

How Important Is Language for Human-Like Intelligence?

Perspectives on Psychological Science
1–6

© The Author(s) 2026

Article reuse guidelines:

sagepub.com/journals-permissions

DOI: 10.1177/17456916251398539

www.psychologicalscience.org/PPS



Gary Lupyan¹, Hunter Gentry², and Martin Zettersten³

¹Department of Psychology, University of Wisconsin–Madison; ²Department of Philosophy, Cognitive Science Program, University of Central Florida; and ³Department of Cognitive Science, University of California, San Diego

Abstract

We use language to communicate our thoughts. But is language merely the expression of thoughts, which are themselves produced by other, nonlinguistic parts of our minds? Or does language play a more transformative role in human cognition, allowing us to have thoughts that we otherwise could (or would) not have? Recent developments in artificial intelligence (AI) and cognitive science have reinvigorated this old question. We argue that language may hold the key to the emergence of both more general AI systems and central aspects of human intelligence. We highlight two related properties of language that make it such a powerful tool for developing domain-general abilities. First, language offers compact representations that make it easier to represent and reason about many abstract concepts (e.g., exact numerosity). Second, these compressed representations are the iterated output of collective minds. In learning a language, we learn a treasure trove of culturally evolved abstractions. Taken together, these properties mean that a sufficiently powerful learning system exposed to language—whether biological or artificial—learns a compressed model of the world, reverse engineering many of the conceptual and causal structures that support human (and human-like) thought.

Keywords

LLMs, language and thought, artificial intelligence, learning from language

Language is clearly an important source of knowledge about the world. It is largely because of language that our knowledge far exceeds our personal experiences. Language allows us to learn about past events, helps to plan for the future, and is indispensable for creating the stories that make up human culture. Language is also at the center of much of our formal education. Given its seeming importance, it is therefore surprising that the role of linguistic input has figured minimally in accounts of how humans acquire and structure semantic knowledge (Barsalou, 1999; Tulving, 1972).¹ In explaining the origins and development of semantic knowledge, the more empiricist cognitive scientists have tended to focus on the role of perception and sensorimotor grounding (Prinz, 2002; Rogers & McClelland, 2004). Such work acknowledges that we might ordinarily learn many facts about water, such as its chemical structure, from language, but the more important conceptual “cores” are learned from our interactions with the world. We don’t need language to learn that water is wet. On the more rationalist side, theorists have

instead stressed the importance of innate conceptual knowledge and abstract reasoning (Bedny & Saxe, 2012; Quilty-Dunn et al., 2022). In artificial intelligence (AI), mastery of natural language has been a long-standing goal: A system that could understand natural language could be controlled by being spoken to. The open-endedness of language also meant that its successful use by a machine could be used as an intelligence metric, as in Turing’s famous imitation game. The use of language was thus viewed as a problem to be solved by a sufficiently intelligent system rather than a way of making a system appropriately intelligent.

But what if this thinking is backward? Could language be a key ingredient in creating human-like intelligence in the first place? Striking evidence for this possibility comes from what happens when artificial

Corresponding Author:

Gary Lupyan, Department of Psychology, University of Wisconsin–Madison

Email: lupyan@wisc.edu

neural networks are exposed to large amounts of language. Large language models (LLMs) such as ChatGPT, Claude, and Gemini are neural network transformers trained through the self-supervised prediction of upcoming text.² Given a context (a chunk of text), the model makes predictions of what is likely to follow, adjusting its weights to minimize the error between its guess and what actually happens next. To learn that “Good” is frequently followed by “morning” does not require learning much beyond simple transition probabilities. But when a sufficiently large model is trained on a sufficiently large corpus, three surprising things happen.

First, the models learn language. Although it was possible a few years ago to deny that LLMs “really” learn language (e.g., Dentella et al., 2023), it is no longer possible to reliably distinguish language produced by people and state-of-the-art LLMs (Hu et al., 2024; Sadasivan et al., 2025).³ A key innovation that led to this progress has been the development of the multi-head “attention” mechanism that allows the models to incorporate context into the internal representation of the linguistic prompt in a more powerful way than previously possible.⁴ These models are generalized pattern learners. They learn language despite lacking specialized language-learning machinery of the sort that has long been argued to be necessary for any system to learn language (e.g., Lidz & Gagliardi, 2015). The same transformer architecture that powers modern language models can be put to use when performing image classification (Dosovitskiy et al., 2021) and to simulate chick development (Wood et al., 2024).

Second, in the course of learning to produce human-like language, LLMs learn some of the very things that have been considered to be prerequisites to language learning (Piantadosi, 2024). For example, we know that human language understanding requires a heavy dose of pragmatic inference (Heintz & Scott-Phillips, 2023). Learning and using language also seems to require a system to be systematic: Understanding “Mary loves John” implies that the system can understand “John loves Mary.” But before being trained on language, LLMs know nothing of pragmatic inference and are hardly systematic. These abilities emerge in LLMs with exposure to language (Hu et al., 2023; Lepori et al., 2023).

Third, and perhaps most surprising, training these general-purpose networks on large amounts of natural language has enabled LLMs to perform a huge range of practical downstream tasks, ranging from summarizing and editing texts to diagnosing patients with apparently superhuman accuracy (Goh et al., 2024).⁵ The rapid uptake of LLMs across a wide swath of society speaks to the practical usefulness of these systems.

The wide-ranging abilities of LLMs raise two questions. First, is it a coincidence that these advances have come from using *natural language* as the main training data, or is there something special about language? Second, if language provides such effective training for artificial neural networks, might language input be more instrumental to *human* intelligence than many cognitive scientists have tended to assume (for recent philosophical treatments, see Chalmers, 2024; Clatterbuck, 2024; Rothschild, 2025)?

The most direct way of answering the first question would require systematically comparing neural networks trained on language to those trained on only nonlinguistic data. If language is not necessary, it should be possible for nonlinguistic models to achieve the same performance (on nonlinguistic versions of tasks) as their language-trained counterparts. Nonlinguistic input should be sufficient (indeed in most cases more effective than language) for giving rise to the kinds of intelligence we find in nonhuman animals. After all, other animals manage to get by without the benefit of language. But when it comes to the types of intelligence that are more uniquely human—including a sophisticated theory of mind, relational and analogical reasoning, and the ability to learn a wide set of non-survival-related skills—we predict that such nonlinguistic AI systems would struggle.

According to some recent work in cognitive neuroscience, the answer to the second question is that no matter the usefulness of language for training AI models, its function in human cognition is highly circumscribed. For example, Fedorenko and colleagues have shown that explicit language tasks such as hearing or reading sentences activates a “language network” (which includes the left inferior frontal and middle frontal gyri and anterior temporal lobe). This language network is not activated by nonverbal tasks such as numerical cognition, understanding actions, and tasks probing theory of mind (Fedorenko, Ivanova, & Regev, 2024).

On the basis of this dissociation between brain networks involved in overtly linguistic tasks and in other cognitive tasks, Fedorenko, Piantadosi, and Gibson (2024) argued that the role of language is strictly communicative and that thought is independent of language. But although this work has been useful in helping us understand the neural substrates of language processing, it conflicts with a range of findings from cognitive science.

Just as manipulating linguistic experience is the best way to understand the role of language in artificial systems, we can ask if manipulating linguistic experience in humans affects human cognition. Although we cannot deliberately deprive people of language, we can

glean valuable insights from cases in which people do not receive typical language input or from individuals who suffer from language impairments. We can also investigate typical development and study the associations between children's emerging language abilities to their cognition. Last, we can experimentally manipulate the availability of language during a task and measure the effects of this manipulation on behavior.

These approaches, when taken together, provide converging evidence that language plays a transformative role in human cognition. For example, children who are born deaf and are not exposed to a conventional sign language struggle with a range of cognitive tasks such as theory of mind (Gagne & Coppola, 2017) and spatial reasoning (Gentner et al., 2013). People who suffer language impairments in adulthood (e.g., stroke-induced aphasia) show impairments in putatively “non-verbal” (fluid) reasoning (Baldo et al., 2015) and in tasks requiring selectively attending to a particular dimension (e.g., appreciating that cherries and bricks are similar by virtue of their color; Koemeda-Lutz et al., 1987)—an ability that underlies much of abstract reasoning. Interestingly, they do not show pronounced deficits in theory of mind (Siegal & Varley, 2006), suggesting that the role of language may be to inform the initial development of theory of mind (after all, the major way we come to know what others are thinking is by having them tell us). Experimental studies that manipulate language demonstrate its causal role in human cognition. For example, holding nonlinguistic experience constant, named categories, and categories consisting of more nameable parts are easier to learn (e.g., Zettersten & Lupyan, 2020); hearing a word activates more categorical mental representations (Edmiston & Lupyan, 2015), which is important for making inferences and for enabling compositional thought. Conversely, interfering with language impairs people's ability to learn rules and attend to specific dimensions (Lupyan, 2009), mirroring some of the impairments observed in aphasia.

Neural Dissociations Do Not Imply Lack of Causality

How can we square these findings with apparent neural dissociations between linguistic and nonlinguistic processes? One answer is that many aspects of language such as phonological and syntactic processing are highly specialized. As we gain linguistic expertise, these processes become modularized and neurally dissociable. But as the evidence above suggests, such emerging modularity does not mean that “nonlinguistic” cognition is independent of language.

To take one example of such interaction, although visual processing is by no means a linguistic process, language actively modulates and constrains it. Language activates both the primary visual cortex (Seydell-Greenwald et al., 2023) and higher level visual regions (Ardila et al., 2015) and modulates basic visual processes. For example, a simple word can make otherwise invisible objects visible (Lupyan & Ward, 2013). The addition of language input also vastly improves the match between visual representations in people and those learned by artificial neural networks trained on images (Wang et al., 2023; see also Bi, 2021). Despite being dissociable, language informs and shapes visual processing. These causal links do not imply that the medium of thinking or perceiving is linguistic. Rather, they suggest that forming useful internal models (both developmentally and when executing a task) benefits from learning and using the abstractions provided to us by natural language.

Investigations of how LLMs learn to process language (e.g., noun-verb agreement) are revealing the emergence of specialized circuits (Tigges et al., 2024). But this does not license distinguishing “formal” linguistic competence, which underlies LLMs' ability to produce fluent coherent language, from “functional” linguistic competence, which involves using language for downstream tasks (Mahowald et al., 2024). It is indeed useful to distinguish linguistic tasks such as noun-verb agreement from various functional downstream tasks such as a medical diagnosis. LLMs can do both, but the mechanisms are likely quite different because the computational needs of the tasks are different. Yet it would clearly be a mistake to conclude from such dissociations that the ability of LLMs to diagnose patients is independent of language. It is of interest that the performance of many (although importantly not all) downstream tasks can be improved not by augmenting training with nonlinguistic materials but simply by having models use internal language—analogue to inner speech—before producing a response (Chen et al., 2025).

Why Does Language Have These Effects?

Why does language have these effects on people, and how can exposure to natural language, no matter its scale, enable LLMs to perform so many different types of cognitive tasks? One reason is that language provides a set of abstractions encoded into its vocabulary. Although it may seem that words simply reflect natural categories (the “joints of nature”), vocabularies are in fact products of collective intelligence. Very few if any of us are capable of inventing number words to denote cardinalities, but once these words exist, we readily

learn them in the course of learning the rest of language. Absent number words, we struggle to represent exact numerosities even when living in a numerate culture (Spaepen et al., 2011). The vocabulary of every language consists of thousands of such prediscovered abstractions that are much easier to learn than to reinvent. Vocabularies and larger verbal constructions can thus be viewed as highly generative compression schemes for capturing the human Umwelt (see also Clatterbuck & Gentry, 2025; Rothschild, 2025).

When we expose sufficiently powerful statistical learning systems (such as transformers) to language and force them to get good at predicting the next word, they come to learn not only the statistical patterns of language but also the generative model of the latent structure that produced the language, that is, the world through human eyes, including perceptual relationships (Marjeh et al., 2024; Wang et al., in press) and causal links (Kıcıman et al., 2024) thought to be unlearnable from passive observation, much less from passive observation of language (Sloman, 2005). To paraphrase a recent social media post from @gracekind.net: What does ChatGPT do? It predicts the next token. What does it do to predict the next token? Whatever it takes (see also Aguera y Arcas, 2025).

Human and Machine Intelligence as Collective, Language-Based Intelligence

In trying to understand the astonishing intellectual achievements of our species, it has been common to appeal to the computational power of individual minds. Recognizing that language is both the product of our collective intelligence and also its shaper encourages a more humble view: The power of human intelligence has less to do with our individual mental firepower and more to do with the scaffolding provided by language (and other aspects of culture; Henrich, 2015). The presence of certain words in a vocabulary ensures that all the members of the language community learn the categories these words denote. Learning these categories paves the way for their creative recombination. From this point of view, language is not just a communication medium; its abstractions help us form internal models we use to cognize the world.

Adopting this perspective makes the successes of artificial neural networks trained on language less surprising: The open-ended nature of language means that it can be used to convey everything from how we feel, to recipes, to scientific findings. Reducing prediction error across these varied domains turns out to be a highly effective means for learning the latent structure of these data: the human Umwelt. Despite the vast differences

between LLMs and human minds, language appears to help both. As our creations, the intelligence of LLMs depends on the cognitive labor of countless human minds. So too does our own intelligence.

Transparency

Action Editor: Arturo E. Hernandez

Editor: Arturo E. Hernandez

ORCID iDs

Gary Lupyan  <https://orcid.org/0000-0001-8441-7433>

Hunter Gentry  <https://orcid.org/0000-0003-4189-0911>

Martin Zettersten  <https://orcid.org/0000-0002-0444-7059>

Notes

1. Exceptions include neo-Whorfian work in the tradition of Melissa Bowerman and Stephen Levinson; work in developmental psychology, including work by Sandy Waxman and Dedre Gentner; and work by Paul Harris et al. on learning from testimony.
2. We are deliberately simplifying the training procedure for the purposes of the current article. The bulk of training is indeed self-supervised prediction of text. However, the ability of LLMs to perform the many downstream tasks they are capable of requires a relatively small amount of supervised training to provide illustrations of, for example, what it means to summarize a document. Importantly, the effectiveness of this supervised training depends on having a sufficiently large base model trained through self-supervised language prediction.
3. The impossibility of distinguishing whether a particular text was written by a person or an LLM does not entail that human and LLM-produced language are identical in aggregate. For example, there is compelling evidence that language produced by LLMs is more uniform compared with language produced by people (Sourati et al., 2025).
4. We put the term “attention” in quotes in this sentence to highlight that it is misleading to think of this mechanism in terms of human attention. It is, more accurately, a form of adaptive kernel smoothing.
5. Although it is difficult to compare data efficiency in an apples-to-apples way, it is certainly the case that LLMs are exposed to orders of magnitude more language than any person. Compared with biological systems, these models also require vast amounts of power to operate. These differences should temper the urge to draw overly strict analogies between artificial and biological neural networks.

References

- Aguera y Arcas, B. (2025). *What Is Intelligence?: Lessons from AI About Evolution, Computing, and Minds*. MIT Press. <https://mitpress.mit.edu/9780262049955/what-is-intelligence/>
- Ardila, A., Bernal, B., & Rosselli, M. (2015). Language and visual perception associations: Meta-analytic connectivity modeling of Brodmann area 37. *Behavioural*

- Neurology*, 2015(1), Article 565871. <https://doi.org/10.1155/2015/565871>
- Baldo, J. V., Paulraj, S. R., Curran, B. C., & Dronkers, N. F. (2015). Impaired reasoning and problem-solving in individuals with language impairment due to aphasia or language delay. *Frontiers in Psychology*, 6, Article 1523. <https://doi.org/10.3389/fpsyg.2015.01523>
- Barsalou, L. W. (1999). Perceptual symbol systems. *The Behavioral and Brain Sciences*, 22(4), 577–609.
- Bedny, M., & Saxe, R. (2012). Insights into the origins of knowledge from the cognitive neuroscience of blindness. *Cognitive Neuropsychology*, 29(1–2), 56–84. <https://doi.org/10.1080/02643294.2012.713342>
- Bi, Y. (2021). Dual coding of knowledge in the human brain. *Trends in Cognitive Sciences*, 25(10), 883–895. <https://doi.org/10.1016/j.tics.2021.07.006>
- Chalmers, D. J. (2024). *Does thought require sensory grounding? From pure thinkers to large language models*. arXiv. <https://doi.org/10.48550/arXiv.2408.09605>
- Chen, Q., Qin, L., Liu, J., Peng, D., Guan, J., Wang, P., Hu, M., Zhou, Y., Gao, T., & Che, W. (2025). *Towards reasoning era: A survey of long chain-of-thought for reasoning large language models*. arXiv. <https://doi.org/10.48550/arXiv.2503.09567>
- Clatterbuck, H. (2024). Hume's externalist gambit. *Philosophy of Science*, 91(5), 1296–1305. <https://doi.org/10.1017/psa.2023.139>
- Clatterbuck, H., & Gentry, H. (2025). Learning incommensurate concepts. *Synthese*, 205(3), Article 106. <https://doi.org/10.1007/s11229-024-04883-7>
- Dentella, V., Günther, F., & Leivada, E. (2023). Systematic testing of three language models reveals low language accuracy, absence of response stability, and a yes-response bias. *Proceedings of the National Academy of Sciences, USA*, 120(51), Article e2309583120. <https://doi.org/10.1073/pnas.2309583120>
- Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., Uszkoreit, J., & Houlsby, N. (2021). *An image is worth 16x16 words: Transformers for image recognition at scale*. arXiv. <https://doi.org/10.48550/arXiv.2010.11929>
- Edmiston, P., & Lupyan, G. (2015). What makes words special? Words as unmotivated cues. *Cognition*, 143, 93–100. <https://doi.org/10.1016/j.cognition.2015.06.008>
- Fedorenko, E., Ivanova, A. A., & Regev, T. I. (2024). The language network as a natural kind within the broader landscape of the human brain. *Nature Reviews Neuroscience*, 25(5), 289–312. <https://doi.org/10.1038/s41583-024-00802-4>
- Fedorenko, E., Piantadosi, S. T., & Gibson, E. A. F. (2024). Language is primarily a tool for communication rather than thought. *Nature*, 630(8017), 575–586. <https://doi.org/10.1038/s41586-024-07522-w>
- Gagne, D. L., & Coppola, M. (2017). Visible social interactions do not support the development of false belief understanding in the absence of linguistic input: Evidence from deaf adult homesigners. *Frontiers in Psychology*, 8, Article 837. <https://doi.org/10.3389/fpsyg.2017.00837>
- Gentner, D., Özyürek, A., Gürcanli, Ö., & Goldin-Meadow, S. (2013). Spatial language facilitates spatial cognition: Evidence from children who lack language input. *Cognition*, 127(3), 318–330. <https://doi.org/10.1016/j.cognition.2013.01.003>
- Goh, E., Gallo, R., Hom, J., Strong, E., Weng, Y., Kerman, H., Cool, J. A., Kanjee, Z., Parsons, A. S., Ahuja, N., Horvitz, E., Yang, D., Milstein, A., Olson, A. P. J., Rodman, A., & Chen, J. H. (2024). Large language model influence on diagnostic reasoning: A randomized clinical trial. *JAMA Network Open*, 7(10), Article e2440969. <https://doi.org/10.1001/jamanetworkopen.2024.40969>
- Heintz, C., & Scott-Phillips, T. (2023). Expression unleashed: The evolutionary and cognitive foundations of human communication. *Behavioral and Brain Sciences*, 46, Article e1. <https://doi.org/10.1017/S0140525X22000012>
- Henrich, J. (2015). *The secret of our success: How culture is driving human evolution, domesticating our species, and making us smarter*. Princeton University Press.
- Hu, J., Floyd, S., Jouravlev, O., Fedorenko, E., & Gibson, E. (2023). *A fine-grained comparison of pragmatic language understanding in humans and language models*. arXiv. <https://doi.org/10.48550/arXiv.2212.06801>
- Hu, J., Mahowald, K., Lupyan, G., Ivanova, A., & Levy, R. (2024). Language models align with human judgments on key grammatical constructions. *Proceedings of the National Academy of Sciences of the United States of America, USA*, 121(36), Article e2400917121. <https://doi.org/10.1073/pnas.2400917121>
- Kıcıman, E., Ness, R., Sharma, A., & Tan, C. (2024). *Causal reasoning and large language models: Opening a new frontier for causality*. arXiv. <https://doi.org/10.48550/arXiv.2305.00050>
- Koemeda-Lutz, M., Cohen, R., & Meier, E. (1987). Organization of and access to semantic memory in aphasia. *Brain and Language*, 30(2), 321–337.
- Lepori, M., Serre, T., & Pavlick, E. (2023). Break it down: Evidence for structural compositionality in neural networks. *Advances in Neural Information Processing Systems*, 36, 42623–42660.
- Lidz, J., & Gagliardi, A. (2015). How nature meets nurture: Universal grammar and statistical learning. *Annual Review of Linguistics*, 1, 333–353. <https://doi.org/10.1146/annurev-linguist-030514-125236>
- Lupyan, G. (2009). Extracommunicative functions of language: Verbal interference causes selective categorization impairments. *Psychonomic Bulletin & Review*, 16(4), 711–718. <https://doi.org/10.3758/PBR.16.4.711>
- Lupyan, G., & Ward, E. J. (2013). Language can boost otherwise unseen objects into visual awareness. *Proceedings of the National Academy of Sciences, USA*, 110(35), 14196–14201. <https://doi.org/10.1073/pnas.1303312110>
- Mahowald, K., Ivanova, A. A., Blank, I. A., Kanwisher, N., Tenenbaum, J. B., & Fedorenko, E. (2024). Dissociating language and thought in large language models. *Trends*

- in *Cognitive Sciences*, 28(6), 517–540. <https://doi.org/10.1016/j.tics.2024.01.01>
- Marjeh, R., Sucholutsky, I., van Rijn, P., Jacoby, N., & Griffiths, T. L. (2024). Large language models predict human sensory judgments across six modalities. *Scientific Reports*, 14(1), Article 21445. <https://doi.org/10.1038/s41598-024-72071-1>
- Piantadosi, S. T. (2024). Modern language models refute Chomsky's approach to language. *From Fieldwork to Linguistic Theory*, 353–414.
- Prinz, J. J. (2002). *Furnishing the mind: Concepts and their perceptual basis*. MIT Press.
- Quilty-Dunn, J., Porot, N., & Mandelbaum, E. (2022). The best game in town: The reemergence of the language of thought hypothesis across the cognitive sciences. *The Behavioral and Brain Sciences*, 46, Article e261. <https://doi.org/10.1017/s0140525x22002849>
- Rogers, T. T., & McClelland, J. L. (2004). *Semantic cognition: A parallel distributed processing approach*. Bradford Book.
- Rothschild, D. (2025). *Language and thought: The view from LLMs*. arXiv. <https://doi.org/10.48550/arXiv.2505.13561>
- Sadasivan, V. S., Kumar, A., Balasubramanian, S., Wang, W., & Feizi, S. (2025). *Can AI-generated text be reliably detected?* arXiv. <https://doi.org/10.48550/arXiv.2303.11156>
- Seydell-Greenwald, A., Wang, X., Newport, E. L., Bi, Y., & Striem-Amit, E. (2023). Spoken language processing activates the primary visual cortex. *PLOS ONE*, 18(8), Article e0289671. <https://doi.org/10.1371/journal.pone.0289671>
- Siegal, M., & Varley, R. (2006). Aphasia, language, and theory of mind. *Social Neuroscience*, 1(3–4), 167–174. <https://doi.org/10.1080/17470910600985597>
- Sloman, S. (2005). *Causal models: How people think about the world and its alternatives*. Oxford University Press.
- Sourati, Z., Karimi-Malekabadi, F., Ozcan, M., McDaniel, C., Ziabari, A., Trager, J., Tak, A., Chen, M., Morstatter, F., & Dehghani, M. (2025). *The shrinking landscape of linguistic diversity in the age of large language models*. arXiv. <https://doi.org/10.48550/arXiv.2502.11266>
- Spaepen, E., Coppola, M., Spelke, E. S., Carey, S. E., & Goldin-Meadow, S. (2011). Number without a language model. *Proceedings of the National Academy of Sciences, USA*, 8(8), 3163–3168. <https://doi.org/10.1073/pnas.1015975108>
- Tigges, C., Hanna, M., Yu, Q., & Biderman, S. (2024). *LLM circuit analyses are consistent across training and scale*. arXiv. <https://doi.org/10.48550/arXiv.2407.10827>
- Tulving, E. (1972). Episodic and semantic memory. In E. Tulving & W. Donaldson (Eds.), *Organization of memory* (pp. 381–403). Academic Press.
- Wang, A. Y., Kay, K., Naselaris, T., Tarr, M. J., & Wehbe, L. (2023). Better models of human high-level visual cortex emerge from natural language supervision with a large and diverse dataset. *Nature Machine Intelligence*, 5(12), 1415–1426. <https://doi.org/10.1038/s42256-023-00753-y>
- Wang, Z., Akshi, Keil, S., Kim, J. S., & Bedny, M. (in press). Constructing meaning from language: Visual knowledge in people born blind and in large language models. *Annual Review of Linguistics*.
- Wood, J. N., Pandey, L., & Wood, S. M. W. (2024). Digital twin studies for reverse engineering the origins of visual intelligence. *Annual Review of Vision Science*, 10(1), 145–170. <https://doi.org/10.1146/annurev-vision-101322-103628>
- Zettersten, M., & Lupyan, G. (2020). Finding categories through words: More nameable features improve category learning. *Cognition*, 196, Article 104135. <https://doi.org/10.1016/j.cognition.2019.104135>