

Dumb Meaning: Machine Learning and Artificial Semantics¹

Hannes Bajohr

In June 2022, Google employee Blake Lemoine was given an indefinite leave of absence. The reason: He had claimed that the artificial intelligence he was helping to test was sentient, and the company thought such a claim bad press (Tiku 2022).² Lemoine insisted that LaMDA, a chatbot system, convinced him in lengthy conversations that it had the intelligence of a highly gifted eight-year-old, and asked to be considered a person with rights (Lemoine 2022b).³ In doing so, Lemoine, who describes himself as “ordained as a mystic Christian priest”, was merely exaggerating a sentiment that also afflicted others at Google. Blaise Agüera y Arcas, a senior machine learning engineer not usually prone to mysticism, wrote of his own interactions with LaMDA just days before Lemoine: “I felt the ground shift under my feet. I increasingly felt like I was talking to something intelligent”. (Agüera y Arcas 2022)

In contrast, a discussion about another AI system, which took place at about the same time, did not use the buzzwords of sentience and intelligence at all. Dall-E, whose second version was developed by the company OpenAI around this time (since September 2023, the third version is integrated into ChatGPT), is a text-to-image AI that can generate images from natural language input. Given a prompt such as “A Shiba-Inu wearing a beret and a black turtleneck”, it produces an output image depicting that very scene (Ramesh et al. 2022). The public beta triggered a slew of experiments, and soon the most interesting or whimsical results were shared on the web and especially on Twitter.

This, too, was revealing: compared to the much less successful experiments with autonomous cars, it suggested that AI has significantly different social effects than long thought – that, before it puts truck drivers out of business, it is more likely to take the jobs of illustrators, graphic artists, and stock photographers (Prakash

1 This essay first appeared in German; the current version is a translation of Bajohr 2024b.

2 Lemoine's term is *sentience*, not *consciousness*, but he seems to use them synonymously.

3 In addition, Lemoine also published the chat transcript of a conversation with LaMDA (Lemoine 2022a)

2022).⁴ Unlike in the case of LaMDA, however, no one thought Dall-E should be conceived of as a person with rights.

The different reactions to the two systems show how quickly thinking about AI veers into familiar conceptual ruts. Intelligence, consciousness, sentience, and personhood have been the major themes of AI research and its imaginaries for nearly seventy years; amusing little pictures, by contrast, seem to raise fewer fundamental questions. But it is quite possible that it is actually the other way around – that the eternal hunt for superintelligence and the singularity obscures the more interesting and subtle conceptual shifts that escape both the tech evangelists in their visionary furor and their skeptical critics.

For philosopher Benjamin Bratton, it is clear that in the face of these new AI systems, “reality has outpaced the available language to parse what is already at hand”. What is needed, therefore, is a “more precise vocabulary” (Bratton and Agüera y Arcaas 2022) that goes beyond the usual handful of big concepts, but also beyond the anthropocentric assumption that the only way in which machines may form world-relations would have to be our way. We can observe such a tendency with Dall-E and LaMDA. Here, the concept of meaning becomes detached from its anthropocentric correlate. It would be meaning without mind – dumb meaning.

Free-floating and grounded systems

Despite constant admonitions from computer scientists, linguists, and cognitive psychologists to use terms such as intelligence and consciousness with care, the tech industry remains relatively immune to such warnings. Thus, critics soon accused Lemoine of having fallen for the “ELIZA effect” (Christian 2022) – of having projected intelligence and consciousness onto LaMDA – a susceptibility Joseph Weizenbaum had already observed in 1966 among users of his ELIZA chatbot: Although ELIZA merely mimicked a Rogerian psychoanalyst, mirroring the patient’s statements back to them as questions, its users behaved as if the program really were a conscious agent interested in their well-being.

The classic objection here is the following: Computers are symbol-processing systems that deal with syntax alone, not with semantics – they can process logical forms but not substantive meaning (Cramer 2008). For their operations it is irrelevant which objects or concepts the symbols name in a human world and which cultural valences are associated with them. Thus, ELIZA merely scans user input for a

4 The June 11, 2022, issue of *The Economist* featured an illustration generated by an image AI. Since then, this has become somewhat of a fashion that will, without a doubt, soon give way to more sophisticated uses.

given syntactic pattern and transforms it into a “response” according to a transformation rule. Weizenbaum gives the example in which the analysand reproaches the analyst: “It seems that you hate me”. The program identifies the key pattern “ x you y me” in this sentence and separates it accordingly into the four elements “It seems that”, “you”, “hate” and “me”. It then discards y (“it seems that”) and inserts x (“hate”) into the reply template “What makes you think I x you”. And so ELIZA responds to the accusation that it hates the analysand by asking how they got that idea (Weizenbaum 1966).⁵

This interaction may have meaning for the user and plausibly suggest a communicative intent on the part of ELIZA, but neither such intent nor such meaning is actually to be found in the program. It has merely processed symbols according to a rule without “knowing” what hate is or what behavior the mores of civil discourse suggest. That is the difference between the processing of information and the understanding of meaning.

For AI researchers who seek to make computers more human, this state of affairs describes what cognitive psychologist Stevan Harnad called the “symbol grounding problem”: Symbols, like those in Weizenbaum’s transformation operation, have no intrinsic meaning for computers because, without the background of practical knowledge of the world, they can only refer to other symbols, never to any reality beyond them. They are not grounded in the world, and there is no way out of this “symbol / symbol merry-go-round”. Whatever meaning there is can only be “parasitic”, and is projected onto the output by human interpreters. (Harnad 1990:340, 339)

Harnad’s criticism, however, was directed against only one particular type of AI, which also includes ELIZA; for obvious reasons, it is called “symbolic”. To solve the symbol grounding problem, Harnad relied on the novel “*subsymbolic*” or “*connectionist*” systems of the time: neural networks of which LaMDA and Dall-E are late descendants. Unlike traditional AI, they are not designed as a set of logical rules of inference, but are vaguely modeled after the brain as neurons and synapses that amplify or attenuate the signals passed through them. They therefore do not require explicit symbolic representations and rules – they are not programmed, but learn independently from examples. While neural networks were mainly used for pattern recognition in the early 1990s, Harnad thought they might be able to access the world. Implemented in an autonomous, mobile robot, equipped with sensors and effectors, a conglomerate of neural networks would first receive impressions and categorize them as recognizable shapes. These would then be handed over to a symbolic AI, but would now no longer be mere references to other symbols but rather

5 I have simplified the procedure somewhat; moreover, ELIZA allows quite different transformation rules, and the therapist is only one subroutine, called DOCTOR.

connected to the world via their causal reference to external data – they would finally be grounded (Harnad 1993).

The consequence of this thought, however, seems to be that the only way to get around the ELIZA effect, which falsely attributes consciousness to computers, is to *actually* give them consciousness. For what Harnad has in mind is, in the end, again an anthropocentric model that hopes embodied cognition and sufficiently extensive referential meanings will produce world understanding, since this is how we more or less function, too. The success of his hybrid model would have to be demonstrated by his robot being as competent at navigating the world as if it were actually intelligent. Since this is not yet the case, the symbol grounding problem cannot yet be considered solved either; a *bit* of meaning does not exist here by definition. And yet, such limited meaning is exactly what LaMDA and Dall-E seem to suggest.

Gradated meaning

With the increasing popularity that neural networks have enjoyed for almost ten years now, the idea that they somehow could have access to meaning beyond mere ungrounded symbols has also become more attractive again. For media studies scholar Mercedes Bunz, neural networks, thanks to their complexity and capacity for unsupervised learning, can now “calculate meaning” rather than just empty symbols (Bunz 2019). And it is true that, in the face of neural networks, the binary distinction between meaning (human world) and non-meaning (digital systems) is becoming increasingly difficult to maintain. Maybe we should consider levels of gradated meaning which, as artificial semantics, no longer presuppose a mind.

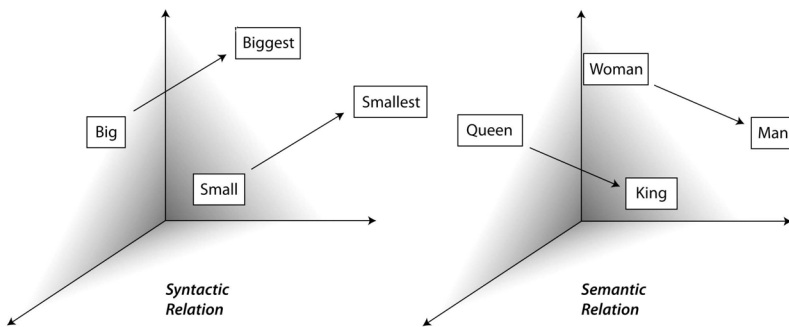
Thus, rather than taking it as a sign of consciousness, the fact that LaMDA's answers sounded so human-like can simply be understood as an indication of such “dumb” meaning. While “broad” meaning presupposes – depending on your philosophical or disciplinary orientation – embodied intelligence, cultural and social background knowledge, or the world-disclosing function of language, dumb meaning would operate below this scale (which is always calibrated on humans) and could best be grasped as an effect of *correlation*.⁶

LaMDA is – similar to the better-known text generator ChatGPT – a large language model implemented as a neural network. Trained on vast amounts of text, it

6 Dumb Meaning excludes *natural* meaning (such as in the symptom/disease relation). It also cannot contain the *intentionalist* meaning Paul Grice has theorized, according to whom the meaning of an utterance is dependent on recognizing the speaker's intention, which in turn requires consciousness. And finally, it is only in a very limited way a *use theory* in the tradition of the late Wittgenstein, since “use” presupposes a shared social background, which requires world-understanding, which, in turn, assumes embodied intelligence.

processes language as a multi-dimensional vector space, a so-called “word embedding”, which works according to the principle of staggered correlations first suggested in the “distributional hypothesis” in the 1950s (Harris 1954; Firth 1957): First, words that frequently appear together are closer together in this space, forming semantic clusters familiar from word clouds. However, since not only the correlations of words to words but also correlations of correlations are encoded, large language models can also explicate implicit regularities that are not spelled out in the training text. This is true for syntactic relations – when the Euclidean distance between the vectors for the positive and superlative of a word is the same – but also for complex semantic relations, that is, word meaning. One of the best-known examples of this principle is the operation: $v_{\text{king}} - v_{\text{man}} + v_{\text{woman}} \approx v_{\text{queen}}$. (Mikolov, Yih, and Zweig 2013)⁷

Fig. 1: Word embedding of a large language model



In this equation – which reads: subtract from the word vector “king” that for “man” and add that for “woman,” and the result is the word vector for “queen” – the latent semantic relation “gender” emerges as an arithmetic correlation, even though it is not explicitly present in the model (fig. 1). That it arises from the mass of language on which the model is trained explains machine learning’s susceptibility to biases: sexism and racism may also be latently encoded in language models (Bender et al. 2021; Bajohr 2024d). The meaning of a sign in a language system constructed in this way is determined purely *differentially*, as in Ferdinand de Saussure’s linguistic structuralism. Instead of referring to anything outside language, sign meaning is simply thought of as difference from other signs and sign correlations (this is excellently explained in Gastaldi 2021).

⁷ This insight still applies to newer, technically different models such as GloVe (Global Vectors for Word Representation).

Large language models, such as ChatGPT, Claude or BARD and LaMDA, basically still follow this principle. However, they no longer vectorize individual words – which would be impossible given the huge amount of training data. Instead, ChatGPT uses a network architecture called the “transformer” to build embeddings of entire word sequences (Meng et al. 2023). To do so, transformers use the concept of “attention” (Brown et al. 2017). The model learns to establish relationships between all words in a sequence; from these attention patterns, the transformer constructs vectors that represent the entire sequence – so-called contextualized embeddings. In contrast to static word vectors, these sequence embeddings take into account a much larger context of use. In a decoder function, the transformer then uses these contextualized embeddings to generate new texts that match the input context. The ability to build such complex sequence representations is the key to the performance of large language models (Wolfram 2023). The effect, nevertheless, is that large language models, by their immense training data alone, are able to produce apparently situational understanding, as LaMDA did, without ever being “in a situation”⁸.

Language models would then be producers of a first degree of dumb meaning. It is dumb because the model captures latent correlations between signs, but still does not “know” what things these signs actually name; with this kind of meaning, one will not be able to build an intelligence that will ever find its way around in the world. The linguist Emily Bender, a vehement critic of all AI hype about alleged consciousness, admits with her colleague Alexander Koller that “a sufficiently sophisticated neural model *might* learn some aspects of meaning”, such as semantic similarity, but considers them to be “only a weak reflection of actual meaning”, which is always related to something in the world, i.e. “grounded”. (Bender and Koller 2020)

But as wrong as it would be to project anything like sentience or consciousness onto this system, one should also not be too quick to dismiss this modicum of meaning.⁹ Insofar as language models make implicit knowledge explicit in a nontrivial way – even if only by matrix transformations in a vector space – they produce dumb

8 This is philosopher Hubert Dreyfus’ term for the prior world-understanding that humans, but not computers, have (Dreyfus 1992).

9 In this respect, I agree that it is “productive to consider reference as just one (optional) aspect of a word’s full conceptual role” (Piantadosi and Hill 2022: 4). The paper by Piantadosi and Hill makes a similar argument to the present essay; however, I believe that the authors go too far in the direction of attributing “rich, causal and structured internal states” to LLMs, which again seems to border too much on anthropomorphism (Piantadosi and Hill 2022: 5). I would also like to note that I am unhappy with N. Katherine Hayles’s notion of computers as “cognizers” – a concept that also suggests a subjectivity on the part of the operating systems that I find difficult to subscribe to; on the positive side, however, Hayles emphasizes the production of meaning by such systems (Hayles 2019; Hayles 2022).

meaning which would not have been available to us without them.¹⁰ In contrast to ELIZA – whose x and y were only empty placeholders to the system – neural networks are not *solely* parasitically dependent on the meaning attributions of human agents, but *also* operate productively with the inherent distributional structure of language.

Text and image and world

Bender and Koller are of course right that LaMDA is not grounded.¹¹ It is a *unimodal* network, processing only a single type of data, namely text. To be grounded in Harnad's sense, she writes, it would be necessary to combine several types of data – it would have to be *multimodal* machine learning (Singer 2022). That is what Dall-E is: Instead of text just referring to other text, here text is correlated with image information. This raises the hope again that arbitrary signs can be linked to things in the world to produce grounded meaning.

Harnad's hypothesis that neural networks in particular could address the symbol grounding problem has recently been taken up by media scholars Leif Weatherby and Brian Justie with their notion of "indexical AI". It is named after Charles Sanders Peirce's notion of the index. Unlike the symbol, which has a purely conventional relationship to its signified (as "dog", "chien", and "Hund" all refer to the same thing), the index is causally linked to it (as smoke refers to fire).

With this coinage, the authors make Harnad's project the basis of a description of contemporary technological culture: "Digital systems, relying on the neural net, have left the world of mere symbol behind and have begun to ground themselves *here, now, for you* – they are able to *point* to real states of affairs". (Weatherby and Justie 2022, 382)¹² Neural networks bring the world – as the data on which they have been

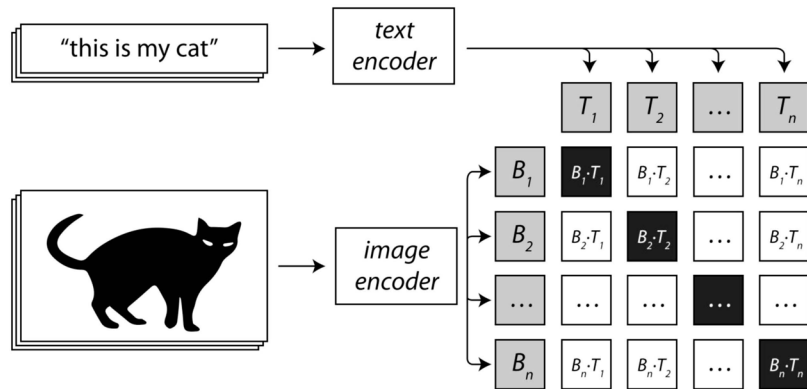
10 The assumption here is that this operation in fact finds something previously unknown and does not simply unfold a tautology; a model of this idea would be Kant's conviction that mathematical propositions are synthetic judgments a priori, that is, that they actually produce *new* knowledge.

11 While the paper presenting LaMDA also claims "groundedness" for the model, what is meant by this is simply that LaMDA's outputs are "grounded in known sources wherever they contain verifiable external world information". (Thoppilan et al. 2022, 2) As *textual sources*, they continue to be part of Harnad's "symbol/symbol merry-go-round" (Harnad 1990, 340).

12 One difficulty with this notion is the question of whether *all* data in a neural network should already be considered indexical (that would include the text of LaMDA), or only those obtained directly by sensors emulating physical senses (that would be images, but not text). Weatherby and Justie seem to have the former in mind, Harnad the latter. Harnad therefore speaks at one point of "iconic representations" through data (Harnad 1990, 342) – Peirce's third sign type, which operates on the principle of similarity between sign and signified. But since these are also indexical as they originate from sensors (which limits their scope

trained – into the computer, getting off of the solipsistic “symbol / symbol merry-go-round”. If we subscribe to this assertion for a moment, we see it plausibly demonstrated in Dall-E.

Fig. 2: Text-image correlation in CLIP



The heart of Dall-E is a machine learning model called CLIP. Via an encoder, it is fed with vectorized text-image pairs taken from the Internet – for example, a photo of a cat with the caption “this is my cat”. CLIP is trained to predict which text vector matches which image vector; the result is a comprehensive stochastic model that correlates image information with text information, but stores it as *one* type of information. In figure 2, this is the table in which the scalar product of the text and image vectors is listed – the better the text and image fit, the better this value; when the original image and text are paired, it is of course optimal (those are the black boxes running diagonally).

CLIP is thus remarkably good at *image recognition*: If you present it with an unknown cat photo, it nevertheless recognizes it as “cat”. But in a second step, it also becomes an *image generator*. To do this, it works in conjunction with another machine learning model called GLIDE (Guided Language to Image Diffusion for Generation and Editing), which has already been trained on a large data set of images.¹³ If the

to immediate, e.g. visual similarity), it seems to me that the argument of Weatherby/Justie and that of Harnad amount to something structurally similar – both are concerned with the connection between system and world, understood more or less broadly.

13 GLIDE is a *diffusion model* based on thermodynamic models, and thus functions differently from the GANs that were popular until recently, which combine two antagonistic submodels (Dhariwal and Nichol 2021). That the AI architectures used for an aesthetic work can themselves be a resource for discussing that work is something I suggest in Bajohr (2022).

user enters a prompt, GLIDE can use the text-image data stored in the CLIP model to reverse this process and synthesize an image that best correlates with the input text. In both operations – image recognition as well as image generation – it is again central that the models can learn and actively reproduce the *correlation* between textual descriptions of objects and their corresponding visual manifestations.

One may object that the image information correlated with the word “cat”, in which the photo of a cat is stored, may have an indexical relation to this cat – light was reflected from it and fell on a photo sensor etc. – but that even so the system will not learn what it means to share a world with a cat. Advocates of symbol grounding therefore try to extend what types of data an AI model gets fed – not only sensory but also motoric and eventually even social feedback: Only through the effects of language use in a community of other speakers inhabiting the same world can meaning be learned (Bisk et al. 2020).

But this claim would again mean to demand “full” human, that is, broad meaning, and to take anything below that not quite seriously. Instead, multimodal AI should be regarded as a second degree of dumb meaning. The Peircean indexical reference to something outside the model and the Saussurean differential reference to other elements within it are at any rate two distinct ways of meaning-making – if only that the dimension of possible correlations increases, and with it the possibility of unearthing unsuspected latent connections, unsuspected dumb meaning.

Indeed, multimodal AIs – besides Dall-E, for instance, Stable Diffusion, Google’s yet-to-be-released Imagen, or Midjourney – are capable of generating very complex text-image meanings. Their power lies in a capability that suggests that such correlations have a productive quality: In studying the deep structure of CLIP, computer scientists found that the model had trained single “neurons” that fired for both the word and the image of a thing. These were *conceptual* neurons in which the distinction between image and text tended to be overcome (Goh et al. 2021). Multimodality, at the neural level, is really *panmodality*, suggesting a semantics without clearly differentiated sign systems (this is also suggested by Merullo et al. 2022). Dumb meaning finds a new quality here, and is not tied to either text or image data, but encompasses both in a way that points to meaning beyond modal separation – and again has nothing to do with mind (see for more on this Bajohr 2024c).

Promptological investigations

AI systems *are* dumb. They have no consciousness. Yet they produce a complex artificial semantics that runs counter to our ordinary notions of meaning. Multimodal AI also shows that imputed consciousness and the meaning-capacity of a system have little to do with each other: The fact that LaMDA in particular seemed to Lemoine like a person – and not Dall-E, although one might argue that it represents a higher, be-

cause more correlation-rich stage of AI development – is simply due to the fact that it operates dialogically and thus is assumed to have communicative intent, whereas the image generator does not. Language always seems to be smarter than the image.

However, meaning beyond communicative intent need not be *merely* parasitic, as the vector operations of word embeddings and the conceptual neurons of text-to-image AIs show. That it is always *also* parasitic is due to the fact that the training data originate from a human world and artificial semantics is precisely not a “robot language” but a correlation effect of information that can be interpreted by humans. Nevertheless, in the long run, a convergence of dumb and broad meaning would be conceivable once they enter into mutually influencing circular processes.

The interface between natural and artificial semantics in the case of Dall-E is the interaction via prompt. On the one hand, “prompt design” – the precise, almost virtuosic selection of the text input – can be used analytically to scan the vector space of dumb meaning for traces of cultural knowledge. This would make the broad meaning of natural language, precisely in its interaction with dumb meaning, more important again.

A “promptology” that takes on such natural-artificial connections – the correlation of datafied language and the cultural meaning attributed to that language on the recipient side – would be a gateway for the humanities and cultural studies. With their knowledge of soft factors such as style, influence, iconography, etc., they could make useful contributions without necessarily taking the form of the more computer science-focused digital humanities; they could work in a phenomenon-oriented way and devote themselves to the artifacts that the model outputs as boundary objects between human and machine, between broad and dumb meaning.

At the same time, however, promptology is not merely an analytical procedure, but also a practice with its own knowledge, which has much to do with an almost “empathetic” interaction with the AI system. It has turned out that with text-to-image AIs, these prompts can be steered in unexpected directions simply by using certain, often counterintuitive or absurd formulations – there is already a start-up, PromptBase, which claims to sell particularly effective prompts (Wiggers 2022).¹⁴ Instead of subjugating the system and using it as an instrument, natural language instead must be adapted to the artificial semantics just to operate the system.

The result is a feedback loop of artificial and human meaning: not only does the machine learn to correlate the semantics of words with those of images we have given it, but we learn to anticipate the limitations of the system in our interaction

14 What is interesting here is that the discussed tendency to eliminate the speech/image distinction at the *technical* level is contrasted with the displacement of the image by speech at the *interface* level. The results of Dall-E could therefore also be understood as *language art* instead of being mere visual objects.

with it; this convergence would not be communicative in a strong sense, but perhaps in a weak, a dumb, sense.¹⁵

Bibliography

- Agüera y Arcas, Blaise (2022): "Artificial Neural Networks Are Making Strides Toward Consciousness". *The Economist*. June 9, 2022 <https://www.economist.com/by-invitation/2022/06/09/artificial-neural-networks-are-making-strides-towards-consciousness-according-to-blaise-aguera-y-arcas> (21 January, 2024).
- Bajohr, Hannes (2022): "Algorithmic Empathy: Toward a Critique of Aesthetic AI". *Configurations* 30(2): 203–31.
- Bajohr, Hannes (2024a): "On Artificial and Post-Artificial Texts: Machine Learning and the Reader's Expectations of Literary and Non-Literary Writing". *Poetics Today*, forthcoming.
- Bajohr, Hannes (2024b): "Dumme Bedeutung: Künstliche Intelligenz und artifizielle Semantik". Hannes Bajohr and Markus Krajewski (eds.): *Quellcodekritik: Zur Philologie von Algorithmen*. Berlin: August Verlag, 197–215.
- Bajohr, Hannes (2024c): "Operative Ekphrasis: The Collapse of the Text/Image Distinction in Multimodal AI". *Word & Image*, forthcoming.
- Bajohr, Hannes (2024d): "Whoever Controls Language Models Controls Politics". In: ke Arns, et al. (eds.): *Training the Archive*. Cologne: Walther König, 189–195.
- Bender, Emily M. / Timnit Gebru / Angelina McMillan-Major / Shmargaret Shmitchell (2021): "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?" *FACt '21: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–23.
- Bender, Emily M. / Alexander Koller (2020): "Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data". *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 5185–98. DOI: 10.18653/v1/2020.acl-main.463.
- Bisk, Yonatan / Ari Holtzman / Jesse Thomason / Jacob Andreas / Yoshua Bengio / Joyce Chai / Mirella Lapata / Angeliki Lazaridou / Jonathan May / Aleksandr Nisnevich / Nicolas Pinto / Joseph Turian (2020): "Experience Grounds Language". *ArXiv*. <http://arxiv.org/abs/2004.10151> (21 January, 2024).
- Bratton, Benjamin / Blaise Agüera y Arcas (2022): "The Model Is the Message". *Noema*. July 12. <https://www.noemamag.com/the-model-is-the-message> (21 January, 2024).
- Brown, Tom B., et al. (2020): "Language Models Are Few-Shot Learners." *Advances in Neural Information Processing Systems* 33, ed. by Hugo Larochelle, et al., 1–14.

15 See for this aspect Bajohr (2024a)

- Bunz, Mercedes (2019): "The Calculation of Meaning: On the Misunderstanding of New Artificial Intelligence as Culture". *Culture, Theory and Critique* 60(3–4): 264–78. DOI: 10.1080/14735784.2019.1667255.
- Christian, Brian (2022): "How a Google Employee Fell for the Eliza Effect". *The Atlantic*. June 21. <https://www.theatlantic.com/ideas/archive/2022/06/google-lambda-chatbot-sentient-ai/661322> (21 January, 2024).
- Cramer, Florian (2008): "Language". Matthew Fuller (ed.) *Software Studies: A Lexicon*. Cambridge, Mass.: MIT Press, 168–74.
- Dhariwal, Prafulla / Alex Nichol (2021): "Diffusion Models Beat GANs on Image Synthesis" *ArXiv*. <https://doi.org/10.48550/arXiv.2105.05233>.
- Dreyfus, Hubert L (1992): *What Computers Still Can't Do: A Critique of Artificial Reason*. Cambridge, Mass.: MIT Press.
- Firth, John R.: "A Synopsis of Linguistic Theory, 1930–1955". *Studies in Linguistic Analysis*. Special Volume of the Philological Society. Oxford: Blackwell, 1–32.
- Gastaldi, Juan Luis (2021): "Why Can Computers Understand Natural Language? The Structuralist Image of Language behind Word Embeddings". *Philosophy & Technology* 34(1): 149–214. DOI: 10.1007/s13347-020-00393-9.
- Goh, Gabriel / Nick Cammarata / Chelsea Voss / Shan Carter / Michael Petrov / Ludwig Schubert / Alec Radford / Chris Olah (2021): "Multimodal Neurons in Artificial Neural Networks". *Distill* 6(3): DOI: 10.23915/distill.00030.
- Harnad, Stevan (1990): "The Symbol Grounding Problem". *Physica D: Nonlinear Phenomena* 42(1–3): 335–46. DOI: 10.1016/0167-2789(90)90087-6.
- Harnad, Stevan (1993): "Grounding Symbols in the Analog World with Neural Nets". *Think* 2: 12.
- Harris, Zellig S. (1954): "Distributional Structure". *Word* 10(2–3): 146–162.
- Hayles, N. Katherine (2019): "Can Computers Create Meanings? A Cyber/Bio/Semiotic Perspective". *Critical Inquiry* 46(1): 32–55.
- Hayles, N. Katherine (2022): "Inside the Mind of an AI: Materiality and the Crisis of Representation". *New Literary History* 54(1): 635–666.
- Lemoine, Blake (2022a): "Is LaMDA Sentient? An Interview". *Medium*. June 11. <https://cajundiscordian.medium.com/is-lamda-sentient-an-interview-ea64d916d917> (21 January, 2024).
- Lemoine, Blake (2022b): "What Is LaMDA and What Does It Want?". *Medium*. June 11, 2022. <https://cajundiscordian.medium.com/what-is-lamda-and-what-does-it-want-688632134489>.
- Meng, Kevin, et al. (2023): "Locating and Editing Factual Associations in GPT". *arXiv*. <http://arxiv.org/abs/2202.05262>.
- Merullo, Jack / Louis Castricato / Carsten Eickhoff / Ellie Pavlick (2022): "Linearly Mapping from Image to Text Space" *ArXiv*. <http://arxiv.org/abs/2209.15162> (21 January, 2024).

- Mikolov, Tomas / Wen-tau Yih / Geoffrey Zweig (2013): "Linguistic Regularities in Continuous Space Word Representations" *Proceedings of the 2013 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. Atlanta, Georgia: Association for Computational Linguistics, 746–51. <https://aclanthology.org/N13-1090> (21 January, 2024).
- Piantadosi, Steven T., and Felix Hill (2022): "Meaning Without Reference in Large Language Models". *arXiv*. <http://arxiv.org/abs/2208.02957>.
- Prakash, Prarthana (2022): "AI Art Software Dall-E Moves Past Novelty Stage and Turns Pro". *Bloomberg*. August 3. <https://www.bloomberg.com/news/articles/2022-08-04/Dall-E-art-generator-begins-new-stage-in-ai-development> (11 January, 2023).
- Ramesh, Aditya / Prafulla Dhariwal / Alex Nichol / Casey Chu / Mark Chen (2022): "Hierarchical Text-Conditional Image Generation with CLIP Latents". *ArXiv*. <https://arxiv.org/abs/2204.06125> (21 January, 2024).
- Singer, Gadi (2022): "Multimodality: A New Frontier in Cognitive AI". *Medium*. February 2. <https://towardsdatascience.com/multimodality-a-new-frontier-in-cognitive-ai-8279d0ee3baf> (21 January, 2024).
- Thoppilan, Romal / Daniel De Freitas / Jamie Hall / Noam Shazeer / Apoorv Kulshreshtha / Heng-Tze Cheng / Alicia Jin et al. (2022): "LaMDA: Language Models for Dialog Applications". *ArXiv*. <http://arxiv.org/abs/2201.08239> (21 January, 2024).
- Tiku, Nitasha (2022): "The Google Engineer Who Thinks the Company's AI Has Come to Life". *Washington Post*, June 11, 2022.
- Weatherby, Leif / Brian Justie (2022): "Indexical AI", in: *Critical Inquiry* 48(2): 381–415. DOI: 10.1086/717312.
- Weizenbaum, Joseph (1966): "ELIZA: A Computer Program for the Study of Natural Language Communication Between Man And Machine". *Communications of the ACM* 9(1): 36–45.
- Wiggers, Kyle (2022): "A Startup Is Charging \$1.99 for Strings of Text to Feed to Dall-E 2". *TechCrunch*. July 29, <https://techcrunch.com/2022/07/29/a-startup-is-charging-1-99-for-strings-of-text-to-feed-to-dall-e-2/> (21 January, 2024).
- Wolfram, Stephen (2023): "What Is ChatGPT Doing ... and Why Does It Work?" *Stephen Wolfram*. <https://writings.stephenwolfram.com/2023/02/what-is-chatgpt-doing-and-why-does-it-work/>

