



Universiteit
Leiden
The Netherlands

Using a language model to unravel semantic development in children's use of a Dutch perception verb

Dijk, B.M.A. van; Duijn, M.J. van; Kloostra, L.; Spruit, M.R.; Beekhuizen, B.; Zock, M.; ... ; De Deyne, S.

Citation

Dijk, B. M. A. van, Duijn, M. J. van, Kloostra, L., Spruit, M. R., & Beekhuizen, B. (2024). Using a language model to unravel semantic development in children's use of a Dutch perception verb. *Proceedings Of The Workshop On Cognitive Aspects Of The Lexicon (Cogalex@lrec-Coling 2024)*, 98–106. Retrieved from <https://hdl.handle.net/1887/3761710>

Version: Publisher's Version

License: [Creative Commons CC BY-NC 4.0 license](https://creativecommons.org/licenses/by-nc/4.0/)

Downloaded from: <https://hdl.handle.net/1887/3761710>

Note: To cite this publication please use the final published version (if applicable).

Using a Language Model to Unravel Semantic Development in Children's Use of a Dutch Perception Verb

Bram van Dijk¹, Max van Duijn¹, Li Kloostra², Marco Spruit¹, and Barend Beekhuizen³

¹Leiden Institute of Advanced Computer Science, ²Utrecht University, ³University of Toronto

Corresponding author: b.m.a.van.dijk@liacs.leidenuniv.nl

Abstract

We employ a Language Model (LM) to gain insight into how complex semantics of Dutch Perception Verb (PV) *zien* ('to see') emerge in children. Using a Dutch LM as representation of mature language use, we find that for ages 4-12y 1) the LM accurately predicts PV use in children's freely-told narratives; 2) children's PV use is close to mature use; and 3) complex PV meanings with attentional and cognitive aspects can be found. Our approach illustrates how LMs can be meaningfully employed in studying language development, hence takes a constructive position in the debate on the relevance of LMs in this context.

Keywords: language development, language models, computational modelling, semantics, pragmatics

1. Introduction

Recent Language Models (LMs) based on Transformer architectures (Vaswani et al., 2017) reflect semantic knowledge present in a language community. BERT vectors (Devlin et al., 2019), for example, are able to distinguish different senses of the same word (Rogers et al., 2020; Vulić et al., 2020; Wiedemann et al., 2019). These LMs implement the distributional hypothesis that words with similar meanings tend to occur in similar contexts, and they represent both word type and word token meanings with real-valued vectors (Lenci and Sahlgren, 2023). The latter allows LMs to encode polysemy and different usages of words.

Despite this, LMs' relevance in the context of language development is disputed: their architecture and volume of training input have been argued to make them incomparable to children (e.g. Bunzeck and Zariw, 2023; Prystawski et al., 2022; Warstadt and Bowman, 2022). Yet, others argue that LMs can show which linguistic phenomena are *in principle* learnable from distributional information, bearing on learnability debates (Contreras Kallens et al., 2023; Piantadosi, 2023; Wilcox et al., 2023).

Here we leverage LMs' rich semantic information to gain insight in children's semantic and pragmatic development. Addressing the question whether children's pragmatic use of lexical items develops over time or, conversely, is adult-like from the start, we use a Dutch LM as representation of mature language use and study the Dutch Perception Verb (PV) *zien* ('to see'). We find that children's use of *see* is close to mature use across the 4-12y age range, and that for all ages the familiar mature usage patterns of the verb can be identified. As such, the paper further illustrates the relevance of LMs in studying language development, by reflecting on LMs as representations of mature language use and setting up appropriate tasks and metrics.

2. Background

Little empirical work employs modern LMs in language development, the exception being work comparing word acquisition in children and LMs (Chang and Bergen, 2022; Laverghetta and Licato, 2021). This is understandable given the debate on the validity of LMs in the child context: LMs and children differ in key respects including word exposure (Warstadt and Bowman, 2022) and learning mechanisms (Bunzeck and Zariw, 2023).

Still, LMs are arguably useful representations of *mature language use* by being trained on corpora of adult language, and are therefore of value in modelling language understanding. LMs can be viewed as an incremental methodological step compared to earlier corpus studies comparing children's verb use to mature use, that relied on manual annotation or feature engineering to identify different senses of mature verb use (e.g. Adricula and Narasimhan, 2009; Parisien and Stevenson, 2009), but different senses, as we will show, can also be conveniently retrieved from LLMs. These and other considerations have led to increasing acknowledgement of LMs' relevance for analysing language development (Contreras Kallens et al., 2023; Lappin, 2023), and efforts to make LMs more comparable to the child context (Warstadt et al., 2023).

Here we address the relevance of LMs in the developmental context by analysing children's lexical semantic development with LMs. We target children's use of Dutch PV *zien* ('to see') as a case study, which has been frequently analysed in language development (e.g. Davis, 2020; Davis and Landau, 2021). Studies of perception verbs across languages have shown that visual perception verbs have extended meanings beyond their denotational meaning 'entity X visually perceives object or event Y', that involve additional aspects of e.g. *attention* ('Let's see if I can find the keys') and *cognition* ('I

see what you mean') (San Roque et al., 2018; San Roque and Schieffelin, 2019). Such meaning extensions are salient for children with a limited lexicon, where meaning extension of known words allows children to express new meanings efficiently (Nerlich and Clarke, 1999). In addition, since visual perception is argued to have strong metaphorical mappings to knowledge and understanding (e.g. Johnson, 1999), *see* can be a window onto how children learn to represent (socio-)cognitive content with language (Sweetser, 1990).

This work addresses the question when meaning extension occurs. Some argue that literal understandings of PVs emerge first in young children (e.g. Davis, 2020; Davis and Landau, 2021; Elli et al., 2021; Landau and Gleitman, 2009), while others argue pragmatic meanings are likely present early due to the social situatedness of language learning (e.g. Enfield, 2023; San Roque and Schieffelin, 2019). In the latter case, the discursive relation between the visual perception event and the events surrounding it may be more salient for a language learner than the encoding of visual perception per se. For example, a young child's utterance *see ball* may be followed by the caregiver showing the ball, or focusing its attention on the ball — further attentional aspects that are likely relevant components of the message for the child beyond the denotational content of visual perception having taken place. While focusing on a single verb may seem limited, we believe as a case study, visual perception verbs are well-chosen as a starting point for generalising the proposed approach, since their acquisitional pathway and pragmatic usages (as described above) are well understood.

We focus on children's use of *see* in ChiSCor, a corpus of freely-told stories by Dutch children (4-12y) in classroom settings (van Dijk et al., 2023b), since complex PV meanings can be especially relevant in the narrative domain. For example, that character X *sees* entity Y may not only imply that X literally perceives Y, but also that X *evaluates* Y or *discovers* Y. Such information, which may be crucial for the 'tellability' of the story (Labov and Waletzky, 1967), can be efficiently transmitted through PVs. Narratives are 'natural' sandboxes for children to challenge their language competence in various ways (Frizelle et al., 2018), including the development of lexical pragmatics.

3. Methods

Language data – We extracted all 308 occurrences of *see* from 619 stories of 442 children (4-12y) in ChiSCor. We manually inspected these occurrences and removed unintelligible usages (mainly transcription errors) as well as stories exceeding a context window larger than 512 tokens, resulting in

210 occurrences. We assigned occurrences to a Young (4-6), Middle (6-9) or Old (9-12) age group, following the age binning in Dutch primary education, and included only PV occurrences from one story per child, resulting in 30 Young, 82 Middle and 42 Old PV occurrences. To balance the sample across age groups, we randomly sampled 30 occurrences from the Middle and Old age group.

A known problem with LMs is that data contamination can lead them to solve tasks by memorisation (Deng et al., 2023). ChiSCor is likely not in the train data of recent LMs, as the corpus is recent and 'hidden' behind view-only links in research papers. Further, ChiSCor's free storytelling is unlike other available Dutch corpora that involve language elicitation and as such constitutes language that tests LMs' generalisation capabilities.

LMs as benchmark models – Using LMs as representation of mature language use requires evidence that the LM models the linguistic phenomenon and domain at issue reliably. We draw on findings that word representations in BERT encode rich semantic information about word polysemy (Garí Soler and Apidianaki, 2021; Wiedemann et al., 2019), although not perfectly. Also, Dutch LMs are for a large part trained on narrative texts (e.g. De Vries et al., 2019; Delobelle et al., 2020), and LMs in general have been shown to model coherence in written narratives well (Laban et al., 2021). In sum, earlier work supports the idea that LMs encode mature PV use in narratives.

Choice of LMs – For reasons of computational efficiency, validity with respect to the child context, and reproducibility, we chose RobBERT-2023-dutch-large, a Dutch BERT-like LM (Delobelle et al., 2020). RobBERT has 455M parameters trained on 19.5B tokens and is more in line with the 100M training input a 10-year-old has seen (Warstadt and Bowman, 2022), compared to often employed large LMs like GPT-3 (175B parameters, 500B tokens (Brown et al., 2020)). RobBERT is accessible through the HuggingFace Transformers ecosystem (Wolf et al., 2019).

Recent work on making LMs relevant to human language acquisition in the BabyLM challenge (Warstadt et al., 2023), highlighted smaller LMs with optimised architectures and train objectives, and curated train data for training developmentally plausible models (Samuel et al., 2023). However, such Dutch LMs are not yet available and training models from scratch is generally not feasible for researchers studying language acquisition. RobBERT was a fitting resource as it is optimised compared to BERT and has a simpler training objective (masked language modelling only) (Liu et al., 2019). These aspects go some way towards the findings of the BabyLM challenge (Samuel et al., 2023; Warstadt et al., 2023).

Task design and metrics – To use LMs as representations of mature language use, zero-shot evaluation settings as described by Laban et al. (2021) are preferred. This means using LMs of-the-shelf without further pre-training on the target domain or fine-tuning to stay close to the mature language use encoded in the LM, similar to how factual knowledge can be retrieved from LMs without fine-tuning (Petroni et al., 2019).

We use various possibilities available through LMs to assess whether and how children’s use of *see* differs from mature use.

Our first task consists of predicting *see* in children’s narratives. We present RobBERT with stories containing a masked instance of *see*, as in the (translated) excerpt in (1):

- (1) [...] one time robot was travelling. and all of a sudden he <mask> a wolf. and he ran away quickly. [...] (Story ID 052301)

In our experiment we provided full stories as context to RobBERT, which varied in number of words ($\bar{x} = 187, \sigma = 108$). If children’s usage differs from adults, the LM might have difficulty predicting the PV correctly.

As a second measure, we compute the negative log-likelihood *NLL* or surprisal for a prediction for a masked token w_m with $NLL(w_m) = -\log p(w_m | w_{1...m-1}, w_{m+1...n})$ with the `fill-mask` pipeline from HuggingFace Transformers. This measure provides further context to the predictive accuracy measure presented above: lower *NLL* implies that the predicted token is less surprising i.e. closer to mature use as encoded in the LM, and more generally indicates how well a given context supports a specific token on the masked position (PV or other).

Lastly, we use the tokens in RobBERT’s top-5 predictions for masked instances of *see* as ‘near neighbours’ that can reveal the additional discursive meanings that the usage of PVs support. Our data and notebooks are available at <https://shorturl.at/jquVX>.

4. Results

Predictive accuracy – First, we assessed RobBERT’s overall performance in predicting *see* at masked positions in all 90 PV occurrences. Accuracy is overall high (.83, Table 1), and although lower for Young (.70) we found no significant difference in accuracy between ages with an ANOVA ($F_{2,87} = 2.974, p = .056$). This shows that RobBERT models children’s PV use in the narrative domain well. The 15 errors were mainly in Young and showed confusion of *seeing* with ‘finding’, ‘having’, ‘looking’ and ‘getting’, meaning that contexts under-constrained the use of *see*. Although these other

verbs can be valid tokens on masked positions (e.g. ‘found’ in (1)), here our aim was to see if RobBERT adequately models that *see* can subsume such other possible meanings in narratives.

Metric	Young	Middle	Old	Overall
Accuracy	.70 (30)	.90 (30)	.90 (30)	.83 (90)
Surprisal	.40 (21)	.23 (27)	.32 (27)	.31 (75)
Top-5	1.00 (30)	1.00 (30)	.97 (30)	.99 (90)

Table 1: Metrics for RobBERT. Accuracy: percentage that *see* was predicted. Surprisal: *NLL* computed for predictions of *see*. Top-5: proportion that *see* was in top-5 predictions. Number of PV occurrences (i.e. observations) in parentheses.

Surprisal – Second, we analysed potential age effects in mean surprisal for 75 correct predictions of *see*. For example, RobBERT may be less surprised by PV use for Old compared to Young or Middle, indicating PV use of Old children is closer to mature use than Young. Interestingly, surprisal distributions are close to 0 for all ages (Figure 1), and although mean surprisal between Young, Middle, and Old differs (Table 1), pairwise comparisons with Tukey’s HSD (Tukey, 1949) revealed no significant age effects. This suggests that PV use by children of all ages is about equally close to mature use.

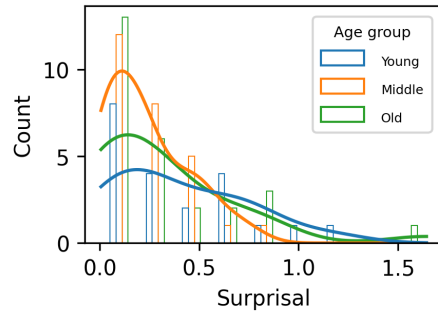


Figure 1: Surprisal distributions.

Top-5 alternative predictions – For virtually all age groups, *see* is in the top-5 predictions (Table 1), which supports the idea that by examining top-5s we get insight in extended meanings of *see*. For 90 PV occurrences and their top-5s (450 tokens) we lemmatised tokens and removed *see* and lemmas that were not verbs (e.g. ‘many’, ‘and’, ‘at’), resulting in 304 lemmas. We then took the set and classified 65 lemmas as having roughly ‘external’, ‘internal’, or ‘other’ meaning. *External* implies a meaning pertaining to plain action (e.g. ‘to go’, ‘to come’, ‘to carry’, ‘to throw’); *internal* a meaning pertaining to an attentional (e.g. ‘to notice’, ‘to meet’) or cognitive state (e.g. ‘to think’, ‘to know’). *Other* pertains to auxiliary verbs and PVs not the focus of the current study (e.g. ‘to have’, ‘to hear’). The

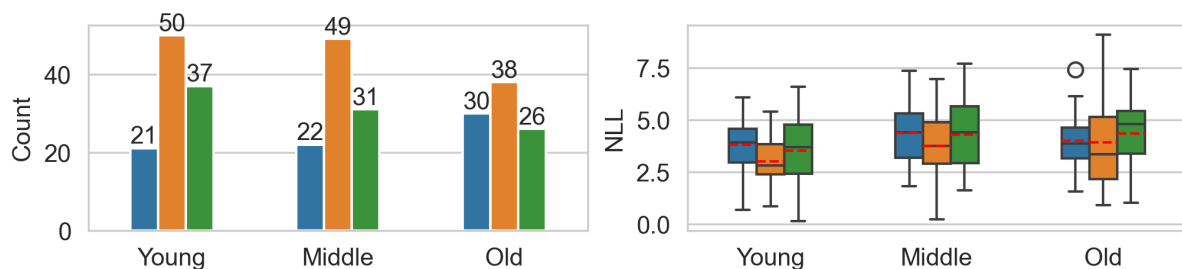


Figure 2: Frequencies (left) and surprisal dist. (right) of internal (blue), external (orange) and other (green) meanings of 304 top-5 lemmas. Bars (left) stack to 100%; dashed red lines (right) indicate means.

Age	Ex.	PV context
Young	(2)	.. and when he returned. then he saw/ <u>knew</u> that the princess was gone. and they lived happily ever after. (102901)
	(3)	.. and then they were lost again. and then they saw/ <u>saw</u> the castle. and then they went in the castle. (122901)
	(4)	.. but then the teacher came and then she was already too late. the teacher had seen/ <u>caught</u> them. and then you get a punishment from the teacher. (033401)
Middle	(5)	.. but then they lost each other all of a sudden. and then Wergje saw/ <u>met</u> another rabbit. and it asked how are you called. (072301)
	(6&7)	.. because when he was home. then he saw/ <u>noticed/discovered</u> that he had the other scales. but then he went to fly on it and he wanted to find his own dragon again. (022301)
Old	(8)	.. once arrived at the cave Puta completely forgot that you were not allowed to touch the big diamond. Puta saw/ <u>checked out</u> the diamond and found it so beautiful. and he touched it accidentally. (034801)
	(9)	.. so then the fat little king went on his fat broom to the cry for help. and what did he see/ <u>think</u> . the cry came from a little fat guinea pig that looked very much like the king. (023801)
	(10)	.. and he ever wanted one time to try it with his eyes closed. to see/ <u>test</u> can I grab that donut well with my eyes closed. (034501)

Table 2: Translated PV contexts with top-5 internal lemmas (underlined) with lowest surprisal. Story IDs given in parentheses. All excerpts were translated by the first author.

idea is that top-5 lemmas indicate what possible meanings PV contexts support, even if these lemmas are not necessarily intuitive substitutions. For example, substituting ‘threw’ for <mask> in (1) renders the excerpt less intuitive. Yet, this immediate context as a sequence of *external* actions better supports understanding *seeing* also as a causal part of a sequence of external actions, than as *seeing* as part of narrative components reflecting a character’s attentional or cognitive *internal* states (cf. examples in Table 2).

We assessed frequencies of external, internal and other meanings, and their mean surprisal over age groups to identify potential age differences in occurrence and closeness to mature use. Regarding frequency, although external and other meanings decrease over age while internal meanings increase over age (Figure 2, left), we found no significant age effects with a χ^2 test $\chi^2(4, N = 304) = 5.044, p = .283$, suggesting that all the different meanings are about equally frequent in Young, Middle and Old groups. Regarding surprisal (Figure 2, right), distributions for external, internal and other meanings are relatively similar both within and between age groups. Pairwise

comparisons with Tukey’s HSD found only a significant difference at the $p < .05$ level between mean surprisal for external meanings for Young and Old.

We illustrate complex meanings of *see* present in all age groups, by providing the three internal meanings that were closest to mature use (i.e. with lowest surprisal) and their PV contexts in Table 2. We make three observations. First, internal meanings with attentional and cognitive aspects can be but are not exclusively cued by surface linguistic frames such as complementation that RobBERT simply picks up, as example (4) and (9) show. In (4) ‘caught’ implies that the teacher knows what the ‘she’ character is up to; in (9) ‘think’ renders the realisation where the cry of help is coming from a representation in the mind of the king. Second, internal meanings are varied: from more purely attentional where characters simply become aware of something or find something out as in (6&7), to more social (5), and evaluative attentional aspects (8). Third, although internal meanings with cognitive aspects have the most abstract lemmas (‘think’, ‘know’) that are argued to be harder to master (Barak et al., 2012), cognitive meanings were found in both Young (2), (4) and Old (9) children.

5. Discussion

Our results show that complex meanings of the Dutch perception verb *zien* ('to see') are about equally frequent in all age groups and that children's use of the PV is overall not significantly different from mature use. This contrasts with earlier work that has argued that children initially acquire more literal meanings of PVs (Adricula and Narasimhan, 2009; Davis, 2020; Davis and Landau, 2021; Elli et al., 2021; Landau and Gleitman, 2009) (Section 2), although we note that children in our sample are older (four years and older) than children in earlier studies (typically between two and four years).

Our result aligns with the idea that it is the social context that cues various complex senses of *see* in children (e.g. Enfield, 2023; San Roque and Schieffelin, 2019), and with the idea that (young) children may employ PVs like *see* as linguistic devices for learning to represent cognitive and attentional states (Johnson, 1999; Sweetser, 1990). We argue that our finding can be explained by the social context provided by live storytelling. PVs like *see* are linguistic devices for efficiently communicating about characters' attentional and cognitive states that are key to understanding the story, as PVs can compress redundant information that would make the story tedious. Earlier work has shown that, in children's live storytelling, contexts of PV like *hear* and *see* are coherent and clear, as evidenced by the rich PV vectors that can be trained from limited amounts of narrative data (van Dijk et al., 2023b).

Narrative language data may explain the contrast between our and earlier findings as storytelling has been argued to solicit 'maximal behaviour' in that it challenges children's linguistic competence (Frizelle et al., 2018; Southwood and Russell, 2004), more than the speech produced by children in child-caregiver interactions would do, which typically take place in mundane contexts. Some earlier work contrasting with our results relied on language data from such child-caregiver interactions (e.g. Davis and Landau, 2021; Adricula and Narasimhan, 2009). The latter work also employed smaller sample sizes with less unique children and more PV use per child compared to the current study, which compresses the variation in complex semantics we find in our analysis.

Interestingly, RobBERT accurately predicted *see* in narratives of children of all ages; we argue that this is not a mere frequency effect (i.e. *see* being more frequent in train data than alternatives), given that top-5 predictions often reveal RobBERT's correct mapping of the nuanced senses of PVs. Also, RobBERT's aptitude in handling PV use in narratives is interesting insofar children's stories are not obvious regarding wording, characters and themes.

One issue pointed out by a reviewer is whether LMs with Transformer architectures are the best fit for representing linguistic knowledge of a mature Dutch language user, or whether other models should be used, e.g. from the BabyLM challenge (Warstadt et al., 2023). The best-performing LMs in this challenge employed Transformer architectures that are essentially optimised versions of vanilla BERT models regarding training objective, architecture and dataset (Samuel et al., 2023). With our choice for RobBERT we aimed to make the comparison to the human case as valid as possible with an existing resource (see Section 3).

In any case, from the BabyLM challenge we learn that the Transformer architecture is also in more modest training setups a powerful encoder of linguistic information. Our claim is not that Transformers are therefore good (cognitive) models of human language users, which is still debated (e.g. Paape, 2023; van Dijk et al., 2023a). Rather, when it comes to specific linguistic aspects such as mature semantic and pragmatic knowledge, LMs as sophisticated distributional learners represent this information in a convenient fashion. For using such computational models as representations of mature language use, the primary question is if their *behaviour* for a specific linguistic phenomenon is sufficiently complex, which for many modern BERT-like models seems the case. But representations of mature use could also be created in other ways, e.g. by clustering different verb senses with features based on verb argument structure in a large corpus of mature language use. Thus, LMs are more of an analytical tool here than direct models of humans. That said, it is still worthwhile and necessary to make LMs more similar to the human context.

6. Conclusion

This paper provided a case study on Dutch children's (4-12y) use of *zien* ('to see') and the emergence of complex semantics in this perception verb. We showed that 1) a recent Dutch LM can predict use of *see* in narratives for different ages reliably; 2) children's use of *see* is close to mature use for all ages; and 3) complex meanings of *see* with attentional and cognitive aspects can be found across all ages. Our results align with work that argues that meaning extension occurs early in children and with the idea that via perception verbs, children may learn to represent socio-cognitive content.

We also showed how LMs can be meaningfully leveraged in developmental contexts. We hope to provide future researchers with useful reflection on how to proceed when using LMs as representations of mature language use, choosing models, and setting up tasks and metrics.

7. Limitations and Ethical Considerations

A limitation of this study is that we provided the whole story as context for predicting a masked occurrence of *zien* ('to see'), but for space limitations only could discuss complex meanings with smaller story excerpts as in Table 2. This may suggest that complex PV meanings can be determined from small pieces of narrative after all. Yet, when doing the same task with smaller PV contexts as in Table 2, i.e. a sentence before and after the sentence featuring an occurrence of *see*, RobBERT's overall accuracy drops from .83 to .57 and overall surprisal increases from .31 to .59, (see Table 1) which suggests that RobBERT needs to take the whole story into account to model PV use adequately. This means that there is more relevant information in the context beyond what we show in the immediate PV context that render RobBERT's predictions of masked tokens accurate and support additional meanings of *see*.

Another limitation is that we had to translate story excerpts to English, as also providing Dutch excerpts required too much space. Some awkwardness in translations could not be avoided. For example, Dutch has a verb 'betrappen' that always has a cognitive meaning similar to 'catching somebody red-handed', whereas 'catching' in English can also have a more obvious action-related meaning. 'Betrappen' was a token prediction in RobBERT's top-5 with low surprisal that we had to translate as 'caught' in example (4) in Table 2.

In this study we used the ChiSCor story corpus and we refer to van Dijk et al. (2023b) for further details regarding ethical considerations and approval that was obtained for collecting language data from children. Regarding computational efficiency, we chose a relatively small, open and free to use language model that can also be employed with limited computational resources.

8. Acknowledgements

This research was done in the context of Max van Duijn's research project 'A Telling Story' (with project number VI.Veni.191C.051), which is financed by the Dutch Research Council (NWO). The collaboration with Barend Beekhuizen from the University of Toronto was made possible by funding from The Leiden University/Swaantje Mondt Fund (reference number W232234-1-097).

9. Bibliographical References

- Norielle Adricula and Bhuvana Narasimhan. 2009. 'understanding is understanding by seeing': Visual perception verbs in child language. In *Proceedings of the 44th Boston University Conference on Language Development*, pages 18–27, Boston. Somerville, MA: Cascadilla Press.
- Libby Barak, Afsaneh Fazly, and Suzanne Stevenson. 2012. *Modeling the Acquisition of Mental State Verbs*. In *Proceedings of the 3rd Workshop on Cognitive Modeling and Computational Linguistics (CMCL 2012)*, pages 1–10, Montréal, Canada. Association for Computational Linguistics.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Bastian Bunzeck and Sina Zarriß. 2023. *GPT-wee: How Small Can a Small Language Model Really Get?* In *Proceedings of the BabyLM Challenge at the 27th Conference on Computational Natural Language Learning*, pages 35–46, Singapore. Association for Computational Linguistics.
- Tyler A. Chang and Benjamin K. Bergen. 2022. *Word Acquisition in Neural Language Models*. *Transactions of the Association for Computational Linguistics*, 10:1–16.
- Pablo Contreras Kallens, Ross Deans Kristensen-McLachlan, and Morten H. Christiansen. 2023. Large language models demonstrate the potential of statistical learning in language. *Cognitive Science*, 47(3):e13256.
- E. Emory Davis. 2020. *Does seeing mean believing? The development of children's semantic representations for perception verbs*. Ph.D. thesis, The Johns Hopkins University.
- E. Emory Davis and Barbara Landau. 2021. Seeing and believing: the relationship between perception and mental verbs in acquisition. *Language Learning and Development*, 17(1):26–47.
- Wietse De Vries, Andreas van Cranenburgh, Arianna Bisazza, Tommaso Caselli, Gertjan van Noord, and Malvina Nissim. 2019. BERTje: A Dutch BERT model. *arXiv preprint arXiv:1912.09582*.
- Pieter Delobelle and François Remy. 2023. *RobBERT-2023: Keeping Dutch Language Models Up-To-Date at a Lower Cost Thanks to Model Conversion*. *The 33rd Meeting of Computational Linguistics in The Netherlands (CLIN 33)*.
- Pieter Delobelle, Thomas Winters, and Bettina Berendt. 2020. *RobBERT: a Dutch RoBERTa-based Language Model*. In *Findings of the Association for Computational Linguistics: EMNLP 2020*,

- pages 3255–3265, Online. Association for Computational Linguistics.
- Chunyuan Deng, Yilun Zhao, Xiangru Tang, Mark Gerstein, and Arman Cohan. 2023. Investigating data contamination in modern benchmarks for large language models. *arXiv preprint arXiv:2311.09783*.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Giulia V. Elli, Marina Bedny, and Barbara Landau. 2021. How does a blind person see? developmental change in applying visual verbs to agents with disabilities. *Cognition*, 212:104683.
- Nick J. Enfield. 2023. Linguistic concepts are self-generating choice architectures. *Philosophical Transactions of the Royal Society B*, 378(1870):20210352.
- Rita E. Frank and William S. Hall. 1991. Polysemy and the acquisition of the cognitive internal state lexicon. *Journal of Psycholinguistic Research*, 20(4):283–304.
- Pauline Frizelle, Paul A. Thompson, David McDonald, and Dorothy V.M. Bishop. 2018. Growth in syntactic complexity between four years and adulthood: Evidence from a narrative task. *Journal of Child Language*, 45(5):1174–1197.
- Aina Garí Soler and Marianna Apidianaki. 2021. Let’s play mono-poly: BERT can reveal words’ polysemy level and partitionability into senses. *Transactions of the Association for Computational Linguistics*, 9:825–844.
- Christopher Johnson. 1999. Metaphor vs. Conflation in the Acquisition of Polysemy: The Case of SEE. In Masako K. Hiraga, Sherman Wilcox, and Chris Sinha, editors, *Cultural, Psychological and Typological Issues in Cognitive Linguistics: Selected papers of the bi-annual ICLA meeting in Albuquerque*, pages 155–169.
- Philippe Laban, Luke Dai, Lucas Bandarkar, and Marti A. Hearst. 2021. [Can Transformer Models Measure Coherence In Text: Re-Thinking the Shuffle Test](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 2: Short Papers)*, pages 1058–1064, Online. Association for Computational Linguistics.
- William Labov and Joshua Waletzky. 1967. Narrative analysis: Oral versions of personal experience. In J Holm, editor, *Essays on the Verbal and Visual Arts*, pages 12–44. University of Washington Press.
- Barbara Landau and Lila R. Gleitman. 2009. *Language and experience: Evidence from the blind child*. Harvard University Press.
- Shalom Lappin. 2023. Assessing the Strengths and Weaknesses of Large Language Models. *Journal of Logic, Language and Information*, pages 1–12.
- Antonio Laverghetta and John Licato. 2021. Modeling age of acquisition norms using transformer networks. In *The International FLAIRS Conference Proceedings*, volume 34.
- Alessandro Lenci and Magnus Sahlgren. 2023. *Distributional semantics*. Cambridge University Press.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. Roberta: A roustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692*.
- Christopher D. Manning, Kevin Clark, John Hewitt, Urvashi Khandelwal, and Omer Levy. 2020. Emergent linguistic structure in artificial neural networks trained by self-supervision. *Proceedings of the National Academy of Sciences*, 117(48):30046–30054.
- Brigitte Nerlich and David D. Clarke. 1999. Elements for an integral theory of semantic change and semantic development. In *Meaning Change—Meaning Variation. Workshop held at Konstanz*, volume 1, pages 123–134.
- Dario Paape. 2023. When Transformer models are more compositional than humans: The case of the depth charge illusion. *Experiments in Linguistic Meaning*, 2:202–218.
- Christopher Parisien and Suzanne Stevenson. 2009. Modelling the acquisition of verb polysemy in children. In *Proceedings of the CogSci2009 Workshop on Distributional Semantics beyond Concrete Concepts*, pages 17–22, Austin, Texas. Cognitive Science Society.
- Fabio Petroni, Tim Rocktäschel, Sebastian Riedel, Patrick Lewis, Anton Bakhtin, Yuxiang Wu, and Alexander Miller. 2019. [Language Models as Knowledge Bases?](#) In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 2463–2473, Hong Kong, China. Association for Computational Linguistics.

- Steven Piantadosi. 2023. Modern language models refute Chomsky’s approach to language. *Lingbuzz Preprint*, lingbuzz, 7180.
- Steven T. Piantadosi and Felix Hill. 2022. Meaning without reference in large language models. *arXiv preprint arXiv:2208.02957*.
- Ben Prystawski, Erin Grant, Aida Nematzadeh, Spike W.S. Lee, Suzanne Stevenson, and Yang Xu. 2022. The emergence of gender associations in child language development. *Cognitive Science*, 46(6):e13146.
- Anna Rogers, Olga Kovaleva, and Anna Rumshisky. 2020. [A Primer in BERTology: What We Know About How BERT Works](#). *Transactions of the Association for Computational Linguistics*, 8:842–866.
- Magnus Sahlgren and Fredrik Carlsson. 2021. [The Singleton Fallacy: Why Current Critiques of Language Models Miss the Point](#). *Frontiers in Artificial Intelligence*, 4:682578.
- David Samuel, Andrey Kutuzov, Lilja Øvrelid, and Erik Velldal. 2023. [Trained on 100 million words and still in shape: BERT meets British National Corpus](#). In *Findings of the Association for Computational Linguistics: EACL 2023*, pages 1954–1974, Dubrovnik, Croatia. Association for Computational Linguistics.
- Lila San Roque, Kobin H. Kendrick, Elisabeth Norcliffe, and Asifa Majid. 2018. Universal meaning extensions of perception verbs are grounded in interaction. *Cognitive Linguistics*, 29(3):371–406.
- Lila San Roque and Bambi B Schieffelin. 2019. Perception verbs in context: Perspectives from Kaluli (Bosavi) child-caregiver interaction. *Laura Speed, C. O’Meara, Lila San Roque, and Asifa Majid (eds.) Perception Metaphors*, pages 347–368.
- Frenette Southwood and Ann F. Russell. 2004. [Comparison of Conversation, Freeplay, and Story Generation as Methods of Language Sample Elicitation](#). *Journal of Speech, Language and Hearing Research*, 47(2):366–376.
- Eve Sweetser. 1990. *From etymology to pragmatics: Metaphorical and cultural aspects of semantic structure*, volume 54. Cambridge University Press.
- John W. Tukey. 1949. Comparing individual means in the analysis of variance. *Biometrics*, pages 99–114.
- Bram van Dijk, Tom Kouwenhoven, Marco Spruit, and Max Johannes van Duijn. 2023a. [Large Language Models: The Need for Nuance in Current Debates and a Pragmatic Perspective on Understanding](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 12641–12654, Singapore. Association for Computational Linguistics.
- Bram van Dijk, Max van Duijn, Suzan Verberne, and Marco Spruit. 2023b. [ChiSCor: A Corpus of Freely-Told Fantasy Stories by Dutch Children for Computational Linguistics and Cognitive Science](#). In *Proceedings of the 27th Conference on Computational Natural Language Learning (CoNLL)*, pages 352–363, Singapore. Association for Computational Linguistics.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. *Advances in neural information processing systems*, 30.
- Ivan Vulić, Edoardo Maria Ponti, Robert Litschko, Goran Glavaš, and Anna Korhonen. 2020. [Probing Pretrained Language Models for Lexical Semantics](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 7222–7240, Online. Association for Computational Linguistics.
- Alex Warstadt and Samuel R. Bowman. 2022. What artificial neural networks can tell us about human language acquisition. In *Algebraic Structures in Natural Language*, pages 17–60. CRC Press.
- Alex Warstadt, Aaron Mueller, Leshem Choshen, Ethan Wilcox, Chengxu Zhuang, Juan Ciro, Rafael Mosquera, Bhargavi Paranjabe, Adina Williams, Tal Linzen, and Ryan Cotterell. 2023. [Findings of the BabyLM Challenge: Sample-Efficient Pretraining on Developmentally Plausible Corpora](#). In *Proceedings of the BabyLM Challenge at the 27th Conference on Computational Natural Language Learning*, pages 1–34, Singapore. Association for Computational Linguistics.
- Gregor Wiedemann, Steffen Remus, Avi Chawla, and Chris Biemann. 2019. Does BERT make any sense? Interpretable word sense disambiguation with contextualized embeddings. *arXiv preprint arXiv:1909.10430*.
- Ethan Gotlieb Wilcox, Richard Futrell, and Roger Levy. 2023. Using computational models to test syntactic learnability. *Linguistic Inquiry*, pages 1–44.
- Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, et al. 2019. Huggingface’s transformers: State-of-the-art natural language processing. *arXiv preprint arXiv:1910.03771*.

10. Language Resource References

van Dijk, Bram and van Duijn, Max and Verberne, Suzan and Spruit, Marco. 2023. *ChiSCor: A Corpus of Freely-Told Fantasy Stories by Dutch Children for Computational Linguistics and Cognitive Science*. Open Science Framework. PID https://osf.io/5h7za/?view_only=705c553e922046058ef9df52b4ac8ed7.