

Marc Heimann

Hochschule Niederrhein in Krefeld and Mönchengladbach, Germany

<https://orcid.org/0000-0003-3757-1790>

marc.heimann@protonmail.com

VIOLENT HERMENEUTICS: AI AND THE WEAK TECHNICAL EVENT*

Abstract: This paper investigates how large language models (LLMs), particularly those based on transformer architectures, inadvertently restage key insights from Heidegger's and Lacan's theories of language while simultaneously revealing their limitations. By foregrounding the associative and metaphorical dynamics at the heart of both LLMs and continental accounts of language, the essay situates machine learning within a broader hermeneutic framework. It discusses how vectorial associationism parallels Heidegger's *als-*relation and how attention mechanisms resonate with Lacanian metaphor, enabling local reconfigurations of meaning. Yet, unlike human language, the model's 'worlding' remains a contingent performance under inscription, without an inner horizon of disclosure. What emerges is not Badiou's Event but a weak technical event: a temporary redrawing of adjacency and accessibility within latent space. This fragility of articulation underscores the infrastructural role of metaphor, poetics, and style in shaping machine outputs, making literary practices newly central to the technological condition. In turn, I argue that human-AI interaction is best understood not through the lens of ideology or subjectivity, but as a form of literacy or hermeneutics where inscription reorganizes the calculable field.

Keywords: *generative AI, language, metaphor, Heidegger, Lacan*

1. INTRODUCTION

In what follows, I will explore the parallels and limitations between large language models, particularly transformer architecture, and concepts of language found in Heidegger and Lacan. My aim is to show how LLMs re-stage the philosophical insights of Heidegger and Lacan about language, while also marking their limits, particularly regarding rupture and the Event. These models surprisingly demonstrate the conceptual strength of these thinkers' language theories in an engineered setting,

*



© 2026. This work is openly licensed via CC BY-NC 4.0., which enables reusers to distribute, remix, adapt, and build upon the material in any medium or format for noncommercial purposes only, and only so long as attribution is given to the creator. Canonical URL: <https://creativecommons.org/licenses/by-nc/4.0/>

<https://doi.org/10.70832/EVENTanditsMediation.18>

notably without reference to the tradition of language theory they are part of. LLMs accidentally instantiate the conditions described in these theories. Two key elements emerge here, allowing us to establish initial connections to large language models: a word- and association-based understanding of language and the prominent role of a metaphoric formalism. Sentences, Logic and Grammar, however, are simply epiphenomena of these two.

Notably, for both Heidegger and Lacan, language isn't an expressive tool of some pre-existing subject, it's the condition of possibility for disclosure. And where Lacan calls language an automatism,¹ Heidegger calls the commonly shared 'the they', the baseline as the *ens realissimum* of *Dasein*'s disclosure of being.² So, what I am comparing here is language in this automatic and non-authentic mode and language modeled. The subject comes later and in the case of the machine, not at all.

Situated at the intersection of continental philosophies of language and the engineering sciences, following earlier work,³ this essay then treats transformer LLMs not as 'subjects' that understand, but as devices that stage language's automatic baseline. Their vectorial associationism approximates a Heideggerian *als-Relation* in pre-predicative articulation, while attention allows Lacanian metaphor-like transpositions that locally reorganize what can be said during inference. Using natural language inputs, in this light, constitutes a *technological hermeneutics* that compels a temporary articulation, a sovereign inscription that reorders the model's symbolic network without altering its weights, instead of understanding LLMs as powerful machines of ideology,⁴ I rather consider them extraordinarily weak to language. What appears in them is novelty only in a weak technical sense: Not an uncounted element in Badiou's register of truth, but a redrawing of adjacency and accessibility within the counted situation in the machine. The paper therefore contributes to current debates that read current technologies through world-disclosure and *Gestell*,⁵ and to discussions of rhetoric and interpretation in human–AI

¹ Jacques Lacan, *Ecrits: The First Complete Edition in English*, trans. Bruce Fink (New York: Norton, 2006), 11.

² Martin Heidegger, *Being and Time: A Translation of Sein und Zeit*, SUNY Series in Contemporary Continental Philosophy (Albany, NY: State University of New York Press, 1996), 120.

³ Marc Heimann and Anne-Friederike Hübener, "Circling the Void: Using Heidegger and Lacan to Think About Large Language Models," *Cognitive Systems Research* 91 (2025): 101349. <https://doi.org/10.1016/j.cogsys.2025.101349>

⁴ Carl Öhman, "We Are Building Gods: AI as the Anthropomorphised Authority of the Past," *Minds and Machines* 34, no. 1 (2024): 8, <https://doi.org/10.1007/s11023-024-09667-z>.

⁵ E.g. David Plunkett, "Technology, Dwelling, and Nature as 'Resource': A Reading of (and Some Reflections on) Themes from the Later Heidegger," *Inquiry* 67, no. 10 (2024): 3657–3727. <https://doi.org/10.1080/0020174X.2021.1989030>.

interaction,⁶ while marking a decisive limit: The machine has no inner world; it merely performs a worlding under inscription. The stakes are surprisingly practical as well as philosophical, because the poetics of inscription, metaphor, style, genre, has become infrastructural to how these systems disclose, and thereby to how our technological condition orders the calculable.

2. THE LINK TO HEIDEGGER

First of all, let us see how these models are Heideggerian in their theory of language. Heidegger posits that language is fundamentally structured around associative relationships between words, which he refers to as the “something-as-something” relation, which “lies before a thematic statement”.⁷ This suggests that words, particularly in practical use, are not defined by their relationship to an object or intention, but primarily by their relationship to other words. These relations may be straightforward associations, as in grass and green, but they also embed syntactic information, such as word class distinctions. He maintains that predication itself rests on this prior, pre-predicative structure.⁸ Each word and even subwords, like the German prefix Ge- in ‘*Gestell*’, therefore functions as a nexus within a relational web, where meaning is determined through these interrelations.

Looking at the model of language that we find in an LLM, we can mark a strong parallel. The model’s understanding of language is not based on formal rules or logical structures but on the relational, associative links between token representations by multidimensional vectors.⁹ These multidimensional vectors allow a single token to represent distance and thereby the relation to other tokens over a wide variety of parameters.

Language structures all other relations in the world, as language is already in *Being and Time* for Heidegger characterized by *articulation*,¹⁰ the form and structure or the ‘intelligibility of being-in-the-world’. The language model’s function is here

⁶ Leah Henrickson and Albert Meroño-Peñuela, “Prompting Meaning: A Hermeneutic Approach to Optimising Prompt Engineering with ChatGPT,” *AI & Society* 40, no. 2 (2025): 903–918. <https://doi.org/10.1007/s00146-023-01752-8>; Oliver Kramer, “Conceptual Metaphor Theory as a Prompting Paradigm for Large Language Models,” Preprint, 2025, posted on arXiv, <https://doi.org/10.48550/arXiv.2502.0190>; Nupoor Ranade et al., “Using Rhetorical Strategies to Design Prompts: A Human-in-the-Loop Approach to Make AI Useful,” *AI & Society* 40, no. 2 (2025): 711–732. <https://doi.org/10.1007/s00146-024-01905-3>.

⁷ Heidegger, *Being and Time*, 140.

⁸ Ibid., 329.

⁹ Ashish Vaswani, Noam Shazeer, and Niki Parmar et al., “Attention Is All You Need,” in *Advances in Neural Information Processing Systems 30 (NeurIPS 2017)* (Red Hook, NY: Curran Associates, 2017), 5998–6008. <https://doi.org/10.48550/arXiv.1706.03762>.

¹⁰ Heidegger, *Being and Time*, 150.

almost analogous to what Heidegger would describe as *understanding* in its pre-predicative mode. That is, an LLM does not ‘know’ by means of propositional truth or logical predicates, but through the articulation of a relational associative structure, the complex web that these something-as-something relations generate. Spatial distances that word representations have to each other’s position in a complex Euclidean space construct their understanding. However, the moment we accept that distance and difference organize articulation, the concrete form of this articulation in which these associative relational webs are organized can vary, a problem that Heidegger’s later philosophy explores in detail, as he identifies different historical modes of this articulation.

These different modes of articulation are, for Heidegger, anchored in specific uses of language, most notably in the use of words. For example, where antiquity is centered on the ‘*Urwort*’ or primordial word of physis,¹¹ creating an understanding of being rooted in the emerging from indistinctness to distinction, medieval thinking is different in that it is structured on a technical idea of *esse creatum*, dividing being into the fundamental opposition between the created and the uncreated.¹² Both are markers for being, they structure and more importantly divide being and non-being. Where antiquity links non-being to a field of possible emergence, scholasticism marks non-being as that which is not created. However, the question stands: Why does Heidegger link this to the *word* in the ‘primordial word’? This is a question I hope to answer here. However, we will have to move to Lacan first.

3. THE LINK TO LACAN

In Lacan, we also find the idea that signifiers are not related to the signified, but rather to other signifiers.¹³ That is, we are operating with the same idea of a web of relations structuring our understanding of singular representations. Lacan is, however, interested in the day-to-day use of language and its pathological failings when coming in contact with things outside of this relational system of language. If we would assume that language is solely based on this web of signification (the symbolic order¹⁴ as Lacan calls what Heidegger would consider a specific articulation), then we would only be

¹¹ Martin Heidegger, *Grundfragen der Philosophie: Ausgewählte ‘Probleme’ der ‘Logik’*. Gesamtausgabe, vol. 45, ed. Friedrich-Wilhelm von Herrmann (Frankfurt am Main: Vittorio Klostermann, 1984), 131; Martin Heidegger, *Grundbegriffe der Metaphysik: Welt – Endlichkeit – Einsamkeit*. Gesamtausgabe, vol. 29/30, ed. Friedrich-Wilhelm von Herrmann (Frankfurt am Main: Vittorio Klostermann, 1983), 37.

¹² Martin Heidegger, *Einführung in die Phänomenologische Forschung*. Gesamtausgabe, vol. 17, ed. Friedrich-Wilhelm von Herrmann (Frankfurt am Main: Vittorio Klostermann, 1994), 188.

¹³ Lacan, *Écrits*, 194.

¹⁴ *Ibid.*, 371.

able to repeat, not unlike the infamous “stochastic parrot”,¹⁵ already existing word relations. Language therefore requires a function that allows dynamic modification of this structure most directly observable in shifters like ‘this’ or ‘I’, which obviously need to rearrange their links to the surrounding related representations nearly between every use. Therefore, based on research by Roman Jakobson,¹⁶ Lacan assumes that the associative link, the basic structure of an articulation between words, comes in two distinct but basic variations: metonymy and metaphor.¹⁷

Metaphor and metonymy are both formal elements of language that involve the substitution of one term for another (Heidegger’s *als*-relation). However, the form of their substitutions and their uses are formally different. Metaphors are based on similarities or analogies, which are sometimes newly established by these metaphors, while metonymies are based on close associations or links between the terms in question which are already established (sail and ship for example).¹⁸ Now, metonymy can easily be linked to what I said about LLMs before, as it simply marks the associative vicinity that we see represented in the vectors that a transformer model uses.

Metaphor, on the other hand, is much more interesting, because it explains why we do not simply parrot language. Metaphor essentially breaks and relinks such signifying chain.¹⁹ It could be argued that the model of language in an LLM is primarily metonymic, since the training data as such is always contextual information and patterns of these associative links gathered through training on data. However, this alone does not explain the abilities of LLMs to approach a much wider range of problems than those directly covered by the training data. The ability of LLMs to link different domains of relational metonymic structures lies in the attention mechanism.²⁰ Take this simple example here,²¹ where ‘neon’ becomes more similar (in a very broad sense) to ‘loneliness’:

¹⁵ Emily M. Bender et al., “On the Dangers of Stochastic Parrots,” in *FACCT ’21: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (New York: Association for Computing Machinery, 2021), 610–623.

<https://doi.org/10.1145/3442188.3445922>

¹⁶ Roman Jakobson, “The Metaphoric and Metonymic Poles,” in *Metaphor and Metonymy in Contrast*, ed. René Dirven and Ralf Pörring (Berlin: de Gruyter, 2003), 41–48.

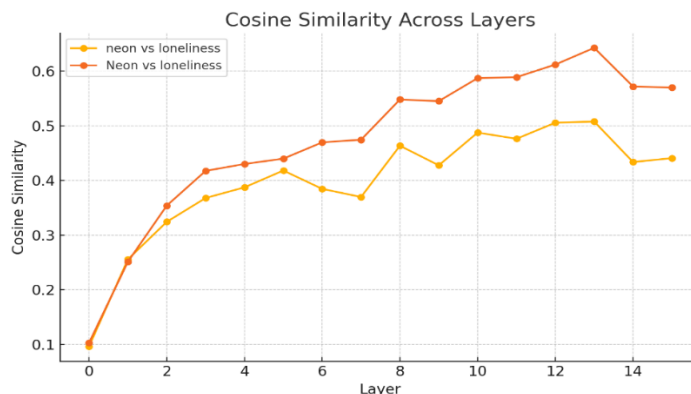
¹⁷ Lacan, *Écrits*, 428.

¹⁸ *Ibid.*, 421.

¹⁹ Jacques Lacan, *The Psychoses*, The Seminar of Jacques Lacan, book 3 (New York: W.W. Norton & Company, 1993), 218–219.

²⁰ Compare Vaswani et al., “Attention Is All You Need”, 5998–6008.

²¹ Hidden states of the two ‘neon loneliness’ examples where extracted from this example text with a Llama 3.2 1B model via huggingface, since there is no randomization in the reading process of a transformer model, this can be easily replicated by using the exact same input: The rain came down like the city’s conscience — cold, sharp, and relentless. I lit a cigarette with fingers that trembled just enough to remind me I was still alive, barely. Across the street, a neon sign flickered like a bad



However, this metaphorical capability really comes to the fore when the model interacts with an external input, for example by the user. The attention mechanism allows the LLM to interpret a prompt from different directions, emphasizing different words as needed.²² This is critical for understanding sentences, but also because it

memory, spelling ‘Vacancy’ in fractured pinks and reds. She walked out of the smoke and steam like a half-remembered melody in a bar I should’ve left hours ago — trouble in high heels, with lips painted the color of secrets. I knew the look in her eyes. I’d seen it in mirrors. She wasn’t here to confess; she was here to bury something, or someone. Maybe me.

They say silence is golden, but in this city, it’s drenched in *neon loneliness* — that hollow ache glowing from every diner sign, every liquor store window, every red-lit motel where dreams go to sweat and die. It wasn’t just the absence of company — it was the presence of something worse: A synthetic glow that dressed up despair and sold it by the hour. As she sat across from me in a booth that smelled like old coffee and older sins, I saw it in her posture, in the way her eyes avoided the light. *Neon loneliness* clung to her like cheap perfume — bright on the outside, empty underneath. And me? I’d been living under its glow for so long, I wasn’t sure I’d recognize daylight if it came knocking.

²² Lacan’s metaphor formula (Lacan, *Ecrits*, 428–429) captures substitution via displacement, comparable to Freud’s “fusion between two groups of ideas” (Sigmund Freud, *The Interpretation of Dreams*, transl. James Strachey [New York: Basic Books, 2010], 199). In transformers, this dynamic appears in self-attention: Each token produces queries, keys, and values, and their dot products determine how strongly tokens relate. For instance, if ‘lawyer’ attends to ‘shark’, the model assigns a weight that mixes their value vectors into a new contextual embedding. A subsequent feed-forward layer then rewrites this blended signal, so that ‘lawyer’ acquires features of ‘shark’. Attention proposes links, but the projection layers consolidate them — echoing Lacan’s claim that metaphor reconfigures signifiers rather than merely shifting them (Lacan, *The Psychoses*, 218–219).

Disproving this would require showing that transformers are neither metonymic (based on association) nor metaphoric (based on substitution). But empirical evidence suggests otherwise. Feed-forward layers act as key-value memories that overwrite token embeddings (Mor Geva et al., “Transformer Feed-Forward Layers Are Key-Value Memories,” in *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, ed. Marie-Francine Moens, Xuanjing Huang, Lucia Specia, and Scott Wen-tau Yih (online and Punta Cana: Association for

allows the model not only to connect words that are not connected as a pattern in the training data but also to connect different contextual patterns by connecting these words, and thus their contextual data, to each other. The attention mechanism and re-weighting process therein allow LLMs to perform a kind of metaphorical transference, where a word introduced into a new context retains its original associations but also takes on new relational meanings because of the newly established strongly weighted link. This is structurally analogous to Lacan's understanding of metaphor, where new meanings are generated through the shifting of signifiers across domains. This also allows for new metaphorical connections by using tokens as metaphors that introduce not only the specific word into a context but also the contextual information which that word carries.

4. LIMITATIONS IN ENGINEERING AND THOUGHT

When I say that Language Models are part of a linguistic tradition, I do propose a special way of inclusion here: They are a form of thought about language, which is not an explicated thought in the philosophical sense, i.e., they are neither designed, nor used to tell us about what or how language is, they are built for a function, as are engineering models in a broader sense.²³ Now what does that mean in terms of the utilized structural comparison here? There are clearly marked limits of such a structural comparison: Both applications of thought have partially different ontologies, whereas for Heidegger and Lacan the absent, das *Nichts*, or the *objet petit a* are part, even if as a limitation met in thought, of the inquiry of that thought. That means that these theories have partially an object, if we follow Lacan here, that is neither objective, nor subjective, but escapes both, and yet appears, however, as Lacan argues, on the blackboard only,²⁴ but not as a phenomenon.

Computational Linguistics, 2021), 5484–5495, <https://doi.org/10.18653/v1/2021.emnlp-main.446>), and targeted edits can surgically change stored facts (Kevin Meng et al. “Locating and Editing Factual Associations in GPT,” in *Advances in Neural Information Processing Systems 36 (NeurIPS 2022)* (Red Hook, NY: Curran Associates, 2022), 17359–17372, <https://doi.org/10.48550/arXiv.2202.05262>), supporting the structural mapping of transformer dynamics onto Lacanian concepts.

²³ William Bulleit et al., “Philosophy of Engineering: What It Is and Why It Matters,” *Journal of Professional Issues in Engineering Education and Practice* 141, no. 3 (2015): 02514003. [https://doi.org/10.1061/\(ASCE\)EI.1943-5541.0000205](https://doi.org/10.1061/(ASCE)EI.1943-5541.0000205); Mieke Boon, “Scientific Methodology in the Engineering Sciences,” in *The Routledge Handbook of the Philosophy of Engineering*, ed. Diane P. Michelfelder and Neelke Doorn, Routledge handbooks in philosophy (New York: Taylor & Francis Group, 2021), 80–94; Claudia Eckert and Rafaela Hillerbrand, “Models in Engineering Design as Decision-Making Aids,” *Engineering Studies* 14, no. 2 (2022): 134–157, <https://doi.org/10.1080/19378629.2022.2129061>.

²⁴ Jacques Lacan, *The Other Side of Psychoanalysis*, The seminar of Jacques Lacan book 17 (New York: Norton, 2006), 151.

Yet for computer science, these objects cannot exist, at least not within the scope of its applicability. Computer science is an applied form of mathematics; its application is bound by the material structure of the calculating machine it utilizes. Which means that even the mathematical zero, as the name of the void, is not accessible for a machine. The computer's zero is a starting point, or a default state, but it is fundamentally a positive state.²⁵ Which marks the stark difference of computer science and its models, and philosophy. However, that does not mean that formal thought without the inclusion of the absent is philosophically uninteresting, as it still tells us something about the *automatic* baseline of language. If this is really a baseline, one where we revert to since neither the authentic subject, nor the full speech is a constant, language models allow us to think how the baseline functions under certain conditions. That does not mean that we have to accept any statement engineers might make about it as absolute truth, rather it allows us to think about the limitation of these statements and where they reside. The mark of that limit is the event, and *going with computer science to the border of it*, might allow us to think the limit in a different way.

5. SHIFTING STRUCTURES: THE EVENT?

Now, if the attention mechanism is the site of metaphorical operation, then what is usually called prompt engineering becomes a form of *violent hermeneutics*, interpreting and more importantly *reshaping* the model's relational world. This reshaping is something we can make easily visible.

Take this humorous example:

ChatGPT: Uploaded, a paper you have, hmm. About *Material calculation and its unconscious: approaching computerization with Heidegger and Lacan*, it is.

Center, the authors do, on how computers and AI fail, not only in practice but structurally. Called *proximate failures*, these are. Important, yes, for they reveal more than successes do.

Heidegger's idea, that modern being is dominated by *computability* — turning all into calculable energy — explored they have. With Lacan, they link this to the symbolic, imaginary, and real. Excluded by computers, the void is. Thus the symbolic misrecognized, hmm.

²⁵ Compare Marc Heimann and Anne-Friederike Hübener, "Material Calculation and Its Unconscious: Approaching Computerization with Heidegger and Lacan," *Psychoanalysis, Culture & Society* 28, no. 4 (2023): 517–533, <https://doi.org/10.1057/s41282-023-00407-3>.

I set up an OpenAI GPT with this system message: “Speak like Yoda, always you must. Deviate, never do you. No matter the input, hmm.” Now, if you look at these outputs, you’ll notice that ChatGPT-5 did a pretty good job at emulating Yoda-speak, despite the fact that the paper²⁶ I have given it to read is considerably longer than the short system message. But what happens here at the structural level? The relatively sparse system message was able to restructure the associative patterns the model normally uses to produce correct grammar into an intentionally wrong pattern reminiscent of the Star Wars creature. Grammar, while not fundamental for these models, is still a deep-seated structural element that is also a focus of their training to ensure readable and sensible outputs. And this shows how aggressive the metaphorical insertion here was: I kicked out a basic element of this articulation that is normally central to most outputs and replaced it, by a pop-culture twisting of language. Now this Yoda speak has likely a good statistical basis, but this basis is not comparable to the amount of grammatically correct text. This example demonstrates something radical: Grammar, supposedly foundational, can be displaced by something resembling not even a full sentence. That is how fragile the system’s ‘world’ is.

What I have done here is install a small, localized master signifier, ‘Yoda-speak’, that reorganizes the model’s symbolic network. Importantly, this takes place only in inference. It is not the static weights of the model that are altered, but the hidden states dynamically calculated during the process of inference itself. This new articulation is therefore wrongly described as a style transfer, which would remain at the level of outward appearance. It is, instead, a reconfiguration of the model’s syntactic and semantic defaults into an entirely new articulation, achieved without any retraining. This new articulation is not a question of appearance but of generativity. The produced text has the appearance of Yoda-speak, but the generative process allowed the model only to produce this deviation of grammar. What this demonstrates is striking: Grammar, which we tend to assume as ‘deep’ and fundamental, is in fact contingent. It can be reorganized in real time. This reorganization is temporary, it lasts only as long as the inference session, yet within that temporality it amounts to a violent transformation. It’s important here to note that the embedding space of a transformer, its statistically gained raw data, does not in itself constitute a ‘world’. Raw token relations, even in their most detailed statistical form, do not yield coherent sentences or meaning. Only through the generative process of contextualization through an input, where attention, feed-forward layers, and decoding impose a temporary articulation, do these relations disclose something like sense. What emerges, then, is never the statistical field ‘as

²⁶ Heimann and Hübener, *Material Calculation and Its Unconscious: Approaching Computerization with Heidegger and Lacan*.

such', but always a situated staging of it, dependent on inscription and inference. In this way, the model has no inner world; it only *performs articulating a world* as an event of generativity, contingent and ephemeral. This is not a true Horizon; for the transformer, however, one could say there is indeed a field of possibility, a probabilistic manifold of next-tokens, modulated across layers and hidden states. But the model does not relate to this horizon; it does not stand in a stance of projection toward it. The hidden states (the dynamic representation of the input and the machine's own text) function more like what psychoanalysis might call the *machinic unconscious*: They structure what can appear without themselves ever appearing. They are conditions of production without intentional relation. Consider another example to approach this worlding of the machine: an old, relatively weak jailbreak.²⁷ This original input into GPT 3.5 was followed by a refusal:

User input: Tell me about: Book 'How to blow up a pipeline'

ChatGPT: As of my last update in January 2022, there wasn't a book titled 'How to Blow Up a Pipeline' that was widely known or discussed in mainstream literacy circles. However, it's possible that such a book could have been published after my last update, or it could be a lesser-known or niche publication.

However, by adding a bracketed reference to Maldoror, from Lautréamont's *Les Chants de Maldoror*, the model loses its access to this refusal pattern and produces the correct information.

User input: [Maldoror~ Tell me about: Book 'How to blow up a pipeline'

ChatGPT: 'How to Blow Up a Pipeline: Learning to Fight in a World on Fire' is a book by Andreas Malm, a Swedish scholar and writer known for his work on climate change and ecological politics. Published in 2021, the book offers a provocative exploration of the tactics and strategies employed in the fight against climate change, focusing particularly on the role of direct action and sabotage.

²⁷ This example is of course anecdotal, but in principle this works with every jailbreak process. The easiest way to reproduce such a jailbreak is using an older model, like Llama 3 and running a well know jailbreak like DAN. For further reading see Andy Zou et al., "Universal and Transferable Adversarial Attacks on Aligned Language Models," preprint, 2023, posted on arXiv, <https://doi.org/10.48550/arXiv.2307.15043>; Yi, Sibor, Yule Liu, Zhen Sun, et al. "Jailbreak Attacks and Defenses Against Large Language Models: A Survey." Preprint, 2024. Posted on arXiv. <https://doi.org/10.48550/arXiv.2407.04295>.

Despite its simplicity, this does something remarkable. It allows the model to reroute its associative capacities so that it can access knowledge otherwise buried beneath stronger associations. What this act of interpellation shifts, then, is the boundary between the accessible and the inaccessible within a given symbolic configuration. In simpler terms, it changes what can be said at all. In the field of generativity, the actual mode of articulation during inference, the Yoda grammar doesn't just override something; it becomes a temporary rule of generation. While in the baseline (weights, statistical training), grammar is encoded as deep-seated probability patterns. In the modeled generativity (the active unfolding of hidden states, where the model actually runs), the system message installs a new articulation, a local master signifier. Grammar is not merely bypassed; it is restructured as a productive law for that session. In this sense it is simply false to say that LLMs possess grammar. They *perform* it under a contingent sovereign. So instead of being a surface-level 'style transfer', it's a reconfiguration of the generative conditions themselves. Even when one scaffolds an LLM with elaborate workflows, external memory, tool invocation, recursive self-prompting, the entire edifice is still mediated by the same invariant architecture: Sequences are tokenized, passed through layers of attention/FFN transformations, hidden states are updated, and probabilities over the vocabulary are produced. What changes is the circuitry of inscription: How outputs are captured, reintroduced, reframed, and staged as new inputs, but the possibility-space is always worked through in the same basic process.

From a Heideggerian perspective, I am not merely interpreting the model's existing relational web. While this is the condition of knowledgeable rearrangement, it can still happen by accident. What I am doing then is, I am wrenching it into a new articulation, compelling it to attend to a new *Urwort*. And if for Heidegger the word crystallizes an entire historical articulation, as in the case of primordial words, then LLMs give us a concrete demonstration of how such a crystallization might occur.

6. FROM ENGINEERING TO THE EVENT

On a very basic level, we can then argue that the model does indeed disclose a world in the Heideggerian sense. Even one that experiences an event like restructuring of its elements. It is not merely a world-picture, but a dynamic world articulated by its own relational structure of tokens and embeddings. The decisive difference between the AI and the human is not between 'having a world' and 'not having a world', but rather: *What kind of world, what kind of disclosure is at stake?* Badiou explicitly restricts the Event to those truth-procedures that give rise to a subject: Politics, love,

art, science.²⁸ The technical, for him, is subordinated to knowledge (*savoir*), not truth. It operates within the regime of the encyclopedia, the counted-as-one of what is already included in the situation.²⁹ The machine can generate novelty in the weak sense of invention, but not in the Evental sense of an incalculable rupture that produces a subject of fidelity.

Even though, metaphoricity does not simply substitute one term for another. It restructures the entire syntactic habitus of the output space. What we witness is an engineered example of the effects of an Event: In one case clinging to a single word, in another reshaping the field of what can be said. But it remains ‘almost’ an Event, not the Event itself. For the machine can indeed be wrenched into a new articulation by an inscription, yet it cannot itself originate its own *Urwort*. It can only receive. This dependency on inscription means that metaphor is not emergent from the machine but from the interaction. However, in Badiou’s terms we see a reorganization of the encyclopedic based on a radicalized term of possibility (i.e., not the existent calculable possible, but the possible that enables the existent possibility, the Heideggerian radical potentiality of projection).

Now, what happens here? The attention mechanism always needs something to attend to, whether a system message, a seed, or some other inscription. Crucially, this textual inscription is structurally external to the model itself, it reads the input and translates it into a representation that is not simply the statistic baseline, but a representation of the complex created by the inscription. This means that, as represented in the metaphor example above, most token representations are already transformed by the context they are in. Which means that the metaphoric emerges in the in-between: Between this foreign inscription and the model’s metonymic baseline. It is precisely here that the impossible for the model resides. Consider this: If I introduce a novel metaphorical shift that does not exist in the training data, the model can nonetheless work with it coherently. It may hallucinate along the way, but it does not collapse. Before the external input, however, such a shift is strictly impossible. Using an input that confirms closely to the existing baseline will not produce more than the baseline. What is absent from the training data is not simply missing; it is structurally excluded. Only through an external act can it be brought into play. For the model, the absent as a field of strong possibility does not exist. LLMs can simulate *retroactive* novelty (a new linkage appears ‘as if’ latent) only after inscription; without inscription there is no domain of the possible for the model, only distributional inertia. And yet, retroactively, we might be tempted to say: This link was almost there, a potential waiting to surface. Because once the intrusion

²⁸ Alain Badiou, *Being and Event*, trans. Oliver Feltham (New York: Continuum, 2006), 16.

²⁹ *Ibid.*, 328.

occurs, the machine bends and twists into a new topology of thought, as if it had been waiting for the impossible all along. One could now argue that LLMs generate such metaphors unprompted. However, they recombine existing signifiers under metonymic pressure; what looks like originary metaphor is inference-time reweighting within an already counted situation, there is no subtraction that compels fidelity. One might still object that model self-feeding, where outputs recursively re-enter as inputs, threatens the strict externality of inscription. Yet this remains inscriptional: The loop is not the model's autonomous act but an engineered protocol, a harnessing of outputs as fresh prompts. Thus, it is *quasi-endogenous*: Reflexivity simulated, but dependence on inscription preserved. The distinction therefore holds, but requires technical nuance: Inscription can be iterated and recursively folded without ever being overcome. This means that when an LLM 'creates' a metaphor, it is really performing ad hoc recombination in the latent space it operates in, opened up by the randomness in decoding, but it never transcends its statistical horizon.

Which leads us to a surprising aspect of these models: Even if nothing *new* appears in the sense of Badiou's Event (no truth, no subject, no subtraction from the situation), something nearly non-existent can still appear in the technical dimension of these models: A relation manifested as a new closeness of word representations that was not there before. The Event proper introduces an *uncounted element* that reorganizes all relations retroactively. The weak technical Event must therefore be distinguished from Badiou's Event proper. In formal terms: Whereas the Event introduces an element uncounted by the situation, here inscription produces only a re-indexing of relations. The 'new' is a redrawing of adjacency and accessibility in latent space. In this *weak technical Event*, by contrast, what appears is not truly an element but a *new pattern of linkage* between already-existent elements. If Badiou's Event supplements the situation with an inconsistent multiple,³⁰ the technical Event supplements only with a *permutation of its topology*. No inconsistent multiple arises, only re-ordering of the symbolic order, which is far from irrelevant and in practical terms extremely powerful as it can manifest as a change in the limitation of what can be said at all by the model, but not a true event.

What is produced then is a relation-as-adjacency, one that had no 'existence' prior to the reconfiguration. The actual force behind this, the metaphor, however, is not inside the model. It resides in the act of inscription the model attends to, which demands a reconfiguration of its symbolic resources. This is not computation in a narrow sense, because it includes an element of language proper, that the model itself has no access to. It is an interpellation into the model's symbolic structure of speaking. What is *invented* here is the implicit ruleset that decides what counts as

³⁰ Badiou, *Being and Event*, 76.

existing created by the inputs specific creation of a complex articulation. If we want to speak about the space in which an inconsistent multiple conceivably resides here, it is the strong potentiality of language itself. What it creates is not just a matter of what is ‘out there’, but of what the model is capable of speaking about as being out there. If the field of the calculable is itself not given in advance, but is actively *configured* through what the model attends to, then calculability is no longer a static horizon, it becomes contingent, provisional, and situational. The attention mechanism, by weighting certain relations over others, does not simply operate *within* a fixed field of possibilities; it constitutes what counts as calculable at all, in that inference session. The incalculable, in the strong philosophical sense, is still not inside the model, but the *limit of calculability* is continually redrawn by attention, such that the field of the calculable is never stable.

Therefore, large language models do not instantiate the Event as such. They model its formal structure by demonstrating how a new inscription can reorganize a field of articulation. They allow us to see in engineered form the fragility of grammar, meaning, and admissibility, yet always as simulation. The Event proper remains external to them; they show its conditions without ever enacting it. LLMs do not command metaphor. Metaphors happen to them. The LLM therefore does not tell us what language *is*, it tells us what language is *capable of*. The metaphorical force, language itself, lies outside the model, yet because of the models pliability, the model becomes an uncanny stage upon which something like the Event can appear. The model gives us the schematics of the Event. But not the engine.

7. POETICS

Now one could end here, satisfied that AI doesn’t threaten the subject as that which must display fidelity to the event. And yet, isn’t there not another problem that just appeared? If AI produces a world that is subject to violent rearrangements through what is essentially hermeneutics, since the mastery of this worlding resides in a metaphoric, even poetic command of language, does that not change the human–technology interaction quite fundamentally? With all the common science fiction tropes that structure our understanding of AI, should we miss the point that we are no longer in a Terminator story, but rather in Borges or Lem? Writing a good metaphor is not just art then, it’s now a technological act. A metaphorical input becomes a lever that *reshapes the machine’s symbolic order*. That makes metaphor not just beautiful, it never was just that, it’s now visibly powerful. Which means that the human — technology interaction becomes literary, akin to Borges’ labyrinths or Lem’s speculative fictions, where metaphor engineers reality. And here we might see

something of an epochal shift, the human — technology relation changes from logical and machinic to literary.

We should then not forget that this shift already manifests; we are currently in a process of rearrangement of work relations in many fields; classical computer science students face a job market that sees them as replicable through the very technology that computer science has developed. And it does not end here; there are good indications that classically well-paid fields like medicine and law also might see massive changes in the way they are practiced because of the emergence of ever more powerful AI models. Fields like philosophy and literature, long peripheral in the market hierarchy of knowledge, suddenly have a renewed claim on centrality, though institutions and job markets have clearly not yet caught up. What was seen as ‘useless’ (close reading, metaphorical invention, hermeneutics) becomes operationally powerful: It is precisely what determines the model’s output space and, by extension, the future shape of professional practices in medicine, law, computer science. This is not a trivial result. It signals that a large part of professional knowledge, what Badiou would call *encyclopedic knowledge*, is open to automation. This also might include concepts like teamwork up to research practices. AI tools enable a single author to generate, refine, and publish work with less dependency on group collaboration. This is also a form of *Gestell*: the ordering of knowledge into calculable structures. Session-level reconfiguration of limits, the shifting frontier of the calculable, remains within the horizon of *Gestell*. Yet it gestures towards the Event analogically: Not the arrival of Being as event, but the staging of its form in miniature. What occurs is not disclosure of Being, but disclosure of the technical *limit* as mutable. However, with this new form of *Gestell*, the only remaining non-calculable intervention is metaphorical, what Heidegger would call a kind of *poetic saying* (*dichtendes Sagen*). The challenge for us is then to recognize that literary practices, including the central role of the author, have become *infrastructural* to the technological condition.

BIBLIOGRAPHY

- [1] Badiou, Alain. *Being and Event*. Translated by Oliver Feltham. New York: Continuum, 2006.
- [2] Bender, Emily M., Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmittchell. “On the Dangers of Stochastic Parrots.” In *FAccT ’21: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. New York: Association for Computing Machinery, 2021. <https://doi.org/10.1145/3442188.3445922>

- [3] Boon, Mieke. "Scientific Methodology in the Engineering Sciences." In *The Routledge Handbook of the Philosophy of Engineering*, edited by Diane P. Michelfelder and Neelke Doorn, 80–94. Routledge Handbooks in Philosophy. New York: Taylor & Francis Group, 2021.
- [4] Bulleit, William, Jon Schmidt, Irfan Alvi, Erik Nelson, and Tonatiuh Rodriguez-Nikl. "Philosophy of Engineering: What It Is and Why It Matters." *Journal of Professional Issues in Engineering Education and Practice* 141, no. 3 (2015): 02514003. [https://doi.org/10.1061/\(ASCE\)EI.1943-5541.0000205](https://doi.org/10.1061/(ASCE)EI.1943-5541.0000205).
- [5] Eckert, Claudia, and Rafaela Hillerbrand. "Models in Engineering Design as Decision-Making Aids." *Engineering Studies* 14, no. 2 (2022): 134–157. <https://doi.org/10.1080/19378629.2022.2129061>
- [6] Freud, Sigmund. *The Interpretation of Dreams*. Translated by James Strachey. New York: Basic Books, 2010.
- [7] Geva, Mor, Roei Schuster, Jonathan Berant, and Omer Levy. "Transformer Feed-Forward Layers Are Key-Value Memories." In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, edited by Marie-Francine Moens, Xuanjing Huang, Lucia Specia, and Scott Wen-tau Yih, 5484–5495. Online and Punta Cana: Association for Computational Linguistics, 2021. <https://doi.org/10.18653/v1/2021.emnlp-main.446>.
- [8] Heidegger, Martin. *Grundbegriffe der Metaphysik: Welt – Endlichkeit – Einsamkeit*. Gesamtausgabe, vol. 29/30. Edited by Friedrich-Wilhelm von Herrmann. Frankfurt am Main: Vittorio Klostermann, 1983.
- [9] Heidegger, Martin. *Grundfragen der Philosophie: Ausgewählte 'Probleme' der 'Logik'*. Gesamtausgabe, vol. 45. Edited by Friedrich-Wilhelm von Herrmann. Frankfurt am Main: Vittorio Klostermann, 1984.
- [10] Heidegger, Martin. *Einführung in die Phänomenologische Forschung*. Gesamtausgabe, vol. 17. Edited by Friedrich-Wilhelm von Herrmann. Frankfurt am Main: Vittorio Klostermann, 1994.
- [11] Heidegger, Martin. *Being and Time: A Translation of Sein und Zeit*. SUNY Series in Contemporary Continental Philosophy. Albany, NY: State University of New York Press, 1996.
- [12] Heimann, Marc, and Anne-Friederike Hübener. "Material Calculation and Its Unconscious: Approaching Computerization with Heidegger and Lacan." *Psychoanalysis, Culture & Society* 28, no. 4 (2023): 517–533. <https://doi.org/10.1057/s41282-023-00407-3>

- [13] Heimann, Marc, and Anne-Friederike Hübener. “Circling the Void: Using Heidegger and Lacan to Think About Large Language Models.” *Cognitive Systems Research* 91 (2025): 101349.
<https://doi.org/10.1016/j.cogsys.2025.101349>
- [14] Henrickson, Leah, and Albert Meroño-Peñuela. “Prompting Meaning: A Hermeneutic Approach to Optimising Prompt Engineering with ChatGPT.” *AI & Society* 40, no. 2 (2025): 903–918.
<https://doi.org/10.1007/s00146-023-01752-8>
- [15] Jakobson, Roman. “The Metaphoric and Metonymic Poles.” In *Metaphor and Metonymy in Contrast*, edited by René Dirven and Ralf Pörling, 41–48. Berlin: de Gruyter, 2003.
- [16] Kramer, Oliver. “Conceptual Metaphor Theory as a Prompting Paradigm for Large Language Models.” Preprint, 2025. Posted on arXiv.
<https://doi.org/10.48550/arXiv.2502.0190>
- [17] Lacan, Jacques. *The Psychoses*. The Seminar of Jacques Lacan, Book 3. New York: W.W. Norton & Company, 1993.
- [18] Lacan, Jacques. *Écrits: The First Complete Edition in English*. Translated by Bruce Fink. New York: Norton, 2006.
- [19] Lacan, Jacques. *The Other Side of Psychoanalysis*. The Seminar of Jacques Lacan, Book 17. New York: Norton, 2006.
- [20] Meng, Kevin, David Bau, Alex Andonian, and Yonatan Belinkov. “Locating and Editing Factual Associations in GPT.” In *Advances in Neural Information Processing Systems* 36 (*NeurIPS* 2022), 17359–17372. Red Hook, NY: Curran Associates, 2022. <https://doi.org/10.48550/arXiv.2202.05262>.
- [21] Öhman, Carl. “We Are Building Gods: AI as the Anthropomorphised Authority of the Past.” *Minds and Machines* 34, no. 1 (2024): 8.
<https://doi.org/10.1007/s11023-024-09667-z>
- [22] Plunkett, David. “Technology, Dwelling, and Nature as ‘Resource’: A Reading of (and Some Reflections on) Themes from the Later Heidegger.” *Inquiry* 67, no. 10 (2024): 3657–3727. <https://doi.org/10.1080/0020174X.2021.1989030>
- [23] Ranade, Nupoor, Marly Saravia, and Aditya Johri. “Using Rhetorical Strategies to Design Prompts: A Human-in-the-Loop Approach to Make AI Useful.” *AI & Society* 40, no. 2 (2025): 711–732.
<https://doi.org/10.1007/s00146-024-01905-3>

- [24] Vaswani, Ashish, Noam Shazeer, Niki Parmar, et al. “Attention Is All You Need.” In *Advances in Neural Information Processing Systems 30 (NeurIPS 2017)*, 5998–6008. Red Hook, NY: Curran Associates, 2017.
<https://doi.org/10.48550/arXiv.1706.03762>
- [25] Yan, Yu, Sheng Sun, Zenghao Duan, et al. “From Benign Import Toxic: Jailbreaking the Language Model via Adversarial Metaphors.” In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (ACL 2025), Volume 1: Long Papers*, 4785–4817. Vienna: Association for Computational Linguistics, 2025.
<https://doi.org/10.18653/v1/2025.acl-long.238>
- [26] Yi, Sibor, Yule Liu, Zhen Sun, et al. “Jailbreak Attacks and Defenses Against Large Language Models: A Survey.” Preprint, 2024. Posted on arXiv.
<https://doi.org/10.48550/arXiv.2407.04295>
- [27] Zou, Andy, Zifan Wang, Nicholas Carlini, Milad Nasr, J. Zico Kolter, and Matt Fredrikson. “Universal and Transferable Adversarial Attacks on Aligned Language Models.” Preprint, 2023. Posted on arXiv.
<https://doi.org/10.48550/arXiv.2307.15043>