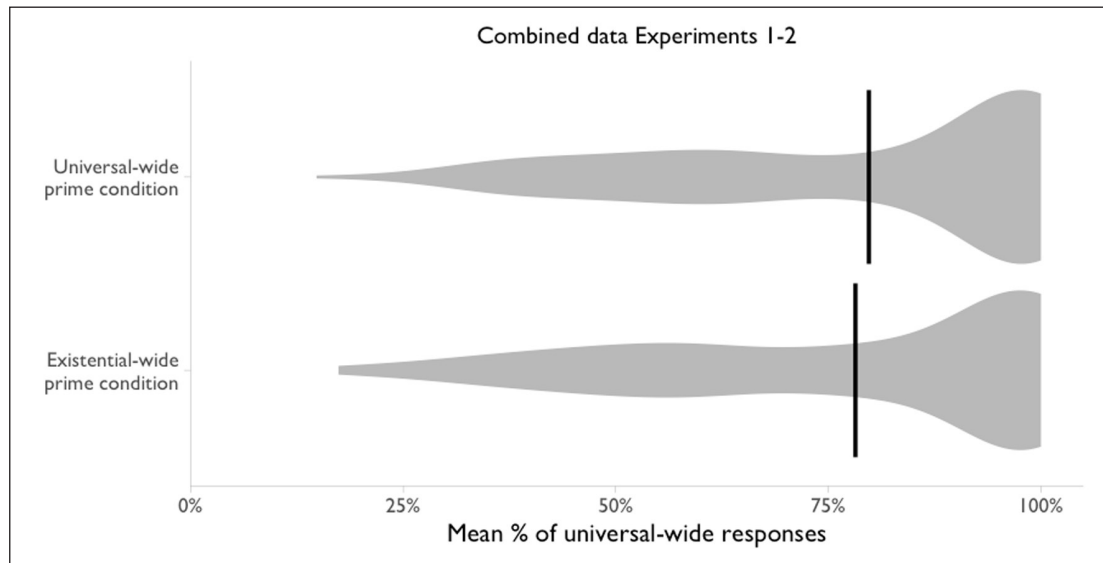


Figure 6: The participants' mean percentage of universal-wide responses on the target trials in Experiments 1–2 combined ($n = 224$), divided between the two prime conditions. The black vertical lines represent the overall mean response rate, and the width of the outlined area represents the proportion of the data located at that point.

We analysed these data using a similar logit mixed-effect model procedure as in Experiments 1-2. First, we constructed a logit mixed-effect model that included Prime Condition (universal-wide vs. existential-wide), Experiment (Experiment 1 vs. Experiment 2), and the interaction between Prime Condition and Experiment as (sum-coded) predictors. The random effects structure was maximal. Identical to Experiments 1–2, the model contained Target Response type as the binomial dependent variable.

The main aim of this combined analysis was to further assess logical representation priming between verbs. Note, however, that this analysis also gained insight into any possible differences in the priming effects in Experiments 1 and 2. Besides the sample size, the only difference between these two experiments is the sample population (for-credit students in Experiment 1 and paid Prolific participants in Experiment 2). Previous studies have shown (descriptively) similar effects of logical representation priming with both student and public populations (e.g., Feiman & Snedeker, 2016; Raffray & Pickering, 2010; Slim et al., 2021), and therefore, we did not predict a difference between these two groups in the strength of logical representation priming.

Similar to the analyses of Experiments 1–2, we calculated p -values by running χ^2 -tests on the log-likelihood values of the full model compared to reduced models. These analyses revealed no main effect of Prime Condition ($\chi^2(1) = 1.28, p = 0.258$), a main effect of Experiment ($\chi^2(1) = 3.94, p = 0.047$), and no interaction between Prime Condition and Experiment ($\chi^2(1) = 0.14, p = 0.712$). Thus, we did not detect a robust effect of logical representation priming between verbs even with a sample size that is larger than that in Experiment 2. The analysis did reveal a difference between the two experiments in the overall target response choices (independent of Prime Condition): participants in Experiment 2 chose the universal-wide response more often than the participants in Experiment 1 (a difference of 5.81%). However, this main effect is difficult to interpret given the difference in sample size between the two experiments. More important is that there was no difference in the priming effect between Experiments 1 and 2, suggesting that priming was unaffected by differences in population between Experiment 1 and Experiment 2.

5.2 Combined analysis 2

In the second combined analysis of Experiments 1 and 2, we tested whether logical representation priming is modulated by verb overlap in the prime-target pairs by combining the results with those of Slim et al. (2023, Experiment 1), in which the prime-target pairs involved the same verb. In Experiment 2, however, the sample size ($n = 160$) is considerably larger than that of the Slim et al.’s within-verb experiment ($n = 63$). Therefore, to further examine the possible modulating effect of verb repetition on logical representation priming in a balanced analysis, we randomly took 63 participants from the combined dataset of Experiments 1–2 and combined that subset of data with the data from our within-verbs experiment. We then used that dataset to run the analysis that tested the interaction between Prime Condition and Experiment. To test the robustness of this interaction, we repeated this analysis 1000 times.

In the first step of each iteration, we tested the data of the 63 randomly selected participants in the between-verb dataset for an effect of Prime Condition by modeling the target-response-type likelihood using a logit mixed-effect model that included Prime Condition as a (sum-coded) predictor. On each iteration, the model only contained a random intercept per Item and Subject, to reduce convergence issues across iterations. Log-likelihood comparisons of the model with a reduced model in which Prime Condition was omitted revealed a significant effect of Prime Condition in 291 of the 1000 iterations ($p \leq 0.05$). Important for our purposes, however, is that the size of the effect of Prime Condition across all 1000 iterations was small compared to the observed effect size in our within-verb dataset (as expressed in log odds ratio, denoted by the estimate coefficient of Prime Condition in the output of each model; Figure 7).

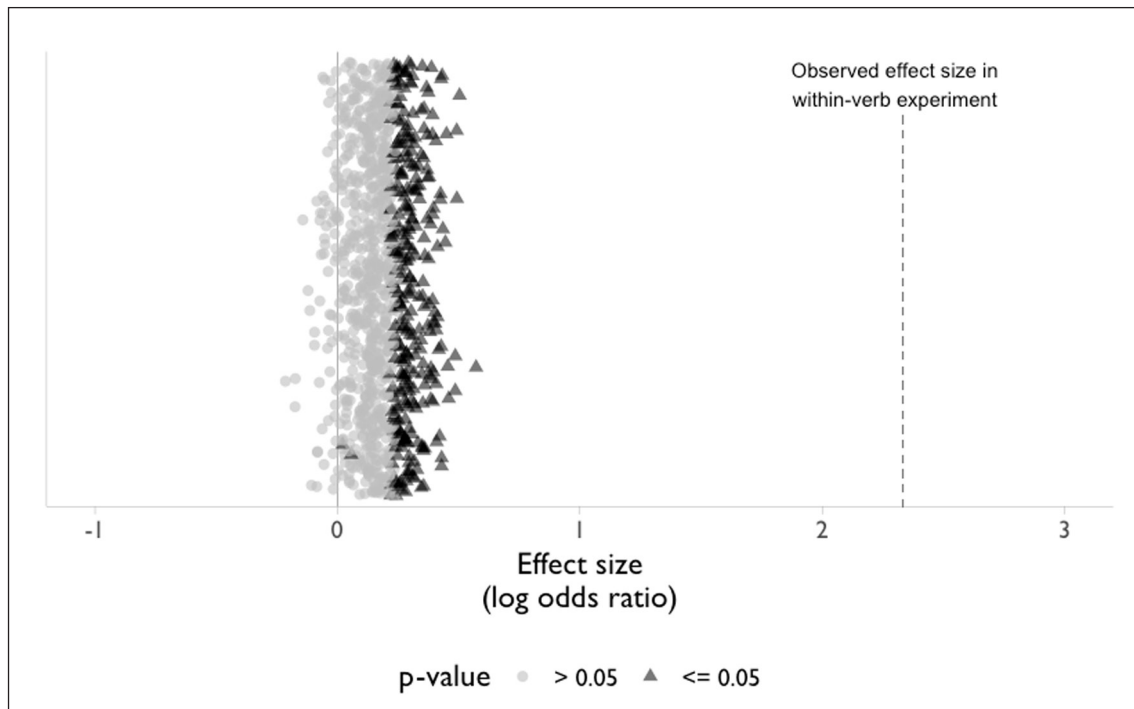


Figure 7: The size of the main effect of Prime Condition in each of the 1000 iterations of our combined analysis. On each iteration, we randomly extracted 63 participants from the total number of participants in the combined dataset of Experiments 1–2 ($n = 224$). The subset of the data selected on each iteration of this analysis was tested for a significant effect of Prime Condition. In this figure, the points on the scatterplot represent the observed effect size (as measured in the log odds ratio in the output of each model) of Prime Condition on each iteration, and the colour/shape indicates whether the effect of Prime Condition was significant on that iteration. The dashed line in the figure indicates the observed effect size in our previously-conducted within-verb experiment (which showed priming of logical representations in the verb is repeated in each prime and target).

In the second step of each iteration, we tested whether priming was modulated by verb repetition by combining each subset of data with our within-verb data and testing this dataset for an interaction between Prime Condition and Verb Overlap. The model constructed on each iteration therefore contained Prime Condition, Verb Overlap, and the interaction between Prime Condition and Verb Overlap as sum-coded predictors. Like in the first step of each iteration, the random-effects structure only contained a random intercept per Item and Subject, to reduce convergence issues across iterations.

These analyses revealed a main effect of Prime Condition in 297 of the 1000 iterations, a main effect of Verb Overlap in 17 of the 1000 iterations, and an interaction between Prime Condition and Verb Overlap in 999 of the 1000 iterations (**Figure 8**). So, taken together, these

analyses revealed that the effect of Prime Condition was significantly stronger if the prime-target pairs contained the same verb in a vast majority of iterations.³

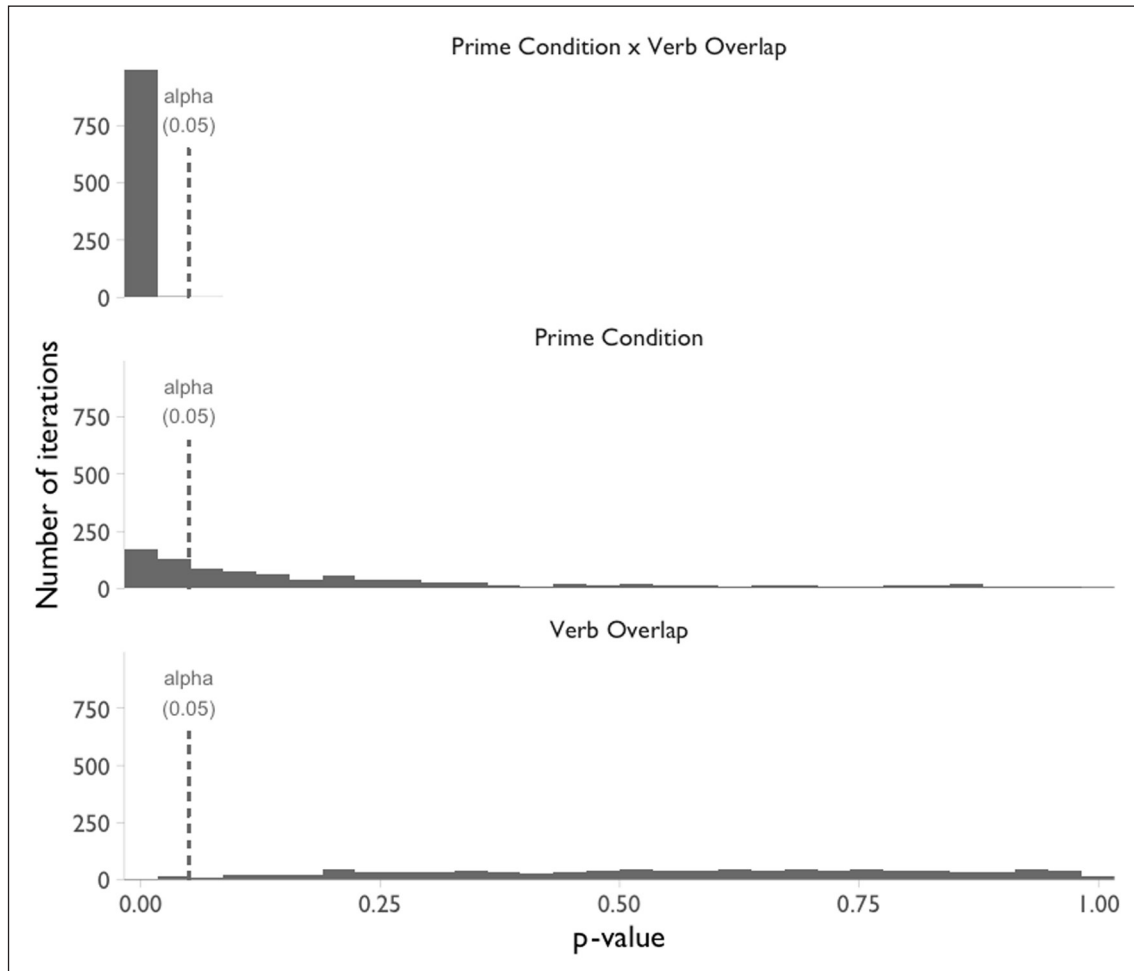


Figure 8: This figure shows the number of iterations in which the Prime Condition x Verb Overlap interaction (top panel), the Prime Condition main effect (middle panel), and the Verb Overlap main effect (bottom panel) were significant. Most important for our research purposes, the Prime Condition x Verb Overlap interaction was significant in all iterations but one.

5.3 Discussion

Our combined analysis of Experiments 1 and 2 did not reveal an effect of Prime Condition between verbs, even though this analysis was highly powered. Additionally, our iterated analyses showed

³ The one iteration that did not reveal an interaction between Prime Condition and Verb Overlap showed a main effect of Prime Condition, so this iteration replicated Feiman and Snedeker's (2016) general pattern of results.

that randomly selecting a smaller group of participants from our combined dataset did result in an effect of Prime Condition between verbs in a minority of cases (but note that such an analysis is underpowered). This effect was virtually always weaker compared to logical representation priming within the same verb.

On the one hand, care is required in interpreting the observed null effects of logical representation priming between verbs. Even though our analysis fails to detect a robust effect of priming between verbs, it is noteworthy that our iteration analysis revealed a descriptive effect in the direction of priming in a clear majority of iterations (**Figure 7**). More importantly, such an effect is very small and therefore not robust. Moreover, our data provide clear evidence for an interaction between Prime Condition and Verb Overlap: Logical representation priming is stronger if the verb is repeated compared to when it is not.

6. General discussion

Feiman and Snedeker (2016) hypothesised that logical representations are content-free abstract representations of the compositional structure of sentence meaning, based on their finding that logical representation priming is not modulated by verb repetition between prime and target. This finding suggests a division of labour between the compositional and the lexical semantic system in sentence interpretation. However, across two experiments we failed to replicate these findings: We consistently observed that logical representation priming between verbs is weaker than logical representation priming within the same verb (even if we increase the sample size to reach more than sufficient power, which we did in Experiment 2 and in our combined analyses). This indicates that logical representations between verbs is not a robust effect.

6.1 Methodological differences: Feiman and Snedeker (2016) and this study

It is possible that the discrepancy in findings between our experiments and Feiman and Snedeker's (2016) experiment is due to methodological differences. However, there are good reasons to assume that this is not the case.

First, although there was overlap between the sets of stimulus materials in our and Feiman and Snedeker's experiments (namely those borrowed from Raffray & Pickering, 2010), we constructed additional materials. Thereby, we increased the number of trials in each experimental condition (27 prime-target pairs per prime condition instead of 12). However, this increase of prime trials does not seem to influence the strength of priming: Previous studies that involved the same materials and contained the same number of items did reveal effects of logical representation priming (Slim et al., 2021). In these studies, however, the verb in the prime and target sentence was repeated (replicating Raffray & Pickering, 2010; and Experiment 2 and Experiment 4 in Feiman & Snedeker, 2016).

Second, the experiments reported here were in Dutch, whereas Feiman and Snedeker's (2016) experiments were in English. It is unlikely that the lack of between-verb priming is due to a difference between Dutch and English. First, the Dutch and English stimuli were close translation equivalents (see <https://osf.io/697wg/> for a list of prime-target pairs). And second, previous studies (Slim et al., 2021) have shown that logical representations can be primed in Dutch language comprehension if the verb is repeated between prime and target (similar to English). In fact, numerous studies have shown that structural priming in general manifests very similarly – if not identically – in Dutch and English (see, e.g. Hartsuiker et al., 2008; Mahowald et al., 2016; Pickering & Ferreira, 2008). Therefore, we see no reason to assume that possible differences between Feiman and Snedeker's results and the current results are driven by differences in comprehension of English and Dutch.

A final difference is that we recruited participants amongst first-year psychology students (Experiment 1) or via Prolific (Experiment 2), whereas Feiman and Snedeker (2016) recruited their participants via Amazon MTurk. Here, it is worth noting that the results of Experiment 1 and 2 reported in the present study are similar, even though there are differences in the demographic characteristics between the undergraduate population tested in Experiment 1 and the more diverse population tested in Experiment 2 (corroborated by our combined analysis). Moreover, the participants from Experiment 2 were paid participants recruited online, similar to the population tested by Feiman and Snedeker. Therefore, we see no reason to assume that demographic and/or recruitment differences affect logical representation priming and caused the discrepancy between the present and Feiman and Snedeker's results.

The most likely explanation for the difference between our findings and Feiman and Snedeker's (2016) findings is the difference in power. Our power analysis revealed that Feiman and Snedeker's experiment was underpowered. Our well-powered Experiment 2 and our highly powered combined analyses again reveal that priming logical representations between verbs is a small effect and therefore not robust, and logical representation priming is consistently enhanced if the verb is repeated in the prime and target.

6.2 Logical representations as event representations

The finding that logical representation priming does not robustly emerge between verbs (but does robustly emerge within verbs; Feiman & Snedeker, 2016; Raffray & Pickering, 2010; Slim et al., 2021) is difficult to reconcile with semantic theories that posit a strict separation between the compositional and lexical system in interpretation. Such accounts propose that comprehenders compute abstract logical representations that only specify structure and not any lexical meaning content (e.g., Feiman & Snedeker, 2016; Heim & Kratzer, 1998; May, 1985). Our findings, however, suggest that comprehenders compute integrated logical representations that specify both compositional structure and conceptual meaning features of (at least) the event denoted in the sentence (as specified by the verb).

Our findings are in line with theoretical accounts that emphasise the role of the lexical system in the computation of logical representations. In Section 1 (the Introduction), we described earlier findings that suggest that conceptual meaning content greatly influences quantifier scope assignment in the form of world knowledge. A sentence like *Every customer returns an item*, for example, is more often assigned a universal-wide interpretation (in which every customer returned a different item) than a sentence like *Every jeweller praised a diamond*, despite the fact that these sentences are structurally identical. Rather, the difference in interpretation is due to our expectations about the typical number of participants and objects involved in “returning” and “praising” events (Dwivedi, 2013; Saba & Coriveau, 2001). Several studies on real-time semantic interpretation have shown that comprehenders immediately integrate such real-world knowledge in the incremental computation of logical representations (e.g., Dwivedi, 2013; Dwivedi et al., 2010; Filik et al., 2004; Urbach et al., 2015; Urbach & Kutas, 2010). These studies indicate that world knowledge is not merely used to choose between alternative abstract analyses, but that expectations based on world knowledge actively guide real-time scope assignment in the developing logical representation. The resulting logical representation therefore captures features of both the compositional structure and lexical meaning content. This account of integrated logical representations fits theories on semantic representation that argue that quantifier scope is specified in conceptual event representations (which specify conceptual meaning features of the event denoted in the sentence and of the number of participants and objects involved in that event; e.g., Fodor, 1982; Jackendoff, 1992; Johnson-Laird et al., 1989).

This description of logical representations explains why logical representation priming is more robust within verbs than between verbs. Recall that structural priming effects emerge because it is easier to reuse representations computed in a preceding instance of language processing (typically explained in terms of *implicit learning*, e.g., Bock & Griffin, 2000; Chang et al., 2006; Jaeger & Snider, 2013). Effects of structural priming are usually stronger if more features of the prime and the target sentence are shared (e.g., Branigan & Pickering, 2017; Ziegler & Snedeker, 2018; Ziegler et al., 2018). Therefore, under the assumption that logical representations specify both structure and event-related conceptual features, it is predicted that abstract logical representation priming (i.e., priming between verbs) is smaller than event-specific logical representation priming (i.e., priming within verbs). This makes abstract logical representation a fragile effect. In addition (and more speculatively), the prime sentences in our experiment may lead participants to update their biases on the typical number of participants and objects involved in a specific event (as previously suggested by Dwivedi, 2013). For example, the ‘climbing’ event in *Every kid climbed a tree* is biased towards a universal-wide interpretation (Dwivedi et al., 2010; Feiman & Snedeker, 2016). Encountering this sentence in a universal-wide prime could strengthen this initial bias and encountering this sentence in an existential-wide prime weakens this initial bias (e.g., Jaeger & Snider, 2013; see also Yildirim et al., 2016). This bias-updating, in turn, could affect the interpretation of a subsequent target sentence if

that target sentence describes a similar event. This is the case in within-verb priming, but not in between-verb priming. Note that priming can simultaneously emerge on multiple levels of processing (see Branigan & Pickering, 2017), and therefore these two proposed loci of priming are not mutually exclusive.

Finally, we should note that we are not arguing that people are *unable* to compute abstract logical representations. In fact, certain contexts require the processing of quantifier relations in the absence of encyclopedic knowledge, such as logic puzzles or syllogistic reasoning tasks (which are often incorporated in standardised tests like the *Law School Admissions Test* or the *Graduate Record Examinations*; AnderBois et al., 2012; Khemlani et al., 2012). In such tasks, quantified sentences are typically presented in a context that is set up in such a way that effects of real-world knowledge are minimised. Nevertheless, people are able to draw inferences from such quantified statements (e.g., Bucciarelli & Johnson-Laird, 1999; Khemlani & Johnson-Laird, 2012). This suggests that, at least in contexts such as logic puzzles, people can compute logical representations that abstract away from lexico-semantic knowledge. At the same time, syllogistic reasoning is notoriously prone to errors, and these are often driven by real-world knowledge (e.g., concluding that *all sparrows are birds* based on the premises *all birds can fly* and *all sparrows can fly*, even though this intuitive response based on real-world knowledge is not supported by the two premises; e.g., Klauer et al., 2000). This suggests that the computation of abstract logical representations is cognitively burdensome, and comprehenders therefore often rely on additional information stored in the lexical system in the assignment of scope (Dwivedi, 2013).

6.3 Is there a lexical boost in logical representation priming?

Structural priming research on syntactic representations has shown the *lexical boost* effect: Syntactic priming is enhanced if the verb is repeated in the prime and target (e.g., Branigan & Pickering, 2017; Hartsuiker et al., 2008; Pickering & Branigan, 1998; see also Mahowald et al., 2016, for meta-analysis). The present results show that logical representation priming is also enhanced if the prime-target pair contains the same verb, and therefore, this finding may be interpreted as evidence for a lexical boost effect in logical representation priming. Recall, however, that the evidence of the lexical boost in priming syntactic representations in comprehension is piecemeal, with some studies showing no indications of a lexical boost in comprehension (e.g., Arai et al., 2007; Thothathiri & Snedeker, 2008; Ziegler et al., 2019; for review, see Tooley, 2022). Since we studied structural priming in comprehension, it is difficult to relate our findings to the known effects from research on syntactic priming.

In addition, there are theoretical considerations that lead one to assume that the lexical boost — as it has been characterised in syntactic priming — does not emerge in logical representation priming. There are two main theoretical accounts on the lexical boost. According to the first explanation, the lexical boost is explained in terms of larger representational overlap between

prime and target. The verb largely determines a sentence's syntactic structure because it specifies argument structure. Verbs are therefore assumed to be connected to syntactic structures in our mental lexicon (Pickering & Branigan, 1998). In case the prime and target contain the same verb, both the abstract syntactic structure and the connection between the verb and the syntactic structure are shared between prime and target. This larger representational overlap enhances priming (e.g., Hartsuiker et al., 2008; Pickering & Branigan, 1998).

The second explanation of the lexical boost effect posits that verb repetition reactivates a short-term memory trace for the structure of the prime sentence, which boosts syntactic priming (e.g., Bock & Griffin, 2000; Hartsuiker et al., 2008; Segaert et al., 2013). This account, however, requires additional assumptions about why the lexical boost effect robustly emerges if the verb is repeated and not if other lexical items are repeated (Carminati et al., 2019; Mahowald et al., 2016; but cf. Kantola et al., 2023; Scheepers et al., 2017). To account for this finding, theorists have put forward the hypothesis that the verb acts as a stronger memory cue for the same reasons as described above: Because it is the head of the syntactic phrase (e.g., Hartsuiker et al., 2008; Tooley & Traxler, 2010).

Importantly, verbs are not necessarily encoded in the same way in syntactic representations and in logical representations. Syntactic representations are grammatical linguistic representations, whereas we described logical representations as non-linguistic representations of conceptual thought (e.g., Fodor, 1982; Jackendoff, 1992, 2002; but cf. Barker & Jacobson, 2007; May, 1985). Therefore, syntactic representations specify how lexical items combine in syntactic phrase structure and logical representations specify how concepts are combined in combinatorial thought (Fodor, 1982; Jackendoff, 1992). In turn, the enhancing effect of the lexical boost in syntactic priming is due to repetition of the same lexical item (which, in the case of verbs, specifies the possible syntactic structures; Pickering & Branigan, 1998) and not due to repetition of the same conceptual content. Moreover, this difference between syntactic and logical representation priming presupposes that the effect of verb repetition is not the same in both types of priming. In syntactic priming, verb repetition enhances priming because the verb specifies syntactic structure. In logical representation priming, we hypothesised that verb repetition enhances priming because the verb specifies much semantic information about the event denoted in the sentence. However, this information is not part of the grammatical features of a verb, but of its (non-linguistic) conceptual semantics (Jackendoff, 1992; Levin & Rappaport Hovav, 2005; see also Fillmore, 2006).

One possible way to gain insight into the assumed distinction in verb representation at the level of syntactic and logical representations is by studying *bilingual* priming. In bilingual priming, the prime and the target trials are presented in different languages to bilingual participants (see Van Gompel & Arai, 2018, for review). Studies on bilingual syntactic priming have shown that the lexical boost is stronger in within-language priming than in between-language priming (in

which case not the same lexical item, but a *translation equivalent* is repeated, e.g., Schoonbaert et al., 2007). This indicates that the lexical boost in syntactic priming is not driven by repetition of conceptual semantics (which are assumed to be shared across translation equivalents, e.g., De Bot, 1992, 2004). However, the first insights that we have on bilingual priming of logical representations reveal a different pattern: A study by Slim et al. (2021) showed that logical representation priming is comparable within the same language (if the verb is repeated) and between languages (if a translation equivalent of the verb is repeated). This suggests that repetition of concepts, rather than repetition of the same word form, enhances logical representation priming. Some care is required in interpreting these results, given the scarce number of studies on bilingual logical representation priming. Nevertheless, the results so far seem to support the assumed dichotomy between the role of verb repetition in syntactic and logical representation priming, and show that this strand of research could provide insight into commonalities and differences in the encoding of verb information in syntactic and logical representations.

Summing up, we hypothesised that the effect of verb repetition manifests differently in logical representation priming than how it is typically described in syntactic priming: Concept repetition (rather than lexical repetition) magnifies effects of logical representation priming. Of course, the arguments presented above are based on theoretical considerations and further empirical research is therefore needed to fully understand the mechanisms underlying the role of lexical repetition in syntactic and logical representation priming.

7. Conclusion

Feiman and Snedeker (2016, Experiment 3) observed that logical representation priming is not sensitive to any effects of verb overlap between prime and target. Two experiments in Dutch, however, showed a different pattern of results. Unlike Feiman and Snedeker, we found that logical representation priming is enhanced when prime and target contain the same verb as opposed to when prime and target involve different verbs. This indicates that logical representation priming between verbs is, at least, not a robust effect. We argued that this finding is difficult to reconcile with prevalent semantic theories that posit a strict distinction between the lexical and the compositional system in the computation of sentence meaning. Rather, we argued that the lexical and the compositional system interact in language processing, which leads to integrated logical representations that specify both content and structure of sentence meaning.

Data accessibility statement

All data and analyses scripts are freely available on <https://osf.io/697wg/>. The materials used in the experiment are available on <https://osf.io/v2w3a/>. Experiment 2 was preregistered (see <https://osf.io/697wg/>).

Ethics and consent

All participants gave informed consent prior to participation. The study reported in this paper has been approved by the Ethical Committee of the Faculty of Psychology and Educational Sciences at Ghent University.

Acknowledgements

We would like to thank Roman Feiman for sharing the data of Feiman and Snedeker's (2016) Experiment 3, which we used in our power analysis.

Competing interests

The authors have no competing interests to declare.

Authors' contributions

Mieke Sarah Slim: Conceptualization, Data curation, Formal analysis, Funding acquisition, Investigation, Methodology, Project administration, Software, Visualization, Writing – original draft, Writing – review & editing; **Peter Lauwers:** Conceptualization, Methodology, Funding acquisition, Supervision, Writing – review & editing; **Robert J. Hartsuiker:** Conceptualization, Funding acquisition, Methodology, Project administration, Supervision, Writing – review & editing

References

- Altmann, G., & Ekves, Z. (2019). Events as intersecting object histories: A new theory of event representation. *Psychological Review*, 126(6), 817. DOI: <https://doi.org/10.1037/rev0000154>
- Altmann, G., & Steedman, M. (1988). Interaction with context during human sentence processing. *Cognition*, 30(3), 191–238. DOI: [https://doi.org/10.1016/0010-0277\(88\)90020-0](https://doi.org/10.1016/0010-0277(88)90020-0)
- AnderBois, S., Brasoveanu, A., & Henderson, R. (2012). The pragmatics of quantifier scope: A corpus study. *Proceedings of Sinn und Bedeutung*, 16(1), 15–28.
- Arai, M., Van Gompel, R. P., & Scheepers, C. (2007). Priming ditransitive structures in comprehension. *Cognitive Psychology*, 54(3), 218–250. DOI: <https://doi.org/10.1016/j.cogpsych.2006.07.001>
- Barker, C., & Jacobson, P. (2007). *Direct compositionality* (Vol. 14). Oxford University Press.

- 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. DOI: <https://doi.org/10.1016/j.jml.2012.11.001>
- Bates, D., Kliegl, R., Vasishth, S., & Baayen, H. (2015). *Parsimonious mixed models* [preprint]. arXiv:1506.04967
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2014). *Fitting linear mixed-effects models using lme4* [preprint]. arXiv:1406.5823. DOI: <https://doi.org/10.18637/jss.v067.i01>
- Bock, J. K. (1986). Syntactic persistence in language production. *Cognitive Psychology*, 18(3), 355–387. DOI: [https://doi.org/10.1016/0010-0285\(86\)90004-6](https://doi.org/10.1016/0010-0285(86)90004-6)
- Bock, K., & Griffin, Z. M. (2000). The persistence of structural priming: Transient activation or implicit learning? *Journal of Experimental Psychology: General*, 129(2), 177. DOI: <https://doi.org/10.1037/0096-3445.129.2.177>
- Bock, K., & Loebell, H. (1990). Framing sentences. *Cognition*, 35(1), 1–39. DOI: [https://doi.org/10.1016/0010-0277\(90\)90035-I](https://doi.org/10.1016/0010-0277(90)90035-I)
- Branigan, H. P., & Pickering, M. J. (2017). An experimental approach to linguistic representation. *Behavioral and Brain Sciences*, 40. DOI: <https://doi.org/10.1017/S0140525X16002028>
- Branigan, H. P., Pickering, M. J., & McLean, J. F. (2005). Priming prepositional-phrase attachment during comprehension. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 31(3), 468. DOI: <https://doi.org/10.1037/0278-7393.31.3.468>
- Brasoveanu, A., & Dotlacil, J. (2019). Quantification. In C. Cummins & N. Katsos (Eds.), *The Oxford handbook of experimental semantics and pragmatics*. Oxford University Press. DOI: <https://doi.org/10.1093/oxfordhb/9780198791768.013.3>
- Bucciarelli, M., & Johnson-Laird, P. N. (1999). Strategies in syllogistic reasoning. *Cognitive Science*, 23(3), 247–303. DOI: https://doi.org/10.1207/s15516709cog2303_1
- Carminati, M. N., van Gompel, R. P., & Wakeford, L. J. (2019). An investigation into the lexical boost with nonhead nouns. *Journal of Memory and Language*, 108, 104031. DOI: <https://doi.org/10.1016/j.jml.2019.104031>
- Chang, F., Dell, G. S., & Bock, K. (2006). Becoming syntactic. *Psychological Review*, 113(2), 234. DOI: <https://doi.org/10.1037/0033-295X.113.2.234>
- Chomsky, N. (1995). *The minimalist program*. MIT Press.
- De Bot, K. (1992). A bilingual production model: Levelt's 'speaking' model adapted. *Applied Linguistics*, 13, 1–24.
- De Bot, K. (2004). The multilingual lexicon: Modelling selection and control. *International Journal of Multilingualism*, 1(1), 17–32. DOI: <https://doi.org/10.1080/14790710408668176>
- De Rosario-Martinez, H. (2015). *Phia: Post-hoc interaction analysis* [R package]. <https://CRAN.R-project.org/package=phia>
- Dwivedi, V. D. (2013). Interpreting quantifier scope ambiguity: Evidence of heuristic first, algorithmic second processing. *PloS One*, 8(11), e81461. DOI: <https://doi.org/10.1371/journal.pone.0081461>

- Dwivedi, V. D., Phillips, N. A., Einagel, S., & Baum, S. R. (2010). The neural underpinnings of semantic ambiguity and anaphora. *Brain Research*, 1311, 93–109. DOI: <https://doi.org/10.1016/j.brainres.2009.09.102>
- Feiman, R., & Snedeker, J. (2016). The logic in language: How all quantifiers are alike, but each quantifier is different. *Cognitive Psychology*, 87, 29–52. DOI: <https://doi.org/10.1016/j.cogpsych.2016.04.002>
- Filik, R., Paterson, K. B., & Livsersedge, S. P. (2004). Processing doubly quantified sentences: Evidence from eye movements. *Psychonomic Bulletin & Review*, 11(5), 953–959. DOI: <https://doi.org/10.3758/BF03196727>
- Fillmore, C. J. (2006). Frame semantics. In D. Geeraerts (Ed.), *Cognitive linguistics: Basic readings*. Mouton de Gruyter. DOI: <https://doi.org/10.1016/B0-08-044854-2/00424-7>
- Fodor, J. D. (1982). The mental representation of quantifiers. In S. Peeters & S. Esa (Eds.), *Processes, beliefs, and questions*. Springer. DOI: https://doi.org/10.1007/978-94-015-7668-0_5
- Frankland, S. M., & Greene, J. D. (2020). Concepts and compositionality: In search of the brain's language of thought. *Annual Review of Psychology*, 71, 273–303. DOI: <https://doi.org/10.1146/annurev-psych-122216-011829>
- Goodwin, G. P., & Johnson-Laird, P. (2011). Mental models of Boolean concepts. *Cognitive Psychology*, 63(1), 34–59. DOI: <https://doi.org/10.1016/j.cogpsych.2011.04.001>
- Harrington Stack, C. M., James, A. N., & Watson, D. G. (2018). A failure to replicate rapid syntactic adaptation in comprehension. *Memory & Cognition*, 46(6), 864–877. DOI: <https://doi.org/10.3758/s13421-018-0808-6>
- Hartsuiker, R. J., Berolet, S., Schoonbaert, S., Speybroeck, S., & Vanderelst, D. (2008). Syntactic priming persists while the lexical boost decays: Evidence from written and spoken dialogue. *Journal of Memory and Language*, 58(2), 214–238. DOI: <https://doi.org/10.1016/j.jml.2007.07.003>
- Heim, I., & Kratzer, A. (1998). *Semantics in generative grammar*. Blackwell.
- Ioup, G. (1975). Some universals for quantifier scope. In J. P. Kimball (Ed.), *Syntax and semantics*. Brill. DOI: https://doi.org/10.1163/9789004368828_003
- Jackendoff, R. S. (1992). *Semantic structures* (Vol. 18). MIT Press.
- Jackendoff, R. S. (2002). *Foundations of language: Brain, meaning, grammar, evolution*. Oxford University Press. DOI: <https://doi.org/10.1093/acprof:oso/9780198270126.001.0001>
- Jaeger, T. F. (2008). Categorical data analysis: Away from anovas (transformation or not) and towards logit mixed models. *Journal of Memory and Language*, 59(4), 434–446. DOI: <https://doi.org/10.1016/j.jml.2007.11.007>
- Jaeger, T. F., & Snider, N. E. (2013). Alignment as a consequence of expectation adaptation: Syntactic priming is affected by the prime's prediction error given both prior and recent experience. *Cognition*, 127(1), 57–83. DOI: <https://doi.org/10.1016/j.cognition.2012.10.013>
- Johnson-Laird, P. N. (1983). *Mental models: Towards a cognitive science of language, inference, and consciousness*. Harvard University Press.

- Johnson-Laird, P. N., Byrne, R. M., & Tabossi, P. (1989). Reasoning by model: The case of multiple quantification. *Psychological Review*, 96(4), 658. DOI: <https://doi.org/10.1037/0033-295X.96.4.658>
- Kantola, L., van Gompel, R. P., & Wakeford, L. J. (2023). The head or the verb: Is the lexical boost restricted to the head verb? *Journal of Memory and Language*, 129, 104388. DOI: <https://doi.org/10.1016/j.jml.2022.104388>
- Khemlani, S., & Johnson-Laird, P. N. (2012). Theories of the syllogism: A meta-analysis. *Psychological Bulletin*, 138(3), 427. DOI: <https://doi.org/10.1037/a0026841>
- Khemlani, S., Orenes, I., & Johnson-Laird, P. N. (2012). Negation: A theory of its meaning, representation, and use. *Journal of Cognitive Psychology*, 24(5), 541–559. DOI: <https://doi.org/10.1080/20445911.2012.660913>
- Klauer, K. C., Musch, J., & Naumer, B. (2000). On belief bias in syllogistic reasoning. *Psychological Review*, 107(4), 852. DOI: <https://doi.org/10.1037/0033-295X.107.4.852>
- Kumle, L., Vö, M. L.-H., & Draschkow, D. (2021). Estimating power in (generalized) linear mixed models: An open introduction and tutorial in R. *Behavior Research Methods*, 53(6), 2528–2543. DOI: <https://doi.org/10.3758/s13428-021-01546-0>
- LeBel, E. P., Berger, D., Campbell, L., & Loving, T. J. (2017). Falsifiability is not optional. *Journal of Personality and Social Psychology*, 113(2), 254–261. DOI: <https://doi.org/10.1037/pspi0000106>
- Levin, B., & Rappaport Hovav, M. (2005). *Argument realization*. Cambridge University Press. DOI: <https://doi.org/10.1017/CBO9780511610479>
- Mahowald, K., James, A., Futrell, R., & Gibson, E. (2016). A meta-analysis of syntactic priming in language production. *Journal of Memory and Language*, 91, 5–27. DOI: <https://doi.org/10.1016/j.jml.2016.03.009>
- May, R. (1985). *Logical form: Its structure and derivation* (Vol. 12). MIT Press.
- Piantadosi, S. T., Goodman, N. D., Ellis, B. A., & Tenenbaum, J. (2008). A Bayesian model of the acquisition of compositional semantics. *Proceedings of the thirtieth annual conference of the Cognitive Science Society*, 1620–1625.
- Pickering, M. J., & Branigan, H. P. (1998). The representation of verbs: Evidence from syntactic priming in language production. *Journal of Memory and language*, 39(4), 633–651. DOI: <https://doi.org/10.1006/jmla.1998.2592>
- Pickering, M. J., & Ferreira, V. S. (2008). Structural priming: A critical review. *Psychological Bulletin*, 134(3), 427. DOI: <https://doi.org/10.1037/0033-2909.134.3.427>
- R Core Team. (2019). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing.
- Raffray, C. N., & Pickering, M. J. (2010). How do people construct logical form during language comprehension? *Psychological Science*, 21(8), 1090–1097. DOI: <https://doi.org/10.1177/0956797610375446>

- Ruys, E. G., & Winter, Y. (2011). Quantifier scope in formal linguistics. In D. Gabbay & F. Guenther (Eds.), *Handbook of philosophical logic*. Springer. DOI: https://doi.org/10.1007/978-94-007-0479-4_3
- Saba, W. S., & Coriveau, J.-P. (2001). Plausible reasoning and the resolution of quantifier scope ambiguities. *Studia Logica*, 67(2), 271–289. DOI: <https://doi.org/10.1023/A:1010503321412>
- Scheepers, C., Raffray, C. N., & Myachykov, A. (2017). The lexical boost effect is not diagnostic of lexically specific syntactic representations. *Journal of Memory and Language*, 95, 102–115. DOI: <https://doi.org/10.1016/j.jml.2017.03.001>
- Schoonbaert, S., Duyck, W., Brysbaert, M., & Hartsuiker, R. J. (2009). Semantic and translation priming from a first language to a second and back: Making sense of the findings. *Memory & Cognition*, 37(5), 569–586. DOI: <https://doi.org/10.3758/MC.37.5.569>
- Schoonbaert, S., Hartsuiker, R. J., & Pickering, M. J. (2007). The representation of lexical and syntactic information in bilinguals: Evidence from syntactic priming. *Journal of Memory and Language*, 56(2), 153–171. DOI: <https://doi.org/10.1016/j.jml.2006.10.002>
- Segaert, K., Kempen, G., Petersson, K. M., & Hagoort, P. (2013). Syntactic priming and the lexical boost effect during sentence production and sentence comprehension: An fMRI study. *Brain and language*, 124(2), 174–183. DOI: <https://doi.org/10.1016/j.bandl.2012.12.003>
- Slim, M. S., Lauwers, P., & Hartsuiker, R. (2023). *Revisiting the logic in language: The scope of each and every universal quantifier is alike after all* [preprint]. PsyArXiv: <https://psyarxiv.com/jgcxy/>. DOI: <https://doi.org/10.31234/osf.io/jgcxy>
- Slim, M. S., Lauwers, P., & Hartsuiker, R. J. (2021). Monolingual and bilingual logical representations of quantificational scope: Evidence from priming in language comprehension. *Journal of Memory and Language*, 116, 104184. DOI: <https://doi.org/10.1016/j.jml.2020.104184>
- Szabolcsi, A. (1997). Strategies for scope taking. In A. Szabolcsi (Ed.), *Ways of scope taking*. Springer. DOI: https://doi.org/10.1007/978-94-011-5814-5_4
- Thothathiri, M., & Snedeker, J. (2008). Syntactic priming during language comprehension in three- and four-year-old children. *Journal of Memory and Language*, 58(2), 188–213. DOI: <https://doi.org/10.1016/j.jml.2007.06.012>
- Tooley, K. M. (2022). Structural priming during comprehension: A pattern from many pieces. *Psychonomic Bulletin & Review*, 1–15. DOI: <https://doi.org/10.3758/s13423-022-02209-7>
- Tooley, K. M., & Traxler, M. J. (2010). Syntactic priming effects in comprehension: A critical review. *Language and Linguistics Compass*, 4(10), 925–937. DOI: <https://doi.org/10.1111/j.1749-818X.2010.00249.x>
- Urbach, T. P., DeLong, K. A., & Kutas, M. (2015). Quantifiers are incrementally interpreted in context, more than less. *Journal of Memory and Language*, 83, 79–96. DOI: <https://doi.org/10.1016/j.jml.2015.03.010>
- Urbach, T. P., & Kutas, M. (2010). Quantifiers more or less quantify on-line: ERP evidence for partial incremental interpretation. *Journal of Memory and Language*, 63(2), 158–179. DOI: <https://doi.org/10.1016/j.jml.2010.03.008>

- Van Gompel, R. P., & Arai, M. (2018). Structural priming in bilinguals. *Bilingualism: Language and Cognition*, 21(3), 448–455. DOI: <https://doi.org/10.1017/S1366728917000542>
- Yildirim, I., Degen, J., Tanenhaus, M. K., & Jaeger, T. F. (2016). Talker-specificity and adaptation in quantifier interpretation. *Journal of Memory and Language*, 87, 128–143. DOI: <https://doi.org/10.1016/j.jml.2015.08.003>
- Zehr, J., & Schwarz, F. (2018). *PennController for Internet Based Experiments (IBEX)*. DOI: <https://doi.org/10.17605/OSF.IO/MD832>
- Ziegler, J., Bencini, G., Goldberg, A., & Snedeker, J. (2019). How abstract is syntax? Evidence from structural priming. *Cognition*, 193, 104045. DOI: <https://doi.org/10.1016/j.cognition.2019.104045>
- Ziegler, J., & Snedeker, J. (2018). How broad are thematic roles? Evidence from structural priming. *Cognition*, 179, 221–240. DOI: <https://doi.org/10.1016/j.cognition.2018.06.019>
- Ziegler, J., & Snedeker, J. (2019). The use of syntax and information structure during language comprehension: Evidence from structural priming. *Language, Cognition and Neuroscience*, 34(3), 365–384. DOI: <https://doi.org/10.1080/23273798.2018.1539757>
- Ziegler, J., Snedeker, J., & Wittenberg, E. (2018). Event structures drive semantic structural priming, not thematic roles: Evidence from idioms and light verbs. *Cognitive Science*, 42(8), 2918–2949. DOI: <https://doi.org/10.1111/cogs.12687>

