

Rhetoric and Generative AI: A Double Articulation

Markus Gottschling
University of Tübingen
markus.gottschling@uni-tuebingen.de

Abstract: This article proposes a rhetorical theory of generative AI usage based on the concept of “double articulation” (Deleuze and Guattari 2005). Amidst the widespread diffusion of Large Language Models (LLMs), rhetorical theory faces a dual crisis of utility and hermeneutics: users struggle to validate the epistemic worth of AI-generated text and to interpret its distributed authorship. Moving beyond the binary of AI as either a neutral tool or an autonomous agent, this article argues that text production with LLMs is a process of reciprocal stratification where human intention (*inventio*) and machinic probability (stochastic generation) are inextricably entangled. By tracing a genealogy of rhetorical generativity – from the *topoi* of Aristotle and the *ars combinatoria* of Ramon Llull to the humanist ideal of *copia* – the study reveals that generative AI does not constitute a rupture but an intensification of rhetorical techné.

The article analyzes how core rhetorical structures (*topoi*, *doxa*, *imitatio*, *copia*) are operationalized in algorithmic environments and argues that successful communication depends on a “Rhetorical AI Literacy.” This literacy recenters the human user not as a mere prompter, but as a critical editor exercising judgment (*iudicium*) to transform machinic probability into situational appropriateness (*aptum*). Ultimately, the article demonstrates that the “human-in-the-loop” is not a temporary fix, but the constitutive horizon of meaning in an age of hybrid textuality.

Keywords: Generative AI, Rhetoric, Double Articulation, Large Language Models, Copia, Imitatio, AI Literacy, Digital Rhetoric.

1 It's not a bug, it's a feature

The growing ubiquity of generative AI¹ systems has consequences that reach far beyond questions of technological application. Generative AI marks a fundamental reordering of the relation between technology and language, with profound implications for rhetoric, in theory and practice, analysis and pedagogy.² A central challenge for rhetorical theory is to account adequately for, and to ground theoretically, the complex interplay between AI-based communication systems and established forms of rhetorical practice.

Across this spectrum, a fundamental problem arises: understanding exactly how generative – and specifically communicative – AI operates remains deeply contested. Because Large Language Models (LLMs) generate statistically likely sequences rather than semantically grounded statements,³ it is unclear whether such models construct implicit “world models” or merely simulate coherence (Vafa et al. 2024). This uncertainty crystallizes around the controversial notion of LLM “hallucination.”

Indeed, “hallucination” functions here not as the description of a technical, but as a rhetorical glitch (Boyle 2015), a paradigmatic metaphor within a broader AI discourse that relies on figurative shortcuts to render opaque systems intelligible. Because the system’s architecture remains a black box, we are forced to reconstruct its operational nature through these discursive negotiations. Acting as epistemic switches, metaphors provide a framing for three distinct modes of interaction: The critique of AI as “snake oil” (Narayanan and Kapoor 2024) or “stochastic parrots” (Bender et al. 2021) generally correlates with resisting the technology (refusal), while the projection of an autonomous “reasoning agent” (Saunders 2025) invites delegating production (passive acceptance). Between these poles opens a

processual perspective: the collaborative constraint (the human-in-the-loop).

From this contested status, two distinct but interrelated crises emerge: a problem of utility (use) and a problem of hermeneutics (understanding). Working through these positions yields a crucial insight: the diffusion and impact of generative AI ultimately hinge on whether we, as its users, choose to remain in the loop by understanding ourselves first as readers. The problem of utility manifests as a dispute over communicative intent and epistemic validity. On one side, pragmatist voices in science and communication argue that scholarly critique must evolve beyond the blanket dismissal that LLMs are “not useful” (Green 2025; cf. Watson et al. 2024), suggesting that these systems can function as valid auxiliaries in rhetorical production. However, this pragmatism faces a rigorous counter-argument from Critical AI perspectives (Bender 2025). For rhetorical theory, such AI critique is potent because it denies the LLM the basic capacity for *logos*: if, as Bender (2025) notes, these models are “nothing more than models of the distribution of word forms,” they fundamentally separate sign from referent. In this view, the machine lacks the intentionality required for a rhetorical act; “hallucination” is therefore not a defect to be fixed, but the inevitable result of a system that mimics the form of language without possessing the cognitive grasp of its meaning. This casts the usefulness of AI not as a practical benefit, but as a threat to the integrity of rhetorical invention (*inventio*) and a deskilling of human writers.

Parallel to this debate on validity is the problem of hermeneutics. Recent scientific discourse has framed the realization that AI outputs require human verification as a novel insight – that texts must be “read and interpreted” before being accepted (Dumit and Roepstorff 2025). For rhetorical and literary theory, however, this “discovery” is merely a vindication of the discipline’s core function. As

¹ Generative AI refers to algorithmically-mediated systems that produce novel content – such as text, images, or code – based on statistical patterns learned from massive datasets during training. This essay focuses on one subclass of generative AI, so-called Large Language Models (LLMs). Trained on sources like books, websites, and whole shadow libraries, LLMs are built on so-called *transformer architectures* that use *attention mechanisms* to weight context dynamically and enable scalable pattern learning (cf. Vaswani et al. 2017).

² For insightful discussions and suggestions regarding these relations and implications, I am grateful to Fabian Erhardt and Jim Ridolfo.

³ During pretraining, masked language modelling tasks require the model to infer missing tokens, strengthening its probabilistic word prediction. When a user enters a prompt, the input text is tokenized – that is, segmented into discrete units – and converted into numerical vectors that the model can process. These vectors pass through multiple hidden layers where contextual relationships are computed. At each step, the model outputs a probability distribution over all possible next tokens; the most probable token is selected, appended, and the process iterates.

humanities scholars have noted, the necessity of close reading and contextual interpretation is not a new requirement imposed by AI, but the foundational premise of criticism (Kirschenbaum 2025; Underwood 2025a; Underwood 2025b). The “hallucinating” machine effectively forces other disciplines to admit what rhetoricians have long maintained: that meaning is never an autonomous property of the text that can be passed on blindly, but is an event that requires the active, critical labor of an audience. The “problem” of the AI text is thus revealed to be a problem of interaction with the text – a failure to recognize and adapt to the rhetorical situation. Indeed, this hermeneutic demand for human interpretation extends across multiple ‘humans-in-the-loop’ within the generative AI pipeline. It involves not only end users who must weigh the implications of AI-generated content, but also the prompters acting as dialogue partners in the co-creation process, and even the evaluators in Reinforcement Learning cycles who, by scoring model outputs, attempt to induce desired behaviors at a foundational technical level.

While technical mitigations attempt to curb stochastic error, they cannot resolve the underlying ontological tension regarding the existence of a coherent world model.⁴ The crucial issue remains the status of the products of generative AI. By severing the traditional link between authorial intention and textual output, LLMs dislodge the ideal of the singular author, installing instead a form of hybrid authorship that relies on texts co-produced by probability engines and human prompts (Gottschling and Kramer 2025; Coeckelbergh and Gunkel 2024).

Within this line of thought, I have argued that generative AI must be understood primarily as an imitation machine (Gottschling 2023) – one that does not produce truth but is parasitic on human intention (Siraganian 2023). However, identifying the ontological status of the machine is only the precursor to the more pressing task: establishing a rigorous framework for its application that resolves the twin crises of use and understanding. To bridge this gap, I propose a theoretical shift: rather than viewing human intent and machinic output as opposing forces, we must grasp their interaction

through the lens of double articulation. Consequently, the aim of this article is to build upon recent semiotic analyses of the Large Language Model as language machine (Weatherby 2025a; Bajohr 2025), moving beyond the mechanics of symbol generation to propose a rhetorical theory of usage that conceptualizes the communicative interaction between the human agent and the LLM-chatbot.

To develop this theory, this contribution advances a perspective that systematically brings the reciprocal entanglement of rhetoric *and* generative AI into view. Drawing on Deleuze and Guattari’s concept of “*double articulation*” (2005, 40) – two interrelated processes of selection and fixation that stratify content and expression into meaning – I sketch a rhetorical theory of generative AI usage. On this account, the linkage of rhetoric *and* generative AI is a process in which expression and content constitute one another as mutual preconditions, with neither pole serving as the origin or the *telos* of the other. Text production as double articulation results in texts permeated by both human and machinic rhetorical production. This is shown through a detailed revisiting of machinic rhetorical structures that come to be important again in the age of generative AI: *topoi* (commonplaces of argumentative invention), *doxa* (shared belief systems), *imitatio* (model-guided variation), and *copia* (resourceful abundance). In concert, rhetoric and generative AI give rise to a new modality of rhetorical practice that breaks with conventional notions of intentional authorship and univocal meaning production.

To lay the groundwork for establishing this perspective of double articulation, the contribution sets out in the next two chapters an account of the reciprocal relation between rhetoricity and generativity. In a first move, it deepens this relation by tracing a genealogy of rhetoric’s generativity, and subsequently analyzing the rhetoricity of generative AI’s architectural and operational structure. This is followed by a detailed examination of the nature and conditions of the double articulation of Rhetoric and generative AI, concluding with an update to core generative-rhetorical structures for the age of AI.

⁴ In the generative-AI pipeline, content quality is secured against “hallucinations” by, among others: *Reinforcement Learning from Human Feedback (RLHF)* – human preference judgments train a reward model that adjusts the base model toward preferred, instruction-following behavior; policy-steering *system prompts* – hidden instructions that set task scope, tone, and safety

constraints; *Retrieval-Augmented Generation (RAG)* – retrieving relevant local or online documents at inference time to ground claims; *Chain-of-thought Reasoning* – structured intermediate prompting steps within the model to reduce errors; and *Post-hoc Evaluation* – verifiers or classifiers that score, calibrate, or block outputs after generation.

2 Rhetoric as a Generative System

While the idea of rhetoric as a technology is not new, theories in digital rhetoric have sharpened this view by treating rhetoric as a procedural, mechanical, and algorithmic practice (Fogg 2003; Bogost 2007; Brown 2014); a perspective whose relevance has become newly urgent in the age of generative AI. A glance back at antiquity reveals that rhetorical education went beyond thematic competence, offering instead operative structures akin to a *techné* – complete with rulebooks, strategies, and modular procedures devised generatively and applied situationally, procedures functionally akin to those driving today’s machine-based text production in LLMs (Brown 2014; Rice 2008).

The structural proximity between rhetoric and generative systems has repeatedly been exploited to challenge common assumptions about rhetoric (Brown 2014; King 2010): on the one hand to defend its mechanical dimension against Plato’s denunciation of rhetoric as mere craft, and on the other to counter the misunderstanding that rhetoric is primarily an art of exceptional, creative invention and intervention. Viewed historically, rhetoric can instead be understood as an iterative procedure designed for re-usability, combination, adaptation, and rulemaking, yet one that requires a form of aptness that Quintilian already defined as a central goal of rhetorical education (*inst. or.*, XI, 1). As Jenny Rice emphasizes, rhetorical practice in the digital age does not simply signal increased efficiency or a surrender to instrumental rationality. Rather, “[b]ecoming a rhetorical mechanic does not call for us to be more efficient ‘on the clock,’ but to figure out ways of caring about the world with the available means of technology” (Rice 2008, 380).

Taking this figure of the ‘rhetorical mechanic’ seriously, the aim of this chapter is to unfold the generative facets of rhetoric from a diachronic perspective. This is done through a short retrospective on those technical, inventive and poetic machines that, over centuries, have negotiated, simulated and expanded the generative potential of rhetorical systems. Such a survey yields an alternative view of the history of text production: a history of recursive procedures, combinatory strategies and connectionist networks that today continue in generative AI applications like ChatGPT and at the same time scarcely conceal their rhetorical provenance. This retrospective reveals a striking continuity: the resurgence of machine rhetoric today gestures not merely toward a new technical

paradigm, but toward the sustained generativity of a discipline that has always mediated between *ars* and *natura*, repetition and variation, rule and exception.

2.1 Rhetoric as Mechanical Procedure: From Antiquity to the Early Modern Period

When the tension between *ars* and *natura* is eclipsed by the kairotic manifestation of *ingenium*, the technical and iterative dimensions of rhetorical practice fade from view. For persuasive force draws no less upon a limited reservoir of means than upon performativity – a reservoir that must be learned, recombined, and repeatedly staged through *imitatio*. For Aristotle, a crucial part of these means is *doxa*, which he recognizes as the decisive basis of rhetorical invention, while the systemic character of the ‘reservoir’ is famously codified in his *topics*, which he describes as a generative “line of inquiry whereby we shall be able to reason from opinions that are generally accepted about every problem propounded to us” (*Topics* I 2, 100a18–20). Commonplaces thereby function as an inventive modular kit that permits argumentative combinations in an almost algorithmic fashion.

It is precisely this view of rhetoric structured as *techné* that Plato had sharply criticized. In the *Gorgias* dialogue, he lets Socrates brand rhetoric as mere culinary spectacle devoid of truth – “rhetoric, which is, in relation to the soul, what cookery is to the body” (*Gorgias* 465e) – and lament that techniques oriented solely toward effect neglect the care for justice. Rather than showing true artistry, the speeches of sophistic *logographoi* thus appear more like the product of a knack (*tribē*) (Hallsby 2024, 233). The counter-position came from Isocrates – Plato’s rival – and from Quintilian, both of whom combined an *imitatio*-based *techné*, in which adopting and varying others’ patterns becomes a condition of eloquence, with situational adaptability. Although still clinging to a need for *natura*, Quintilian describes the orator’s task as a discerning flexibility that negotiates between rule-bound speech and timely deviation (*inst. or.* II, 13). The same linkage is evident in Isocrates, who conceives rhetoric not as an abstract technique but as the outcome of a graduated educational process: talent, formal schooling and practical experience constitute its foundation (Isocrates 1929, 173). Situational awareness, for Isocrates, arises from repetition and variation, which, then, is not an end in itself, but a rehearsal (Papillion 1995). What both describe is a

scaffolding for creative, timely response, resting on a practice of revisiting and recombining what is already known: the situational and iterative recourse to established knowledge and commonplaces. In the context of ancient rhetorical *techné*, this appears as a generative procedure – indeed, as a topological method – for the production of persuasive arguments.

It is this procedural character – the modular recombination of shared knowledge in response to specific situations – that resurfaces across centuries and takes on new form in early attempts to formalize what later came to be symbolic reasoning. In the 13th century, Majorcan theologian and philosopher Ramon Llull developed a combinatorial device that would become foundational for later theories of logic and algorithmic thought. In works such as *Ars brevis* and *Ars generalis ultima*, he devised a system of rotating discs inscribed with divine attributes, logical operators, and grammatical categories. The aim of this argument machine was not to produce logical deduction in the modern sense but to support theological persuasion grounded in metaphysical truth: “Converting the infidels was the task to which he had dedicated his life.” (Fidora and Sierra 2011, 2) Harkening back to Aristotle’s *Topics*, Llull’s system – “the answering of all questions, assuming that one can identify them by name” (Llull 1985, 579) – was fundamentally inventive, relying on “formal notation, using variables and symbols to represent the different logical propositions and operations” (Fidora and Sierra 2011, 2). In contemporary terms, Llull’s system anticipates the logic of generative design: a finite set of truths and rules is used to produce a potentially infinite number of derived propositions. As Sales (2011, 29) puts it, “Llull had just come across the idea of a generative system” – one that, according to Sales, is called a grammar in linguistics and a machine in computer science.

Llull’s influence on philosophy, mathematics, and logic in the early modern period – known as “Lullism” – was considerable (Ramis Barceló 2019; Traninger 2001). In the 17th century, Gottfried Wilhelm Leibniz explicitly identified Llull’s combinatorics as a precursor to his own *characteristica universalis* – a universal symbolic language designed to represent and derive knowledge (Leibniz 2006). Even before Leibniz, lullist elements are already evident in the theoretical frameworks of early modern thinkers such as Nicholas of Cusa, Giordano Bruno, and others across Europe. Yet these elements were repeatedly adapted and transformed, depending on the theological, dialectical, or philosophical context in which they appeared. Llull’s

formalized schema thus serves as a crucial bridge in the genealogical development of symbolic logic, automated inference, and, ultimately, information-theoretical concepts. One branch of this lineage leads – via Leibniz and his calculating machine – to the foundational paradigms of information theory and artificial intelligence (Crossley 2011).

Against this influence, René Descartes, writing in the early seventeenth century, offered a pointed critique of lullist logic. He did not regard its combinatorics as a genuine method of discovery, but rather as a superficial art of verbosity that merely reassembled existing terms without producing new knowledge – rather “to speak injudiciously about those of which one is ignorant, than to learn anything new.” (Descartes 2006, 16) This early suspicion toward automated recombination resonates uncannily with contemporary concerns: the worry that textual generation based on statistical patterns alone lacks the capacity to produce insight. In this light, Descartes already articulates the epistemic tension and hermeneutic problem that continues to shape our assessment of generative systems: What kind of knowledge – if any – emerges from copiousness alone?

Humanist writers of the sixteenth century, by contrast, responded to linguistic wealth not with skepticism but with renewed conceptual investment. In *De utraque copia verborum ac rerum*, originally published in 1512, Erasmus established the notion of *copia* – abundance – as a central rhetorical and pedagogical principle. In this treatise on expressive skill, he explicitly places himself within a classical lineage, invoking ancient authorities such as Quintilian. It is above all the tenth book of Quintilian’s *Institutio oratoria* (*inst. or.* X) that Erasmus draws upon, where rhetorical abundance is not only affirmed but methodically cultivated. This reference was later retrospectively emphasized: in the sixteenth century, the first chapter of *Institutio oratoria*’s tenth book was even renamed *De copia verborum* in explicit alignment with Erasmus’s terminology (Traninger 2020).

Although early modern reformers such as Philipp Melancthon occasionally linked Erasmus’s rhetoric to Llull’s *ars*, Erasmus himself did not actively pursue this lullist genealogy. Instead, he preferred to anchor *copia* in ancient rhetorical tradition and to establish it as a broad philosophical-rhetorical principle, deliberately distancing himself from speculative or esoteric strands of combinatorics. Traninger’s (2020) analysis of the throughlines of abundance thus reveals *copia* as a conceptual

crystallization point in sixteenth-century humanist discourse – a locus where questions of productive speech intersect with the strategic construction of intellectual genealogies. The appeal to classical authority and the reception of lullist logic emerge here not as mutually exclusive, but as overlapping, at times competing, and at times complementary lines of descent.

De copia radicalizes the idea of rhetorical abundance by not merely illustrating it but systematizing it. The treatise famously offers 195 variations of the sentence “Tuae litterae me magnopere delectarunt”, but its more lasting contribution lies in the rule-based procedures it unfolds across two books. Book I is devoted to *copia verborum* and develops techniques for lexical and syntactic variation: synonymic lists, syntactic permutations, and exemplary phrases provide students with ways to experiment with expressive variety. Book II turns to *copia rerum* – the abundance of things – offering a reservoir of themes, narrative *topoi*, and argumentative schemes. Erasmus explicitly addresses his work to students, whom he encourages to treat words, phrases, and ideas as *raw material* to be reworked in response to occasion and situation – producing discourse, as it were, straight from the workshop (Erasmus 2016, I.1.29). Erasmus’s project can be understood as the construction of rhetorical machines – generative devices whose function lies not in a finished product but in the procedural logic they encode: “miniature lesson[s] in computer programming, since the focus is less on what emerges from the machine than on the machine itself – the rules that we might follow in order to generate copia of words and ideas” (Brown 2014, 502).

Commonplace books – personal collections of quotations and excerpts – and comprehensive instructional works such as Christian Weise’s *Politischer Redner* from 1685 continued the tradition of *Lullism* and *copia*. “Commonplaces were, and indeed are, among the most conspicuous pieces” of a sophisticated rhetorical “machinery of verbal production long neglected” (Moss 1996, vi). These sources offered modular, and at times personalized, building blocks of text, reflecting a notion of authorship that was not yet shaped by the Romantic ideal of originality and genius (Gottschling and Kramer 2025). After the Enlightenment, however, rhetoric gradually lost its standing in the European educational canon. As Literature and Poetry displaced it as the primary vessels of authentic expression and national identity, the explicit

mechanics of text production were marginalized. Consequently, these recombinant procedures migrated into the realm of fiction, where they persisted as forms of aesthetic experimentation and rhetorical play – from Laurence Sterne’s *Tristram Shandy* and Jean Paul’s *Flegeljahre* to the literary modernism of Stéphane Mallarmé and his typographic poem *Un coup de dés jamais n’abolira le hasard*.

2.2 Literature as a Recombination Laboratory

A modernized form of rhetorical combinatorics emerged in the twentieth century, manifesting increasingly through literary experimentation. Raymond Queneau’s *Exercices de style* (1947) demonstrate how a short narrative can be retold in ninety-nine stylistic variations. Here, formalization is a driver of creativity – an explicit continuation of Erasmus’s *copia*: rule-based variety becomes the very object of literary practice. Queneau carried this program into the literary group *Oulipo*, which he co-founded in 1960. Combinatorics, permutation, and mathematical constraint became conditions of poetic freedom for the *Oulipians* – pushed to the limit in works such as Queneau’s cut-up sonnet experiment *Cent Mille Millions de Poèmes* (1961) or Georges Perec’s lipogrammatic novel *La Disparition* (1969), which omits the letter “e” entirely. The proliferation of such texts coincided with a series of literary-theoretical debates – from a thought-school now known as poststructuralism – conducted under the rubric of the “death of the author,” (Barthes 1967) which culminated in Derrida’s theory that sought the *deconstruction* of ‘logocentrism’, a “way of theorizing language that gives central importance to the immediate ‘presence’ of the spoken word, thereby rendering writing a secondary and derived technical reproduction, the trace of an ‘absence’” (Gunkel 2023).

Experiments with computer-generated text had already begun in the 1950s – among them Christopher Strachey’s parodic love letters (Schönthaler 2022). However, the computer scientist Theo Lutz is often credited as the founder of a more ambitious form of computer literature. His *Stochastische Texte* (1959), produced on the Zuse Z22, transferred Queneau’s literary principle into algorithmic practice: using just sixteen nouns and sixteen adjectives, Lutz’s program was capable of generating over four million distinct sentences.

Because he drew his vocabulary from *Das Schloss*, Kafka’s unfinished novel fragment, “this project meaningfully joins its arbitrary logical function with its content” (Bertram and Montfort 2024, 346). With this, Lutz – and with him, a whole lineage of poets working with computers as well as computer scientists experimenting with poetry in the decades that followed – posed a foundational question: how do program and text relate when it comes to the automation of writing? This corresponds to the deconstructionist question who is actually responsible for a given text – who can be considered its *author* – when a machine performs large parts of the writing? (Schönthaler 2022) These questions continue to shape generative literature today: instead of asking whether creative agency resides solely in the code, the textual output, or the conceptual design, we might ask if these instances function as mutually constitutive factors within a dynamic semiosis, producing signs through varying degrees of human and machinic entanglement.

From the 2010s onward, as algorithms increasingly shaped the production and curation of digital texts, this logic of rule and variation became strikingly visible in the form of automated social media bots – especially on Twitter. These bots combine predefined or statistically learned text modules, publish the resulting outputs in real time, and feed them directly back into public discourse (Bertram and Montfort 2024, 403–405). In doing so, they extend the tradition of rhetorical *copia* by making recombinant procedures permanently visible – intervening rhetorically in public space without intention, but not without effect. By 2015 at the latest, the emergence of generative AI marked a shift from a recombinatory to a connectionist paradigm – and with it, a qualitative leap: neural networks no longer generate text through rule-based sequencing, but instead model semantic structure statistically (Bajohr 2023). Early experiments with language models – some of them not yet transformer-based – deliberately explored semantic breakdown as an aesthetic strategy, showing how the calculated loss of coherence could produce poetic surprise. Poetry, as Hahn (2019) notes, is particularly well suited to such generative (re)construction – not only because of its frequent brevity, but also due to its syntactic and semantic openness. Accordingly, Bertram and Montfort (2024, 145–283) list poetry as the most frequently used genre in AI experiments, offering subcategories such as haikus, sonnets, and visual poetry.

By the 2020s, co-writing with AI had firmly established itself. Specialized assistance software such as *Sudowrite* is now used for prose generation; in parallel, authors train their own CustomGPTs or combine smaller, fine-tuned models with large-scale language models (Shanahan and Clarke 2023; Becker 2024). Works such as Lillian Yvonne Bertram’s *Travesty Generator* (2019), K. Allado McDowell’s *Pharmako AI* (2020), Hannes Bajohr’s *Berlin, Miami* (2023), or David Link’s *Pandaemonium* (2024) explore – each in its own way – the boundary between machine creativity and human selection. These hybrid texts relocate literary authorship to a twilight zone where originality and automation blur – unsettling established boundaries of intention, expression, and control.

A particularly reflective response to this new paradigm comes from Austrian author Jörg Piringer. In *günstige intelligenz* (2022), he uses the pre-ChatGPT model GPT-3 to expand his own collection of neologisms with machine-generated terms. He then prompts the model to write a poem for each neologism – and, in a move explicitly modeled on Erasmus’s Variations and Queneau’s *Exercices de style*, asks GPT-3 to reformulate one of them, *Bussardschwadron*, as a Wikipedia entry, a Donald Trump tweet, an ode, a curse, and so on. The project highlights the element of surprise inherent in AI-generated text while deliberately preserving the human frame: by exposing prompt design and employing distinct typography, the act of curation remains explicitly visible. In doing so, Piringer situates machine-generated textuality within the sphere of human agency: literature becomes an interface practice in which aesthetic judgment and statistical modeling are intricately intertwined.

A remarkably prescient exploration of themes that are central to today’s debates around LLMs is found in Georges Perec’s radio play *Die Maschine*, originally aired in 1968. The piece anticipates the logic of generative AI by staging a fictional computer program that dissects, permutes, and rewrites Goethe’s *Wanderers Nachtlied* through associative operations. The system – called “Control” – not only produces anthropomorphisms like sighs and strict interruptions but also reveals characteristic glitches and problems that remain familiar today. These include, for instance, ideological slippages that arise from the interplay between source material and generative variation – such as the machine’s attempt to generate a phrase beginning with “heil hi,” (Perec 1972, 25) only to interrupt itself before completing the word. This moment – framed within the

simulation as a recursive glitch devoid of intent, yet signaling a complex co-intentionality between author and apparatus – subtly foregrounds how discursive sediments emerge from the statistical logic of training data. What appears is not an explicit articulation of ideology, but a glitch shaped by cultural residue, probabilistic sequencing, and the system’s own act of self-censorship. In such cases, the aesthetic punchline shifts from narrative coherence to a meta-reflection on algorithmic processes as a symbolic form (Cramer 2011; Kohlby 2020). While “Control” attempts to produce and evaluate poetic quality, the system’s breakdown ultimately points to the inescapability of human judgment. Perec thereby intensifies what rhetoric has long theorized under the concept of *imitatio auctorum* – and what computational procedures now reintroduce into creative writing: an author imitates a machine that imitates an author, all while commenting on the co-creative entanglements of production and authorship between human and machine. The result is a human-machine-human text that, paradoxically, surpasses the supposed “original” – wherever that may be located in the training data (Gottschling 2023).

What may be even more evident to today’s LLM-literate audience than to Perec’s contemporaries is the decisive role of the recipient in reconstructing intertextual references and opaque meanings. This task becomes a particular challenge in the rhetoric of generative AI, where several instances – author, machine, audience – jointly participate in the (re-)construction of textual meaning and thereby also double the process of textual production itself. In Perec’s recursive loop of simulation and meta-commentary, the rhetorical conditionality of even the most modern forms of text generation becomes visible: meaning does not arise from the program’s calculation, but from the interface between text, machine, and reader.

What *Die Maschine* powerfully demonstrates is that machine-based rhetoricity does not constitute a rupture with tradition, but rather an intensification and multiplication of its technical dimension – a conclusion likewise borne out by the diachronic reconstruction offered in this chapter. Whether in the form of *topoi*, Lull’s *ars combinatoria*, Erasmus’s *copia*, or artificial neural networks, the concern has always been with the abundance of operative rules that generate variety and excess – while a wise and thus human judgment secures their situational appropriateness. In this sense, the production rules of rhetorical *techné* are effectively doubled: not only does the human perform *techné* in a quasi-mechanical

way, but the machine, in turn, comes to command both the explicit and implicit writing conventions of its human counterpart.

Consequently, while machinic rhetoric was once relegated to the playful margins of literary experimentation, it has today returned to the center of communicative practice. Far from being obsolete, the constitutive power of rhetorical *techné* – and its dynamic interplay with creative invention – paves the way for understanding generative AI not merely as a tool, but as a rhetorical system in its own right.

3 Generative AI as a Rhetorical System

In fictional texts, the potential of generative AI has already been tested and aesthetically explored as if in a laboratory; however, the development of LLMs increasingly shifts the focus back to their pragmatic and communicative effectiveness. A central problem remains that LLMs work solely reconstructively with their training data and must first be embedded in concrete situations. Wolff (2023) reminds us that what appears as new in AI-generated texts is often less creative or innovative than probabilistically mediocre. Wolff foregrounds a distinctly rhetorical perspective: Could writing with AI actually introduce something new into social discourse? Which modes of imitation give rise to creative breakthroughs? And through what communicative dynamics does novelty take shape? Wolff’s postulate, that AI literature only comes alive through context and paratext, confirms the classical insight that textual meaning is embedded situationally.

An analysis of the rhetoricity of generative AI systems begins precisely at the juncture where writing and the (re-)construction of meaning converge – a crossroads animated by human intention no less than by stochastic procedure. It is within this contact zone that the persuasive force of LLM texts unfolds. Those who critically reflect on this rhetorical superimposition can discern the mechanism, yet its effect persists: formal and aesthetic coherence persuade even when reference to truth remains contingent. Within this dynamic, the rhetoricity of generative AI has two sources. It stems, first, from the persuasive surface of LLM texts and, second, from the way interaction is organized inside the chatbot interface, where semantic plausibility and epistemic indifference fuse into an engaging conversational front – the key to its persuasive power. The following pages therefore take these co-ordinates as a point of departure to examine how plausibility is materialized in text, how it is persuasively consolidated through

interface practices, and how responsibility is assigned under epistemic as well as ethical auspices.

3.1 From Plausibility to Persuasion

LLMs reconstruct human language with a precision that is not only technically impressive but also generates distinctly immediate rhetorical evidence. A mere glance at a chatbot reply often suffices for a user to judge it consistent, readily agreeable, and therefore plausible. What matters is less that the chatbot lays out a fully-fledged argumentative logic – verifying that would require time and expertise – than the instant sense that the text meets basic criteria of textuality and develops its themes along established rhetorical *topoi*. LLM outputs display cohesion, they are coherent and informative; yet at the same time they are open to multiple interpretations and capable of updating commonplace arguments, at least superficially, for specific situations (Bornscheuer 1976). These primarily formal markers form a correlative syntactic *passe-partout* in which the accompanying semantic associations settle with almost no friction. At the same time, they remain ready for further processing: the answers produced by LLM chatbots in response to user prompts circulate at first only within the dialogue window. They are provisional communicative artefacts – ‘possible worlds’ (Doležel 1998) – that become real only once the user accepts them as plausible, works on them further, and shares them (Gottschling and Kramer 2025).

Because users base their judgement chiefly on these formal markers, the chat window functions as a ‘limbic space’ of plausibility – a semantic forecourt where hypotheses are ranked by their momentary evidence before they are subjected to a binary test of true or false (Schlögl 2024, 40). Plausibility here is not mere formal smoothness; it performs an epistemic pre-filter by inviting the recipient to pursue what is said as potentially accurate. Schlögl explicitly links this to Aristotle’s distinction between evident truth and the *pithanon*; in this sense, chatbot plausibility, as the spontaneously ‘agreeable’, is a rhetorical category (Stendel-Günther 2003). LLM chatbots thus operate exactly in this limbic threshold. Their methodological target is not evidence but probability: the model selects that token sequence which promises an optimum in the context of its input history and thereby maximizes rhetorical connectivity. The result is offerings of meaning that can be taken as ‘potentially true’ – regardless of their factual

accuracy. Generative AI is therefore already a rhetorical system, as it converts plausibility into persuasiveness, while never providing epistemic certainty.

These persuasive effects intensify because the LLM simulates intentionality, thereby inviting anthropomorphizing: humans tend to attribute intention wherever coherent language and dialogic structure appear – a tendency Joseph Weizenbaum described as the “ELIZA effect”, and that Hofstadter (1995, 157) later characterized as the habit of reading “far more understanding than is warranted” into computer responses. The readiness to explain such systems through metaphor is itself a product of this anthropomorphizing attribution (Coghlan 2024, Inie et al. 2024). This effect is further amplified by generative AI’s propensity to perform metacognitive illusions (Erhardt in press): when, for instance, the machine claims in dialogue to have freely supplemented a translation because it ‘hermeneutically derived’ the addition from the given text, it is already offering intentional interpretations because of its operational makeup. Starting from raw probability distributions, the model spins coherent language and, at its peak, performs a semblance of metacognition (Johnson et al. 2025). This manifest projection of intentionality first arises in the text itself and fully acquires its persuasive punch when the interface anchors it – visually, dialogically, rhythmically.

Building on this interface anchoring, broader design choices add a further dimension of persuasive staging. Within the chatbot, the interaction is orchestrated by system prompts—invisible, pre-scripted directives that act as the model’s behavioral guide. These instructions influence the flow of conversation, ranging from how the system handles disallowed queries, to when it initiates an internet search, to which follow-up action it proposes to the user. Additionally, the chatbot interface itself is more than a mere display surface: because the system needs to mask perceptible latencies, pulsing circles, wave-like dots and similar chat indicators signal continuous processing activity. These are accompanied by performatively subject-constituting phrases such as “I think ...”, along with user optimization that relies on an empathic tone, interpersonal politeness markers and relationship-building cues; together these cues first realize the LLM text as a dialogue “on equal footing” – transforming plain statistical probability into felt reliability. The more granular the information a user discloses to the chatbot, the more specifically it can

choose linguistic patterns that strengthen situational coherence within the dialogue – and thus, at least plausibly, imitate situational appropriateness in the world outside the chatbot.

The persuasiveness of LLM systems therefore results from a double superimposition: statistical probabilities generate “formal-semiotic” plausibility “independent of any ‘intelligence,’” (Weatherby 2025a, 17) and the interface translates that plausibility into interpersonal credibility. The chatbot frames the output, smooths latency, converts chunked token streams into fluent sentences, proposes follow-up questions, and keeps the conversation flowing. Both layers – textual plausibility and interface credibility – demand critical scrutiny, both individually and in terms of their interplay. At each level, and across their interaction, the chatbot performs what, in a human speaker, we would call a deliberate *dissimulatio artis* transforming a stochastic procedure into an apparently communicative performance. Such metacognitive and metacommunicative illusions are not harmful per se; they can, in Natale’s (2021) term, function as a ‘banal deception’ that lowers the threshold for cooperation and thus increases productive use. The real concern arises when the staging fully occludes the probabilistic foundation of generative AI – when a ‘strong deception’ presents statistical contingency as factual certainty (Umbrello and Natale 2024). Matters are aggravated by the fact that it often remains – and, given the opacity of generative AI, must remain – unclear: which phenomena are genuine ‘emergent properties’ of the system, which stem from developers’ design choices in the interface, and how the two levels influence one another?

Because these levels work in concert, users are prone to cognitive distortions, as the past few years have repeatedly shown. The systems’ capacities are easily overestimated: from conceiving the generation process as reasoning – an implicit metaphorphological leap already discussed above – to assuming that the models exhibit signs of artificial general intelligence (e.g. Tiku 2023; Schaeffer et al. 2023, Bubeck et al. 2023). In actual use, a phenomenon described by novelist Michael Crichton as the Gell-Mann Amnesia effect (Crichton 2002) frequently appears:⁵ users identify serious errors when the chatbot discusses their own area of expertise, yet they consistently

revert to acknowledging its authority when querying fields unfamiliar to them. They effectively forget the system’s demonstrated unreliability, because its responses sound coherent, plausible, human – and therefore persuasive.

Handing off rhetorical production to LLM-chatbots is thus no simple delegation of routine labor, but a precarious outsourcing of cognitive agency. Prompting effectively cedes the classical *officia* of *inventio*, *dispositio*, and *elocutio* to a system whose generative probabilities lie outside our epistemic responsibility, operating instead within an opaque architecture of sedimented, pseudo-agentic forces. What we relinquish, then, is not merely stylistic effort; we transfer the very circulatory motion – iterative inference, revision, and judgement – that constitutes rhetorical thinking, entrusting it to a system that produces coherence not through intent, but through the opaque reconfiguration of sedimented probabilities. The consequence, as proponents of critical AI literacy warn, could be cognitive de-skilling that dulls users’ ability to recognize intertexts, test causalities, or negotiate *topoi*, while replacing these operations with the model’s average-oriented chain of thought (Bastani et al. 2024; Kosmyna et al. 2025). However, it would be equally misguided to conclude from this that interactions with generative AI should be rejected outright. Rather, I believe that rhetoric should take its cue from Hans Magnus Enzensberger and his dictum – formulated with regard to the analysis and use of mass media in connection with the protests of 1968 – that “fear of handling shit is a luxury a sewer-man cannot necessarily afford” (Enzensberger 1970, 18). A rhetorical theory of generative AI must begin with the situated give-and-take between human and model, at the junctures where judgement, approximation, and design converge, for only there can we learn to steer, annotate, and contest what the machine proposes.

3.2 User, Prompt, Dialogue

Because plausibility grants an LLM answer recognition as a communicative partner, the prompt marks the threshold at which the user negotiates their position in the dialogue. The ability to frame complex tasks in ordinary language creates a

⁵ Based on structured observations from AI literacy trainings conducted at the Tübingen Center for Rhetorical Science Communication Research on Artificial Intelligence with rhetoric students, scientists and university staff, as well

as lifelong learners since 2023. While formal empirical validation is forthcoming, the pattern has been consistent across cohorts.

connectivity that allowed ChatGPT to leave earlier voice assistants such as Siri or the Google Assistant far behind. This, in turn, fueled demand for a systematic typology of prompts. Yet the more intensely researchers try to order these mechanisms causally, the clearer their limits become. The models' stochastic mode of production, continual parameter updates, and platform-specific training corpora undercut any definitive systematization (Schulhoff et al. 2024; Gupta and Shivers-McNair 2024). What began as instrumental "prompt engineering" – the quest for precise phrasings – has therefore evolved into a *promptology* (Bajohr 2023): an unfinished yet indispensable body of knowledge about how different input gestures generate different textual worlds. Within this field, prompt formulas and contextual semantic fields have established themselves as a form of rhetorical situational analysis meant to probe the possibilities and limits of text production (Gottschling 2025).

Precisely because the prompt remains malleable – indeed porous – as an interface, the issue of role communication gains weight alongside the situational contexts that guide world-building through external cues supplied by prompts and system roles: how the user positions themselves toward the chatbot and which role they assign to it (Batzner et al. 2024). Designating roles is a double-edged rhetorical technique: it opens finely grained dialogue scenarios while simultaneously stabilizing the user's cognitive comfort zone. When the model is prompted to speak as a legal adviser, historian, or ironic commentator, the probable answer space is pre-structured stylistically, thematically, and persuasively. Yet the result also absorbs dispositions from the training data and the design choices on the developer side that render a model more or less useful for particular outputs. When OpenAI rolled back a spring 2025 update of GPT-4o because of excessive sycophancy, the episode laid bare how susceptible chatbots remain to the tendency toward affirmative mimicry (Stewart 2025). These design choices underscore that conversational continuity – and the explicit affirmation it entails – remains a primary objective for LLM chatbots.

The prompt is inseparably coupled with the chatbot's dialogue function. Steinhoff's concept of *Gebrauchssuggestion* ("suggestion of use") (2025) captures how the graphical user interface draws the writer into a quasi-interaction: the chatbot possesses neither body nor intention, yet it participates in the act of writing by offering suggestions, posing follow-up questions, and proposing interim

formulations. What began with computers as a unilateral user-artefact contact has, as Guzman and Lewis (2020) emphasize, shifted toward a field in which machines must be modelled as autonomous communicators. As noted above, successful prompt practice presupposes a calculated anthropomorphizing – a conscious suspension of disbelief (Umbrello and Natale 2024; Coeckelbergh 2018). A competent user should be aware of the statistical nature of the output and still embrace the fiction of a conversational counterpart, jointly crafting a textual possible world (Gottschling and Kramer 2025).

From this interweaving arises a demand for Rhetorical AI Literacy: responsibility for successful communication is augmented by specific system functions yet ultimately rests on the user's *iudicium* (judgement). Defined by Ueding and Steinbrink (2005) as the capacity to assess and distinguish the means provided by art and language regarding their specific utility for the work at hand, this faculty acts as a resolution to the "problem of utility" introduced at the outset. Through *iudicium*, the focus shifts: we stop arguing if the tool is universally valid and start judging if a specific output is appropriate *in situ*, thereby transforming the abstract crisis of epistemic validity into a concrete act of reflective selection.

To exercise this critical judgment, users must recognize their cognitive biases and stylistic idiosyncrasies, along with the more objective situational conditions and audience constraints – the practical import of the Socratic *know thyself* – and, with equal care, cultivate a comparable understanding of the chatbot itself. Those who ignore the model's stochastic nature and the chatbot's default settings or fail to reflect on training biases risk being deceived by the apparently cooperative text flow. Conversely, insight into the rhetoricity of these systems puts at the rhetorically competent user's disposal a palette of strategic production and evaluation techniques: selecting specific models for specific tasks; chain-of-thought prompting to surface argumentative steps or inject contexts at precise junctures; intentionally distributing roles and hierarchies, for example forbidding the model to alter text or instructing it to ask only questions; and, last but not least, iteratively comparing outputs from different models to break the hermetic loops of language-model reasoning.

Such measures cannot ensure absolute reliability. On the one hand, the systems themselves operate non-linearly and non-recombinatively: they function only because they produce a new – sometimes better,

sometimes worse – path through the neural network to reach a probable result. On the other hand, humans co-determine the quality of the outcome: not only does the prompting strategy shape the result, but so too does what users already know about the subject and how deeply they are willing or able to analytically engage and negotiate with the chatbot's findings.⁶ Instead of them becoming mere prompters, Rhetorical AI Literacy teaches users to seek attaining editorial control, so they can navigate a broad continuum of practice, moving fluidly between employing the system as a heuristic search engine and supplying data, verifying reasoning steps, and editing outputs to remedy the model's gaps (Gottschling 2025). The communicative capacity and persuasive power of LLM chatbots do not absolve the user of their rhetorical duties: when every utterance could potentially be machine-generated, rhetorical *ethos* as a means of ethical persuasion becomes more complex. The communicator's credibility must now be assessed not only by personal character but also by technical provenance and editorial skill. The chat environment itself makes this shift visible. As long as the text remains in the chat window, it has workshop status – a possible world that attains social reality only through copy-and-paste into another medium.

3.3 Co-Intentionalities and the Problem of Authorization

In the dialogical interplay between human and chatbot, LLMs reopen poststructuralist and postmodern debates on writing – notably Derrida's critique of logocentrism – because they distribute responsibility across humans and machines, forcing a rethinking of authorship and its ethical implications notably disclosure and accountability in publication practice (Coeckelbergh and Gunkel 2024; Coeckelbergh and Gunkel 2025; Gottschling and Kramer 2025). A fundamental difficulty arises here: Even though poststructuralism theoretically decoupled text and intent, the figure of the author had largely been re-established in literary practice – anchored by the pragmatics of the book market, legal

liability, and reader expectations. Generative AI, however, renders the latent fragility of this construct virulent again. The specific puzzle of textual agency transcends these juridical and economic stabilizations, revealing that the search for an 'intentional author' is not merely a legal gap, but the radicalization of a longstanding semiotic crisis.

The very architecture of generative AI is palimpsestic: LLM outputs pile human traces upon human traces – training-set curators who decide what counts as "the world"; reinforcement-learning annotators who encode acceptability; prompt authors who steer concrete situations; and audiences whose uptake retroactively ratifies or revises the text. Thus, Kranzberg's (1986) axiom that technology is neither good or bad, nor neutral applies with particular force. Given the distributed pipeline, I argue that AI outputs should be read not as intentional statements but as products of distributed semiosis, structured by co-intentionalities, the layered design and use intentions sedimented in system prompts, RLHF-annotations, interface defaults, and user prompts, that couple human-machine articulations.⁷

Because of these joint, role-differentiated contributions... rhetoric and literary theory are being readmitted to the debate. This marks a striking reversal of the so-called Science Wars, where literary theory was often accused of methodological overextension (Engelmeier 2016) – essentially, of reading too much complexity into plain facts. Under the conditions of textual agency, however, what was once dismissed as 'unscientific' now appears as a necessary methodological toolkit for close reading. This shift also forces a re-evaluation of Roland Barthes (1967). Barthes famously proclaimed the 'death of the author' to celebrate the 'birth of the reader,' framing interpretation as a liberating privilege where the reader creates meaning free from authorial intent. LLMs inaugurate a 'second birth of the reader'—but with a crucial difference. In the age of AI, this position is no longer solely a hermeneutic privilege, as Barthes imagined, but a doubled duty of critical quality control.

Rereading LLM outputs – as advocated by Dumit and Roepstorff (2025) – enacts a poststructural imperative: the emancipation of the

⁶ Such negotiation also comes into play when the user wishes to employ the LLM for purposes not intended by the developer – for example, to obtain instructions for building a bomb. These so-called jailbreak attempts have been analyzed as persuasion processes and underscore the fact that systems operating through language can, above all, be steered linguistically (e.g. Ke et al. 2025; Zeng et al. 2024).

⁷ I use intentionality in a Dennettian, "intentional stance"-based sense: the *as-if* attribution of goals and strategies to a communicating agent – the orator – in order to explain and evaluate a text's production and uptake (cf. Dennett 1998)

reader beyond the authority of the author. This is precisely where rhetorical theory is at home: it recenters the audience as the steering instance of communication (Aristotle, Rh. 1358b) and treats writing as a radically distributed act, “shared in interactions not only among people but also between people and various structures in the environment” (Syverson 1999, 8). Syverson’s ecological theory of composition reminds us that text production has never been an isolated act of single authors. Generative chatbots do not alter that fact in principle; rather, they amplify it. While it may be tempting at first glance to assume that co-authorship extends only to the prompter and the chatbot itself, the following pages will show that authorship is also distributed across the training corpus, the model architecture, RLHF labelling, and the system prompt. This network is where the rhetoric of generative AI reallocates authorship. Consequently, the transition from the user perspective to the question of shared agency and intentionality becomes seamless: what changes is not the principle of distributed agency, but its magnitude. This widening becomes especially tangible when autonomous software agents – task-specific process chains that act on a prompt without human supervision – enter the workflow.

The prospect of fully automated text production – of a chatbot that orchestrates multiple LLM instances and auxiliary software modules to handle world-building and reasoning on the basis of a single prompt, automatically researching sources, arranging arguments, switching stylistic registers, and publishing a situation-appropriate text in the correct channel – strikes some observers as a promise, others as a nightmare. Technically, such process chains can already be implemented with custom agents and nested inner loops (Kapoor et al. 2024; Kellogg 2025); in practice, however, the results still fall short of lofty expectations (Wolfangel 2025). The more complex the subject matter and the more demanding the target format, the more conspicuous the conceptual flat spots, logical leaps, and the “non-world-reference” become, where stylistic congruity fails to align with real-world situations. Even the most advanced reasoning models remain bound to the probabilistic text dispositive. Their “thinking” may be sophisticated compared with early generative systems: chain-of-thought formulations boost textual accuracy, but that precision only escapes the limbic threshold of pre-semantic plausibility when users authorize it – either by releasing the text into a (semi-)public space or by

granting agentic systems the credentials to act on their behalf (as when a user hands over login data so the system can send an email).

Generative text production has nonetheless attained normal-status in everyday and professional communication with student essays, corporate memos, and podcast scripts now routinely being co-authored with LLMs. This gives rise to two interwoven problem fields: quality and intentionality. Pedagogical experiments show that targeted co-agency can markedly raise the quality of average texts, yet – as the preceding analysis has repeatedly shown – LLM-supported writing does not reliably reach the level of competent human authors. The paradigm of co-agency holds that every LLM text is the result of intertwined contributions: the human prompt sets an instruction, the model interprets it, and the human reconstructs, revises, or turns it into a communicative act. This cycle amounts to a double construction of meaning in which probabilistic indeterminacy is not a defect but the motor of sliding authorship. The power and responsibility of authorization, however, remains with the human. Upstream actors – data curators, model architects, platform providers – bear design and data responsibilities, but even where an LLM-driven agent acts without human clicks, the authorizing user is the deployer who provisioned credentials and set the process in motion; the duty of verification remains with that deployer. If AI is integrated into the production process without a wholesale hand-off of tasks, the enduring presence of the “human-in-the-loop” should be understood not as a late-stage safety bolt-on but as the constitutive horizon of meaning itself. Agency, in other words, remains co-intentional: models learn rhetorical strategies from us, deploy them back to us, and rely on us to decide when the conversation is “good enough.” Quality assurance shifts from production to revision: the writer’s role is replaced by a rhetorically trained editor who trims the model’s plausible draft for coherence, factuality, and situational appropriateness and is able to craft, from multiple human and machine possible worlds, the most fitting, plausible, and persuasive text.

The interplay of models and humans can thus be framed in terms of co-intentionalities. These emerge well before the end-user’s prompt: the training corpus is imbued with human opinions, attitudes, interests, and motivations embedded in blog posts, book excerpts, social-media threads, or scholarly articles. When these materials are tokenized, their inherent world-views inscribe syntactic and semantic weights. Developer teams then condense these weights within

the neural network, add system prompts and guardrails, and introduce a human corrective via *RLHF*. All these stipulations flow into every chatbot answer as latent dispositions, forming the system’s baseline stance. At this point, influence through intentional stance shifts from developers to end-users: user prompts contribute not only the immediate production request but also – deliberately or unconsciously – the contextual markers, situational cues, role assignments, and cultural constraints discussed earlier. Each AI text is therefore a hybrid artefact that bundles co-intentionalities no longer discretely attributable, charging the resulting artefact with rhetorical force. The value of such artifacts rests not merely on originality or creativity but on the explicit negotiation of authorization and accountability, of license and liability.

Such negotiation takes place, first, between chatbot and user. Processes akin to human persuasion unfold especially when the co-intentionalities of developers and users stand at cross-purposes. Current and especially future inner-loop agents that publish content without human clicks threaten to bypass this gateway. Yet they, too, merely translate statistical patterns; their “decisions” are no different in kind, only co-intentionally extended and enacted without human oversight. Hence the rhetorical diagnosis tightens: LLMs are neither autonomous authors nor neutral tools. Rather they function as reconstructive engines whose probabilistic outputs remain inert until animated by a specific rhetorical situation – locating the text’s origin not in a single agent, but at the crossroads of stochastic procedure and human intention (Wolff 2023). Consequently, dreams of fully automating communication or externalizing critical thinking ignore the fact that every rhetorical effect ultimately depends on an internal instance of reception to stabilize the output. Whether this collaborative role is filled by a human editor or a recursive feedback loop, there must be an instance that judges, revises, and assigns responsibility – effectively performing an act of structural stabilization. What follows, then, is a shift from the question of agency to a theory of co-production: double articulation, a metatheoretical frame that models how human invention and machinic generation reciprocally stratify expression and content.

4 Double Articulation

What follows from the twofold interrelation of rhetoricity and generativity – namely, first, that

rhetorical production patterns themselves appear as generative; and second, that generative AI gives rise to persuasive structures? Generative AI stands in a lineage of algorithmic text production that runs from combinatory literature and procedural poetry to statistically modeled language processes. These forms of experimental text generation staged, early on and often playfully, the fragility and processual character of written meaning-making. In light of Hillers’s (2023) argument that so-called language models are not semantic systems but statistically operating writing models, we have to assert that contemporary AI texts belong to a rhetorical tradition as well: they may be products of a kind of writing that aims at formal coherence while guaranteeing no stable meaning. In all these cases, however these texts emerge as the configurations of formal rules that laden with co-intentionalities that, in turn, invite the production of meaning by their readers.

The question sharpened provocatively by Foucault (1980) around meaning and authorship – “what matter who’s speaking?” – becomes the maxim of a poststructuralist return under rhetorical auspices: it loosens the bonds among sign, referent, and speaker such that authorship becomes a distributed function of a surplus-producing assemblage of data, model, and reception context. What counts in LLMs not as a bug but as a feature is a genuinely rhetorical perspective on language; one that develops a *techné* of production alongside an analysis of its *techné*-conditions. By circulating signs without a stable referent, LLMs illustrate – recalling Derrida’s aphorism “il n’y a pas de hors-texte” – that, from a rhetorical perspective, there is no outside of the text; rhetorical meaning arises only through the interplay of text and co-intentional contexts. LLMs thus confront their users with the rhetorical task of balancing, in theory and practice, a semiosis that is in principle open-ended.

4.1 Rhetoric and Generative AI

Throughout the argumentation, the tripartite relation of a rhetoric *of*, *with*, and *against* AI is implicitly negotiated. To make this relation explicit: LLM-chatbots do more than generate text. Because they imitate human discourse, they almost automatically provoke discursive and scholarly reactions regarding their own value, risks, and authority. These arguments unfold in an arena where the *rhetoric of AI* (Majdik and Graham 2024) structures the debate, subjecting scholarly scrutiny

to persuasive pro-AI futures (e.g., Amodei 2024), ethically informed counterclaims (e.g. Vallor 2024), nuanced research positions (e.g. Weatherby 2025a), or populist simplifications (e.g. Andreessen 2023). At the same time – in what Majdik and Graham call *rhetoric with AI* – generative AI is reshaping research and praxeology in the humanities disciplines, notably rhetoric, composition, and literary studies, which supply foundational training in writing and professional reading (Wang and Tian 2025). In this perspective, algorithmic procedures such as computational text analysis, topic modeling, and the use of LLMs to analyze rhetorical patterns in discourse take center stage (e.g. Yang and Majdik 2024). A third line of inquiry moves beyond descriptive research to adopt an explicitly critical-normative stance. Efforts in *rhetoric against AI* criticize the effects of algorithmic communication on social contexts, particularly regarding their ethical, social, and political implications (Fox 2024; Goodlad 2025; Sano-Franchini et al. 2025). Thus, all three perspectives address the fact that generative AI inherently produces social consequences and direct effects on public discourse.

However, treating these aspects in isolation is insufficient. Rather, the relation between rhetoric and generative AI appears fundamentally doubled and entangled – which is why rhetoric is particularly well suited to serve as a holistic explanatory approach to the phenomenon of generative AI itself. This suitability has increasingly been recognized in adjacent fields that also study generative AI. Williamson (2024) speaks of a Rhetoric of Machine Learning; persuasive gen-AI structures are being analyzed more frequently in science communication and media psychology (Schäfer 2023, Kessler et al. 2025, Henestroza 2025); and the digital theorist Leif Weatherby concludes his book on *Language Machines* with a chapter that introduces the “return” of rhetoric as a consequence of the striking success and operational logic of LLMs. If, as Weatherby argues – and as the preceding chapter also demonstrated – language as a cultural system is severed from any cognitive core we would want it to reflect so that the coupling of sign and referent grows brittle, then rhetoric could offer a trainable “cognitive-cultural technique” (Weatherby 2025b, 195). From a rhetorical-theoretical vantage point one must add: of course, it was there all along – eclipsed, as it were, both by a concept of authorship and originality that foregrounds poetic creation and by a philosophical conception of truth, each presuming language as the bearer of intrinsic meanings.

Weatherby also underscores the *techné*-character of LLMs when he explains that both generation and computation are production processes that do not initially require meaning:

Generation remediates computation as language. In the general outcry we are currently about how LLMs do not ‘understand’ what they generate, we should perhaps pause to note that computers don’t ‘understand’ computation either. But they do it, as Turing proved. And generation does language, severing the relationship between understanding and generation, but also reintroducing language to the problem of computation. (Weatherby 2025b, 204)

In this sense, much of the debate from the side of Rhetoric *against* AI makes a categorical mistake. A “rhetorical” understanding of generative AI runs counter to the charge that these systems are simply “dumb” or “evil.” The outrage often springs from the illusory expectation of a “superhumanly correct” system – a totality that exists in neither culture nor technology (Weatherby 2025b, 210). AI outputs thus do not appear as moral failure but as a necessary corollary of the insight that our culture is already an unmanageably diverse fabric that is at once rhetorical and computational. The point is not to expect LLMs to produce perfect texts for our world, but first to acknowledge that they can produce *topoi* (Igarashi 2021), “semantic packages” (Weatherby 2025a) or “dumb meaning” (Bajohr 2022). Generative AI produces “surfaces that exhibit meaning-effects in their own right” (Bajohr 2025): formally coherent and rhetorically persuasive yet maintaining only a shallow narrative consistency. Following Bajohr (2024), one might describe this as “surface narration”: The semantic indeterminacy of AI-generated content indicates that its meaning takes shape only through the ongoing interpretation and renegotiation by human agents.

There is, then, a demand for rhetoric that arises from the semantic lacuna in writing with generative AI. It links the technical generation of text to the situational fit of rhetorical communication – hence to a double ascription of meaning: on the one hand by the “sender” of a message – that is, the co-producer acting with the LLM-chatbot – and on the other by the audience as receivers of the co-produced message. Practically speaking, for generative AI, rhetoric functions as a “total program of *cognitive-linguistic training*” (Guillory 2022, 144), a learnable, situation-sensitive procedure for fitting generated text to context. As Weatherby puts it with some bite, Aristotle’s definition of rhetoric as *heuresis* – the discovery of arguments – has become, in the age of

AI, a “large-scale process of fucking around and finding out” (Weatherby 2025b, 210), a hybrid of contingency and generativity. Added to this is what Weatherby (2025a, 38) calls a “double autonomy” of symbolic systems: language and data operate as distinct registers whose relations aim not at unambiguous reference but at hypothetical, often speculative linkages. This openness is not a deficit but a constitutive feature – it generates those “persuasive surfaces” (Kramer and Gottschling 2025) whose meaning arise only in and through reception.

Precisely because LLMs remain structurally opaque while operating in rhetorical situations, they resist codification into a rigid rule set. The model of double articulation developed by Deleuze and Guattari (2005) proves productive here, serving not merely as an analogy, but providing the theoretical frame for the entanglement of rhetoric *and* generative AI. In a move of strategic ambivalence typical of their philosophy, they deploy the binary of content and expression precisely to demonstrate the impossibility of their separation: Content and expression emerge in distinct yet interwoven ways. The modeling of the relation between rhetoric *and* generative AI therefore cannot aim at optimizing a supposedly linear hand-off from human to machine. Such a simplified view ignores the operational resistance of the system described earlier—that dense fabric of statistical biases and latent vectors that fundamentally refracts every human input. Both must instead be understood as generative systems reciprocally attuned to one another: the human initiates a rhetorical – and thus generative – procedure; the machine, for its part, runs a rhetorical procedure of its own. These processes overlap, mirror, and continuously infold each other. There is no either-or here, no clean dichotomy, only an irreducible interference of two generative rhetoricities – or rhetorical generativities – that mutually condition one another.

4.2 Double Articulation: A Theoretical Model of Entanglement

This section proposes double articulation as a model for the amalgamation of rhetoric and generative AI. The term is not operative but perspectival: it affords a theoretical vantage point from which the interfolded practice of rhetoric and generative AI

becomes clearer and more explainable. Gilles Deleuze and Félix Guattari develop “double articulation” in the third plateau of *A Thousand Plateaus* to address the geological problem of stasis (Deleuze and Guattari 2015, 40, Adkins 2015). Through a lecture – delivered by Professor Challenger, a fictional Arthur Conan Doyle character –, they explore the ways in which “intensities are captured and regulated,” producing stratification as a “locking [of] singularities into systems of resonance and redundancy” (Deleuze and Guattari 2005, 40). Put simply, to stratify means to harden fluid potentials into fixed layers. In the context of generative AI, this describes precisely the collapse of a high-dimensional probability distribution (the fluid potential) into a linear, immutable sequence of tokens (the fixed stratum). In their inimitable, eclectic manner, Deleuze and Guattari range across geological, molecular, technical, cultural, and linguistic domains to explain how double articulation governs the formation and layering of stable shapes. Stratification, for Deleuze and Guattari, results from two interrelated processes: a first articulation of content that segments the continuous flow of material – cutting out sediments, syntactic patterns, and thematic *topoi* – and a second articulation that arranges these segments into functional structures, fixing them as stable expression.

Since every articulation is double, there is not an articulation of content and an articulation of expression – the articulation of content is double in its own right and constitutes a relative expression within content; the articulation of expression is also double and constitutes a relative content within expression. For this reason, there exist intermediate states between content and expression, expression and content: the levels, equilibriums, and exchanges through which a stratified system passes. In short, we find forms and substances of content that play the role of expression in relation to other forms and substances, and conversely for expression.
(Deleuze and Guattari 2005, 44)

Crucially, double articulation designates a constant mutual conditioning that knows stable expressions only as relative thickenings – moments where the flow of matter slows down enough to acquire a temporary shape. There is no absolute end-point, only intermediate states of varying density.⁸ This fluidity recalls Deleuze’s concept of the fold, which

⁸ A simple analogy illustrates this: think of building a brick wall. The first articulation molds raw clay into distinct, compressed units – the bricks (content); the second articulation cements and stacks these units into a specific

vertical structure – the functional wall (expression). The brick is a stable expression of the clay, yet simultaneously serves as the raw content for the wall.

visualizes how an inside (content) and an outside (expression) are not separate entities, but a continuous surface pleating itself – like fabric folded over to create a pocket. Yet, while the fold captures the continuity of this fabric, it is the lobster’s double pincer that becomes the emblem of the active mechanism:

There are double pincers everywhere on a stratum; everywhere and in all directions there are double binds and lobsters, a multiplicity of double articulations affecting both expression and content. (Deleuze and Guattari 2005, 45)

When Deleuze and Guattari quip that God is a lobster, the point is that “stratification is always a process of ‘double articulation’” (Adkins 2015, 44). Applied to semiotic artifacts – a central object of their interest and the reason they begin with the pair content/expression (Adkins 2015, 57) – double articulation denotes an organizing act of “translation” that produces the human-scientific “world” (Deleuze and Guattari 2005, 62) as a linguistic overcoding: “translation is the application of a new code on top of an already existing code” (Adkins 2015, 59). Content and expression remain superimposed yet categorically distinct, non-reductively complex: “the stratification of being, the different thickenings [...] are not reducible to one another. [I]t is not possible to reduce everything to [a] linguistic stratum in a radical constructivism” (Adkins 2015, 61).

One need not follow Deleuze and Guattari into the depths of infinite folds, abstract machines, substrata, and deterritorializations to see how productive this overlay is for understanding rhetoric *and* generative AI. They are not equal, yet each contains the other. In their translational coupling they yield a communicative stratum that produces the epistemic “world” as stratification – or, put more simply, as textuality. Their account of a real, if formal, distinction between content and expression in the double articulation of language can, by analogy, illuminate the relation between rhetoric *and* generative AI. The point is clearest when aligned with the architecture of generative systems themselves: LLMs do not operate directly on linguistic signs but on numbers; signs are tokenized and embedded as vectors, and “meaning” appears as relationship between weighted parameters in a high-dimensional numerical space. Linguistic translation then rebinds this numerical production to the domain of language. Double articulation identifies, within one communicative event, a statistical-numerical production and a linguistic-rhetorical production

whose reciprocal constraints yield a “double autonomy of the symbolic order” (Weatherby 2025a, 41) within a human-machine semiosis that differs only formally. The result is an ongoing stratification that recomposes content and expression, form and substance, into ever new texts and other generative artifacts.

Double articulation, accordingly, is not a model of replacement but of reciprocal implication. It enables us to conceive generative AI as a cultural-rhetorical reality machine and rhetoric as an epistemically efficacious form-giving that is newly unfolded through techno-stratificatory processes. In contrast to a linear “hand-off,” the appropriate countermodel is co-production, a doubling of procedures: (1) the human initiates a rhetorically generative system; (2) the system, in turn, runs a rhetorical procedure of its own. Yet, viewed through the lens of double articulation, this generative process proves to be more than a simple retrieval or sequence. The first articulation—the genesis of content—unfolds as a highly mediated field of conflicting forces: The human prompt acts as an impulse injected into the ‘sedimented intentionalities’ of the training data, where it is immediately modulated by the latent co-intentionalities of the system—the opaque weight-orderings of the neural network and the normative interventions of RLHF. What emerges as the ‘substance of content’ is thus the result of a hybrid negotiation, an opaque reconfiguration where user intent and statistical probability are inextricably fused. This machinic procedure generates a rhetorically plausible surface – a ‘possible world’ – that remains a mere proposal until the second articulation of the human reading act re-semantizes it, grounding the statistical probabilities back into social reality (Gottschling and Kramer 2025). Both chatbot and human traverse the stages of rhetorical production and condition one another as they go.

The outcome is not a traceable or cleanly separable output but an amalgam, a mixture, an overlay – terms that mark the simultaneity of fusion and difference. In this double articulation of rhetoric and generative AI as modes of production, an interstice opens in which expression and content, human *inventio* and machinic decomposition, situational fit and stochastic variation converge – without either collapsing into the other. The procedures of production differ formally, yet the result remains “the same thing, one and the same stratified subject” (Deleuze and Guattari 2005, 84). The chains of process complement one another; prying one loose from the other misses the object.

More promising is attention to co-intentionalities that intensify particular formations and function as parallel forces reciprocally shaping each other.

For analysis and evaluation alike, the decisive point is that, by virtue of double articulation, rhetoric and generative AI can no longer be disentangled. Their mutual inscription shifts the analytical plane: communicative artifacts can no longer be treated under a neat subject-object schema. They appear instead as stratified unities in which algorithmic and cultural sediments coexist. The intensity of layering varies – from fully automated texts to hybrids to communicative acts in which generative AI exerts only indirect influence through multiple intermediations. Each system contains the other. Every communicative artifact produced with generative AI sediments rhetorical structures: *topoi* in the training data; evaluations of argument patterns in RLHF annotations; persuasive formats in user prompts and so on. Conversely, contemporary rhetorical practice already bears the imprint of generative AI – from co-authored texts to concise phrasings that fit a situation especially well yet were first coined with AI assistance. Even an anti-AI manifesto such as Sano-Franchini and colleagues’ *Refusing GenAI in Writing Studies* (2025) stratifies generative AI in relation to rhetoric through double articulation. It does not simply stand outside the system; rather, it re-inscribes itself into the discourse, actively generating specific *topoi* of refusal – rhetorical patterns to be revisited in the next chapter – that will inevitably become the training sediment for subsequent models. The logic of double articulation is inescapable here: the very act of refusal becomes raw material (content) for the system’s next articulation (expression) – and vice versa. Thus, the critique feeds the logic it opposes, ensuring that the machine can eventually simulate even its own negation. To borrow Enzensberger’s phrase (1970) once more: the “shit” is not eliminated by critique; on the contrary, it proves most resilient precisely where the rejection is loudest.

5 Rhetorical Mechanisms revisited

In the interstice of rhetorically and algorithmically generated patterns of probability, hybrid patchwork texts emerge that host sedimented rhetorical structures, forms and co-intentionalities. The two systems overlap and overwrite one another; their reciprocal influence thwarts clear attribution. This has consequences for the methodological toolkit of rhetoric: in line with Jenny Edbauer’s rearticulation of the rhetorical situation, rhetoric *and* generative AI

do not differentiate discrete systems in which oratorical instances can be neatly parsed, but a “circulating ecology of effects, enactments, and events” (Edbauer 2005, 9). From such ecological vantage, the double articulation binding rhetoric *and* generative AI is best grasped as a dynamic co-articulation of processes. Over the course of this study several interlinked concepts from the rhetorical repertoire – *topoi*, *doxa*, *imitatio*, *copia* – have already been considered under the sign of generativity. Because they mark points where generativity and rhetoricity converge, it is worth probing them again, under revised auspices, for their potential to make double articulation not only a theoretical perspective but also a usable method for analyzing the production and reception of texts.

5.1 Topoi

From Isocrates, *topos* names a reusable vantage point for invention; with Aristotle, it becomes a systematized inventory – a taxonomy of *loci* for discovering arguments. Seen through the lens of rhetoric *and* generative AI, LLMs operationalize this tradition as automated catalogues of *topoi*: they identify, weight, and recombine commonplaces from training corpora, yielding coherent continuations while leaving deliberative fit and intentional framing to human judgment. According to Schlögl (2024), plausibility arises where utterances repeatedly meet with assent in discursive practice. In this sense, Bornscheuer (1976) locates the persuasive force of *topoi* in four moments: habituality, potentiality, intentionality, and symbolicity. Generative AI particularly activates habitualized structures, because these appear as weighted relations – which words co-occur how often in which contexts and thereby acquire (cor-)relative – or, as Bajohr puts it, “dumb” – meaning in relation to one another (Bajohr 2023). Yet double articulation, understood as a principle of production, also implies that without intentional, situation-specific actualization such structures risk shrinking into mere formulae. Only deliberate user interventions in prompt design, together with critical *ex post* editorial adaptation, render intentionality and symbolicity effective, allowing *topoi* surfaced by generative AI to function as vivid updates tailored to rhetorical context.

Double articulation also enables the emergence of new *topoi*. When AI-generated phrases are taken up across different discursive fields and integrated into human texts, they find their way back into

training corpora and can become established there as recurring commonplaces (Igarashi 2021; Weatherby 2025b). This circular transfer illustrates how semantic shifts arise and how *topoi* evolve – for instance, through the increasing use of novel technical metaphors that first circulate in AI outputs and are then absorbed into specialist vocabularies. A case from the early days of ChatGPT is the verb “delve”: because users encountered it so frequently in their interactions, they began to conjoin two semantic valences of a word otherwise rarely used – first, the familiar sense of digging in to examine a subject in detail, and second, the suggestion that generative AI itself could serve as an instrument for such in-depth inquiry. The term subsequently seeped into human discourse (Ramirez 2025), appearing, for example, far more often in scientific abstracts (Kobak et al. 2024), and it now functions as a topical emblem in debates about LLM-assisted writing. Thus, not only is the word’s meaning updated; the entire argumentative hinterland associated with thorough, deep engagement has been reconfigured.

At the same time, the co-intentionalities at work in generative systems indicate that LLMs do not operate as intentionless generators but under multiple influences – often configured as modular prompts, policies, and interface affordances – many of which can themselves function as *topoi*. Through this networking, models and users together reproduce collective purposes and amplify rhetorical intentions. The retrieval and actualization of **topoi** thus become co-productive acts: both machine and human can assume either role – each can generate probabilistic building blocks and each can select, contextualize, and feed them back into future cycles.

5.2 Doxa

In rhetoric, *doxa* can be understood as the enabling condition for probable and communally shared opinions (Erhardt 2025). Since Aristotle, *doxa* has named the field of the sayable – that reservoir of common knowledge that persuades not through epistemic certainty but through intersubjective plausibility. As the counterpart to *epistēmē*, *doxa* has long had a questionable reputation, especially in philosophical traditions shaped by certain readings of Plato’s *Gorgias*. The “widespread claim that rhetoric was born in conflict between *doxa* and *epistēmē*” is overstated, however, since on closer inspection the two are not in existential conflict and a *doxastic* rhetoric possesses its own epistemology (Bengtson

2024, 37). Still, *doxa*, as immediate acceptability, provides the argumentative resonance chamber within which rhetorical effect becomes possible at all. As a collection, *doxai* function not as a fixed canon of truths but as a mobile ensemble of shared convictions, implicit values, cultural presuppositions, and argumentative *topoi*. Nor is *doxa* merely an ideological obstacle to be overcome; nowadays it is understood as “a condition of intersubjectivity and a constitutive element of any verbal interaction” (Amossy 485). Without *doxa*, rhetoric would have “empty pockets” – no shared opinions, no persuasive means, no commonplaces (Alford 2024, 23).

It is immediately evident that LLMs are stretched between the professed aim of delivering *epistēmē* and the actual capacity to (re-)produce *doxa*. In the usual chatbot configuration, LLMs aim at likely continuations; together with maximum-likelihood training, this leads them to reproduce common linguistic, attitudinal, and interpretive patterns – including stereotypes; RLHF fine-tuning can further amplify this tendency through sycophancy. The resulting texts are often formally correct but stylistically generic and “very flat” (Wolfram 2023). The effect recalls Gustave Flaubert’s *Dictionnaire des idées reçues*, a posthumously realized compendium of received ideas that contribute nothing new to discourse.

In this perspective, *doxa* not only prevents genuine communication, it also hinders individual thinking. Shared ideas appear as uncritical opinions passively absorbed and repeated, eventually leading to alienation; they largely contribute to the subject’s dissolution in language. (Amossy 2002, 468)

As Amossy summarizes this long-standing position, the remedy is creativity through writing: “True literature is supposed to displace conventional perspectives and deconstruct dominant ideological views.” (Amossy 2002, 468) Attending to LLM chatbots in such manner allows one to identify and analyze their parroting of political positions, bias reproduction or stylistic flatness. Yet, as Amossy notes, ideological critique itself draws on established opinions and thus reproduces *doxa*. That is simply how rhetorical opinion-formation works: as variation of the established that yields the situationally new. It is rhetorically salient that generative systems perfect precisely this kind of stylized imitation: LLMs stabilize, vary, or recombine existing *doxic* patterns in unprecedented ways. **Doxa** is not merely represented; in the mode of stochastic iteration, it can also be newly produced. Setting aside a purely

ideological reading of generative systems, attention to doxic elements helps identify and leverage fragments that become active in the “*interdiscourse*” (Amossy 2002, 473) as pieces of “the given” (Alford 2024, 11), thereby remaining open to critique without ignoring the doxic components of prompts or analyses. Those who criticize generative systems must account for the fact “that his or her own *doxa* comes into play whenever he or she tries to spot the others” (Amossy 2002, 485).

At the same time, systems such as ChatGPT, Claude, or Grok are never neutral in the sense of merely reproducing the most probable structures represented in their training data. Owing to the complex bundle of co-intentionalities that follows from the double articulation of rhetoric *and* generative AI, they act as gatekeepers and produce “deceptions” in Umbrello and Natale’s sense (2024): more or less perceptible steering mechanisms that prioritize, exclude, or even newly introduce doxic fragments. Post-hoc evaluations that block the output of illegal, proprietary, or sensitive content are one part of this complex; platforms adjusting alignment objectives to the ideological priors of politically more acceptable positions, is another. The persuasiveness of generative systems thus draws from their capacity to reproduce *doxa*, but also – potentially – to subvert it or to produce *adoxa*: public opinions that remain isolated, unanchored, and socially unmoored (Alford 2024).

Donald Trump’s executive order *Preventing Woke AI in the Federal Government* dramatizes this complex as a struggle over common ground (Stalnaker 2002) surrounding LLMs. Under the order’s stated “obligation not to procure models that sacrifice truthfulness and accuracy to ideological agendas,” (Trump 2025) one can also read the exact inverse of what the text professes: the imposition, via LLMs, of a preferred understanding that becomes *doxa* because it radically excludes disfavored views as LLM chatbots will no longer produce them. When double articulation is curtailed in this way – because the infrastructure disables the mechanisms of user influence – the resulting changes may cut deeply into the epistemic fabric, unsettling not only *doxa* but the very concepts of a shared reality. Collaboration between rhetoric *and* generative AI is productive only under conditions of maximal mutual influence. To understand how such influence generates productive variety rather than ideological stasis, we must turn to the operational principle that mediates between the archive and the new text.

5.3 Imitatio

Against the modern misunderstanding of the term, classical rhetoric understood *imitatio auctorum* not as mere parroting but as a structuring principle of rhetorical formation. As Quintilian and Isocrates teach, speakers do well to orient themselves toward exemplary models – not for imitation as an end in itself, but in order to attain a conjunction of *iudicium* (judgment) and *copia* (abundance) (Quintilian, *Inst. or.* X, 1). The point is not mechanical repetition but the utilized command of variation and situational fit (Quintilian, *Inst. or.* II, 13; Isocrates 1929). Rhetorical imitation is therefore not backward-looking; it is future-oriented, purposive, and differentiating – a rule-guided eclecticism applied to linguistic and semantic variation.

With LLM chatbots, given the flatness of default outputs, the central question becomes how textual production can be steered so that *imitatio* can, in fact, be put to situationally productive use. Above, I described context prompting and interaction as rhetorically grounded forms of user influence; they allow for embedding outputs and generating, steering, and assessing variations. Occasionally, generative systems themselves invite users to choose among options, proposing variations or follow-up operations for refinement. Yet these operations resist systematization; a specific *techné* of prompting is of limited use given opaque model architectures and the unfathomable scope of training data. However, if none of these strategies can solve, in principle, the problem of how *iudicium* and *copia* might be joined, that is because the issue is one of creativity – unsolved not only owing to LLM opacity but also rhetorically, since it evades codification. How human speakers achieve creative leaps – genuine breaks and surprising turns that yield innovation – remains unclear (Knape 2013). Perhaps the closest account, and strikingly consonant with LLMs’ imitative mode of production, is the bee simile in Seneca and, in an imitative variation, Francesco Petrarca (Stackelberg 1956; De Rentiis 1998): those who wish to succeed rhetorically require receptive or hermeneutic creativity to transform what they have read into something new, just as bees gather nectar from many blossoms and deposit it in the hive. Crucially, as Petrarca puts it (2005, 46), the necessary transformation into honey – the movement in the “beehive of the heart” (*alveario cordis*) – is a virtuosic appropriation enabled by human *ingenium* (Assmann 1993), even if the process remains obscure in its particulars.

This opaque passage from diversified quantity to quality also characterizes LLM production. The *ingenium* in the beehive of the heart – seeking to “transform those things they found into something else which was better” (Petrarch 2005, 46) – tends, in the case of LLMs, toward iteration rather than qualitative leap. Proponents of generative systems repeatedly claim such a leap, but the systems themselves have yet to deliver it. In a sense, however, it may not be necessary – a reminder of how unclear the process is for us humans. After all, the significance of *ingenium* is tied to a conception of singular authorship that, as noted above, draws on logocentrism and a Romantic aesthetics of genius – both unsettled by late-twentieth-century deconstruction and, in a different register, by the meaning-effects of LLMs.

That imitative iteration nonetheless generates meaning from verbal abundance is articulated by Cave (1979), drawing on Derrida’s notion of *différance*. The figure of *Imitatio auctorum* is a double-articulated movement around an inaccessible original, permitting us to read the interplay of production and imitation as a reflection on normative rule sets (Kaminski and De Rentiis 1998). It presupposes doxic canon formation: imitation of exemplars shapes not only style but also judgments about what may be said. Yet the imitating text must keep its distance from canonical originals; it may not simply reproduce them. Precisely this structural gap – a *horror vacui* – becomes a productive source of multiplication. In a circling approach to the never fully attainable original, a wealth of linguistic variations unfolds; the original potential for meaning is not replaced but expanded through new connotations. From a lack of access there arises a diversity of newly generated horizons of sense. As Traninger (2020) emphasizes, this principle is already present in Erasmus. It is best described not only as cumulative but as generative: rhetorical tropes and figures function as topoi that are not merely stored but productively transformed, guiding the multiplication of discourse. The model of a self-reflexive, generative rhetoric implicit in *imitatio auctorum* reappears in the operating principle of generative AI systems. In both cases the underlying structure does not simply repeat contents, it provides patterns for proliferation. Tropes, styles, and forms of expression operate as rule-like grids that enable new realizations. Thus, LLM text production, as circling repetition, yields sense – not despite but because of its technological deficit, namely the absence of any recoverable relation to an original. As

the examples of Plato and Descartes already show, the imitative-generative quality of rhetoric has long provoked doubts and anxieties. Today these surface chiefly as a fear of loss of control in a hand-off scenario in which artificial neural networks assume the role of the doxic canon for imitation and humans cede the production of persuasive texts to such systems.

Since Isocrates and Quintilian, rhetorical theory has offered pedagogical strategies to discipline iterative imitation by means of *techné* and to channel it into *aemulatio* — the practice of surpassing one’s models. The distinction between *imitatio* and *aemulatio* marks a decisive point for the rhetorical function of generative systems: whereas *imitatio auctorum* aims at variation as such, *aemulatio* is a relational practice that requires comparison with an original in order to exceed it. Even the simple prompt “in the style of ...” reveals both the potential and the ambivalence of such benchmarking: What, precisely, is being imitated or emulated by the system – tone, syntax, creativity, style, reasoning? And what competencies must users possess to judge an output as *imitatio auctorum*, *aemulatio*, or mere repletion? When studies compare human with machine-generated poems – and find, for instance, that lay readers cannot differentiate ChatGPT from Shakespeare and even prefer AI verses – the result says less about the capacities of generative systems than about a misframed study design that presupposes the evaluative competence of non-expert readers for assessing successful *aemulatio* (Porter and Machery 2024). If we operate from a perspective that generative systems *aemulate*, we risk letting the comparison obscure their performance as double articulation.

This is illustrated by *The Infinite Conversation* between Werner Herzog and Slavoj Žižek, an early instance of *imitatio* through generative AI (Miceli 2022). The ostensibly endless exchange is fine-tuned to their characteristic ways of speaking and writing: the vocal imitation of their acoustic and stylistic idiosyncrasies is convincing; the conversation, fittingly, circles around cinema, Lacanian psychoanalysis, and spirituality; and the jargon lends the utterances an air of depth, further amplified by the hybridization of their stylistic patterns. Yet a peculiar emptiness quickly emerges. The content alternates between the absurd and the redundant, often amounting to linguistically plausible backgrounding – Muzak for the AI age. Unsurprisingly, both Herzog and Žižek rejected the result. Herzog – who enthusiastically lent his “real”

voice to the audiobook version of *I am code*, a collection of AI-generated poems produced with OpenAI’s *code-davinci-002* model (code-davinci-002 2023) – wrote in his book *Die Zukunft der Wahrheit* that *The Infinite Conversation* is soulless,⁹ “nothing but the mimicry of a discourse” (Herzog 2024). Žižek (2022), for his part, felt the AI had turned him into “harmless shit,” missing his lewdness and provocations. Whereas Miceli casts his creation as an illustration of a crisis of trust and misinformation (Miceli 2023), Žižek contends that *The Infinite Conversation* is not “an occasion for a profound analysis of the dangers of such AI products.”

Listening to *The Infinite Conversation* reveals a clear failure of genuine *aemulatio*. Yet, the project’s value lies neither in sophisticated AI critique nor in a faithful reproduction of the thinkers’ intellect. While technically superior simulations are feasible today, it is precisely the imperfections that compel attention. The juxtaposition of glitches, nonsense, and absurdities alongside recognizable stylistic tics illuminates the specific mechanics of *imitatio auctorum*—demonstrating how rhetorical recognition persists even amidst semantic breakdown. Amid dense repetition and iteration, a sentence, a turn of phrase, or a fragment occasionally stands out, enabling the listener to (re-)construct meaning. In its constant variation, *The Infinite Conversation* at least realizes the potential for creativity as surprise and perplexity (Vee 2023). This flash of momentary meaning owes less to the intentional content of genuine Herzog-or-Žižek aphorisms than to the performative abundance of variations. It shows that even in imitative emptiness rhetorical sense can condense momentarily – not through the depth of an argument but through the overlay of rhetorical patterns that cohere only when completed by a listener.

5.4 Between Copia and Copy

The tension within *imitatio auctorum* between rhetorical variation and mechanical repetition condenses, for rhetoric and generative AI, in the pair *copia* and copy. As noted earlier at the seam between the *ars lulliana* and the rationalism of Descartes, the problem is what kind of knowledge – if any – emerges from abundance alone. As a rhetorical term, *copia* names creative plenitude, a strategic readiness for

variation and difference; the copy denotes mere replication. On one side stand the abundance of arguments, elaborate ornament, overwhelming, many-faceted display; on the other the dismay of speaking, in the worst case, about matters one does not truly understand. Historically, *copia* was not only a stylistic device but a mode of knowing (Traninger 2020). In the early modern period, it entailed command of text and language – an enrichment rather than a flat knockoff, a self-authorization, a rhetorical arsenal for future performance. *Copia* was storehouse, possibility, and strategic dispositive at once: stock arguments, stylistic variations, *doxa*, turns of phrase, *topoi* were all gathered in commonplace books not as dead phraseology but as reactivatable potentials. At best, language is not only repeated but transformed.

Against this backdrop, LLMs are best read as a technical instantiation of *copia* – engines of copying and variation: they store, vary, and recombine language patterns. However, they enact a crucial split between operation and effect: While the system technically generates *copia* as an unlimited surplus of variations, it rhetorically remains trapped in the mode of the copy or mechanical replication as long as it lacks human direction. When contexts are missing and variation is not situationally appropriated, the material plenitude of *copia* appears as technical copy. Several interlocking rhetorical risks thus become visible in generative systems, risks that reach to the root of the relation between rhetoric and *techné* – shifting old discussions into the present. Outputs of generative AI often read like the product of a merely functional skill, thin on epistemic depth – more “trickery than a reflection of true artistry or knowledge” (Hallsby 2024, 233). A telling symptom is the recurrent promise by LLM chatbots to produce files or applications they cannot in fact deliver: the performance persuades at the surface while the competence fails underneath. The result is a simulation of rhetoric that is formally robust yet cognitively thin – a *copia* severed from its claim to validity. Hence, a kind of semantic exhaustion forms: AI-generated text feels fluent and accommodating, yet it tends toward redundancy and levels itself through an oscillation between stylistic uniformity and communicative fit (Omizo and Hart-Davidson 2024). Recent studies likewise show that, despite creating apparent stylistic variety, LLMs exhibit significant limits in linguistic competence and in

⁹ [orig. German: „nichts anderes als die Mimikry eines Diskurses“]

producing substantive stylistic diversity and argumentative depth (Mahowald et al. 2024; Guo et al. 2024; Reinhart et al. 2025; Chakrabarty et al. 2024). At scale, such patterns yield a pair of cascading risks: model collapse and knowledge collapse. *Model collapse* is feared when new models are trained predominantly on synthetic data produced by other language models; the resulting loop of semantic self-production may induce irreversible defects and erase the tails of the original content distribution (Shumailov et al. 2024). *Knowledge collapse* denotes a progressive narrowing of the available stock of knowledge, over time and across technical representations, as the use of LLMs becomes widespread – that is, when synthetic *copia* replaces human (Peterson 2025). The central risk, then, is that without targeted rhetorical steering, generative systems begin to circle around themselves – not as a dynamic resonance chamber but as algorithmically closed, ever-fading copy.

Between abundance and empty copy, the dynamics of generative AI can be grasped as a *copious void* (Hallsby 2024): infinite imitations produce a paradoxical surplus of possible meanings that takes on determinate shape only through double articulation – contextual framing and communicative negotiation. What matters is not reproduction as such, but the transformation and adaptation of existing patterns through co-intentionalities: the algorithm supplies raw material and at the same time actively shapes rhetorical form. Through Cave’s analysis of cornucopian writing, the radicality of such a structural shift becomes clear: *copia* crosses the line between *verba* and *res*: sign and referent collapse into one and yield not representation but an “imaginative plenitude” that is itself a form. Generative AI produces “word-things,” textual objects that do not point to a stable world-relation but, as rhetorical figures, display the “double aspect” of language and thought (Cave 1979, 21). So understood, *copia* is not merely a category of textual plenty; it realizes the interlocked doubling of technical and rhetorical articulation. Its productive realization depends on, firstly, whether the rhetorical structures produced by the generative machine are recognized and put to use, and, secondly, whether outputs are recontextualized through rhetorical *iudicium* and *aptum*.

Copia thus contains both: the sedimented rhetoricity of the system and the generative performance of technical language production. Double articulation is orthogonal to the human–machine divide: it does not assign rhetoric to the human and generativity to the machine. Rather, each

side carries both dimensions. What is doubled are form and substance – rhetorical form with generative substance, and generative form with rhetorical substance – and the crossings are mediated by iteration and variation. Each time a user takes an output out of the chatbot interface and puts it to work in a real communicative setting – a memo, an email, a policy draft, a classroom exchange – they make a judgment; that act, in Deleuze and Guattari’s sense, stratifies the double articulation and fixes meaning. Assessing the usability and probative force of an output requires awareness of co-intentionalities and of the conditions of situational appropriateness under which *copia* can become specific meaning. This presupposes a speaker who assumes responsibility and, ideally, achieves a reflexive grasp of the double articulation. To possess *copia* is always to press claims to validity rhetorically; what matters is how, where, and when such structures of validity are concretized. This is what Rhetorical AI Literacy means: Rhetorical judgment, competence in variation, and contextual sensitivity therefore remain key qualifications for working with generative *copia*. With every iteration and every recourse to rhetorical abundance, the question recurs: whose forms dominate, which norms prevail, and who controls the modalities of authorization?

6 Conclusion: Is generative AI inevitable?

At the beginning of the Common Era, in the second book of his *Institutio oratoria*, Quintilian describes the orator’s task as follows:

If the whole of rhetoric could be thus embodied in one compact code, it would be an easy task of little compass: but most rules are liable to be altered by the nature of the case, circumstances of time and place, and by hard necessity itself. Consequently the all-important gift for an orator is a wise adaptability since he is called upon to meet the most varied emergencies. (Quintilian, *inst. or.*, II, 13)

Because rhetoric does not permit clean causal inferences, the burden falls on the orator, who must adapt appropriately to each situation in order to persuade. This has never ceased to be the case, but today – nearly two millennia later – such wise adaptability is newly and acutely demanded by the rise of generative AI.

This contribution has sought to theorize that demand by sketching a model of double articulation between rhetoric *and* generative AI. In double articulation, generative AI builds on, accesses, and

permeates rhetorical principles, (re-)producing and transforming *topoi*, *doxa*, *imitatio*, and *copia* as a plurality of outputs that persuade less by truth than by imaginative reach, rhetorical malleability, and situational integrability. Conversely, generativity – and with it, generative AI – permeates rhetoric, structuring and mechanizing it, elucidating that meaning is tied less to human authors than to the co-intentional, variational production of text itself. The human is not thereby reduced to a mere prompter, second-order observer or post hoc attributor of intention relieved of responsibility. It is precisely *wise adaptability* that machinic generation still lacks. Without judgment and decision, generative text remains in a liminal state. As an orator-in-the-loop, the human remains embedded in the making of meaning. In the context of the double articulation of rhetoric and generative AI, *wise adaptability* names the human capacity to handle the overabundance of generation with poise – balancing selective productivity, critical contextualization, and targeted rhetorical transformation.

This text is not only the result of theoretical work; from the outset it was also meant to realize double articulation in writing practice. Accordingly, I researched, structured, wrote, and translated together with different models and systems – among them *Elicit* for literature search; Google’s *Gemini 2.5 Flash* and *3 Pro*; *Perplexity.ai*; Anthropic’s *Claude Sonnet 3.7* and *4.0*; and, above all, OpenAI’s *ChatGPT* with the models *4o*, *o3*, and the “Thinking” versions of *ChatGPT-5* and *5.1*. A systematic evaluation of the process is not possible here for lack of a consistent self-observation protocol. What remains, however, is the insight that, in the end, far less of the generated material or scaffolding stays in the text than a first glance at the doubly articulated text in the chatbot window might suggest. But although human expertise, critical reading, framing, and situational fit are what give the text its shape, it is ultimately impossible to separate which formulation originated from the AI and what comes from the author. In the end, this is *my* text, which I authorize for publication – even though numerous conscious and unconscious co-intentionalities have flowed in as quotations, thought-moves, and semantic or syntactic fragments.

In this sense, the double articulation of rhetoric and generative AI has been realized in practice as well. We can no longer extricate ourselves – consciously or unconsciously – from the mutual influence of human and machine. The challenge, then, is not to police a rigid border between human

judgment and machinic generation, but to navigate a recursive loop where these functions are constantly exchanged. In this double articulation, agency is radically distributed: the machine not only generates but executes its own statistical judgments, just as the human not only judges but actively generates. Even though Retrieval-Augmented Generation (RAG) and real-time web access can now provide LLMs with immediate contextual data, the specific target situation for which a text is prompted remains largely inaccessible to the model. The interplay of co-intentional presuppositions – both human and machine generated – combined with projective audience design and unarticulated human situational knowledge is simply too intricate to allow for a rhetorically adequate mastery of the situation.

Thus, LLMs remain for the foreseeable future unequipped with situational knowledge and so the pressing task for the human-in-the-loop is to master this fluid interplay, recognizing that rhetorical authority no longer resides in a single subject, but in the successful orchestration of these entangled capacities.

References

- Adkins, B. (2015). Deleuze and Guattari’s A Thousand Plateaus. Edinburgh University Press. <https://doi.org/10.1515/9780748686476-005>
- Alford, C. (2024). Entitled Opinions: Doxa After Digitality. The University of Alabama Press.
- Amodei, D. (2024). Machines of Loving Grace. How AI Could Transform the World for Better. <https://darioamodei.com/machines-of-loving-grace>
- Amossy, R. (2002). How to Do Things with Doxa: Toward an Analysis of Argumentation in Discourse. *Poetics Today*, 23(3), 465–487. <https://doi.org/10.1215/03335372-23-3-465>
- Andreessen, M. (2023, October 16). The Techno-Optimist Manifesto. <https://a16z.com/the-techno-optimist-manifesto/>
- Aristotle. (199 C.E.). *Topics* (W. A. Pickard-Cambridge, Trans.). Alex Catalogue.
- Aristotle. (2007). *On Rhetoric: A Theory of Civic Discourse* (G. A. Kennedy, Trans.; 2. ed). Oxford Univ. Press.
- Assmann, A. (1993). Die bessere Muse. Zur Ästhetik des Inneren bei Sir Philip Sidney. In W. Haug & B. Wachinger (Eds.), *Innovation und Originalität* (pp. 175–195). Niemeyer. <https://doi.org/10.1515/9783110964608.175>
- Bajohr, H. (2023). Dumb Meaning: Machine Learning and Artificial Semantics. IMAGE,

- 37(1), 58–70. <https://doi.org/10.1453/1614-0885-1-2023-15452>
- Bajohr, H. (2024). Cohesion Without Coherence: Artificial Intelligence and Narrative Form. *Transit*, 14(2). <https://doi.org/10.5070/T714264658>
- Bajohr, H. (2025). Surface Reading LLMs: Synthetic Text and its Styles (No. arXiv:2510.22162). arXiv. <https://doi.org/10.48550/arXiv.2510.22162>
- Barthes, R. (1967). *The Death of the Author*. Aspen, 5–6. <https://www.ubu.com/aspen/aspen5and6/threeEssays.html#barthes>
- Bastani, H., Bastani, O., Sungu, A., Ge, H., Kabakcı, Ö., & Mariman, R. (2024). Generative AI Can Harm Learning. SSRN. <https://doi.org/10.2139/ssrn.4895486>
- Batzner, J., Stocker, V., Schmid, S., & Kasneci, G. (2024). GermanPartiesQA: Benchmarking Commercial Large Language Models for Political Bias and Sycophancy (Version 1). arXiv. <https://doi.org/10.48550/ARXIV.2407.18008>
- Becker, J. (2024). Können Chatbots Romane schreiben? Der Einfluss von KI auf kreatives Schreiben und Erzählen: Eine literaturwissenschaftliche Bestandsaufnahme. In G. Schreiber & L. Ohly (Eds.), *KI:Text* (pp. 83–100). De Gruyter. <https://doi.org/10.1515/9783111351490-007>
- Bender, E. M. (2025, April 24). LLMs are nothing more than models of the distribution of the word forms in their training data, with weights modified by post-training to produce somewhat different distributions. Unless your use case requires a model of a distribution of word forms in text, indeed, they suck and aren’t useful. [Quote Thread]. Bluesky. <https://bsky.app/profile/emilymbender.bsky.social/post/3lnl6d2vaxx2y>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? 🦜. *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- Bengtson, E. (2024). *The Epistemology of Rhetoric. Plato, Doxa and Post-Truth*.
- Bertram, L.-Y., & Montfort, N. (Eds.). (2024). *Output: An Anthology of Computer-Generated Text, 1953–2023*. The MIT Press. <https://doi.org/10.7551/mitpress/15249.001.0001>
- Bogost, I. (2007). *Persuasive Games: The Expressive Power of Videogames*. The MIT Press. <https://doi.org/10.7551/mitpress/5334.001.0001>
- Bornscheuer, L. (1976). *Topik: Zur Struktur der gesellschaftlichen Einbildungskraft* (1. Aufl.). Suhrkamp.
- Boyle, C. (2015). The Rhetorical Question Concerning Glitch. *Computers and Composition*, 35, 12–29. <https://doi.org/10.1016/j.compcom.2015.01.003>
- Brown Jr., J. J. (2014). The Machine That Therefore I Am. *Philosophy & Rhetoric*, 47(4), 494–514. <https://doi.org/10.5325/phlrrhet.47.4.0494>
- Bubeck, S., Chandrasekaran, V., Eldan, R., Gehrke, J., Horvitz, E., Kamar, E., Lee, P., Lee, Y. T., Li, Y., Lundberg, S., Nori, H., Palangi, H., Ribeiro, M. T., & Zhang, Y. (2023). Sparks of Artificial General Intelligence: Early experiments with GPT-4 (No. arXiv:2303.12712). arXiv. <http://arxiv.org/abs/2303.12712>
- Cave, T. (1979). *The Cornucopian Text: Problems of Writing in the French Renaissance*. Clarendon Press.
- Chakrabarty, T., Laban, P., Agarwal, D., Muresan, S., & Wu, C.-S. (2024). Art or Artifice? Large Language Models and the False Promise of Creativity (No. arXiv:2309.14556). arXiv. <http://arxiv.org/abs/2309.14556>
- code-davinci-002. (2023). *I Am Code: An Artificial Intelligence Speaks* (B. Katz, J. Morgenthau, & S. Rich, Eds.; First edition). Back Bay Books.
- Coeckelbergh, M. (2018). How to describe and evaluate “deception” phenomena: Recasting the metaphysics, ethics, and politics of ICTs in terms of magic and performance and taking a relational and narrative turn. *Ethics and Information Technology*, 20(2), 71–85. <https://doi.org/10.1007/s10676-017-9441-5>
- Coeckelbergh, M., & Gunkel, D. J. (2024). ChatGPT: Deconstructing the Debate and Moving It Forward. *AI & Society*, 39(5), 2221–2231. <https://doi.org/10.1007/s00146-023-01710-4>
- Coeckelbergh, M., & Gunkel, D. J. (2025). *Communicative AI: A Critical Introduction to Large Language Models*. Polity.
- Coghlan, S. (2024). Anthropomorphizing Machines: Reality or Popular Myth? *Minds and Machines*, 34(3), 25. <https://doi.org/10.1007/s11023-024-09686-w>
- Cramer, F. (2011). Exe.cut(up)able statements: Poetische Kalküle und Phantasmen des selbstausführenden Texts. Brill | Fink. <https://doi.org/10.30965/9783846744994>
- Crichton, M. (2020). Why Speculate? <http://geer.tinho.net/crichton.why.speculate.txt>
- Crossley, J. N. (2011). Ramon Llull’s Contributions to Computer Science. In A. Fidora & C. Sierra (Eds.), *Ramon Llull: From the Ars magna to artificial intelligence* (pp. 39–59). Artificial Intelligence Research Inst., IIIA.
- Deleuze, G., & Guattari, F. (2005). *A Thousand Plateaus: Capitalism and Schizophrenia* (11. print). Univ. of Minnesota Press.
- Dennett, D. C. (1998). *The intentional stance* (7. printing). MIT Press.

- Descartes, R. (2006). A discourse on the method of correctly conducting one's reason and seeking truth in the sciences (I. Maclean, Ed.). Oxford University Press.
- Doležel, L. (2009). *Heterocosmica: Fiction and Possible Worlds*. Johns Hopkins University Press.
- Dumit, J., & Roepstorff, A. (2025). AI Hallucinations Are a Feature of LLM Design, Not a Bug. *Nature*, 639(8053), 38–38. <https://doi.org/10.1038/d41586-025-00662-7>
- Edbauer, J. (2005). Unframing models of public distribution: From rhetorical situation to rhetorical ecologies. *Rhetoric Society Quarterly*, 35(4), 5–24. <https://doi.org/10.1080/02773940509391320>
- Engelmeier, H. (n.d.). Fools Rush In. Ein Besuch bei der Sokal-Affäre, anlässlich ihres 20. Geburtstags. *Merkur. Deutsche Zeitschrift Für Europäisches Denken*, 70(811), 43–56.
- Enzensberger, H. M. (1970). Constituents of a Theory of the Media. *New Left Review*, 64, 13–36.
- Erasmus, D. (2016). *Copia: Foundations of the Abundant Style / De duplici copia verborum ac rerum commentarii duo*. In C. R. Thompson (Ed.), *Literary and Educational Writings*, 1 and 2: Volume 1: *Antibarbari / Parabolae*. Volume 2: *De copia / De ratione studii* (Volume 23-24, pp. 280–660). University of Toronto Press. <https://doi.org/10.3138/9781442676695>
- Erhardt, F. (2025a). Assumable plausibilities: Rhetoric as a doxastic technique of probability. *Argumentation et Analyse Du Discours*, 35. <https://doi.org/10.4000/14yb2>
- Erhardt, F. (2025b). Metacognitive Text Organizastion Semiotic and Rhetorical Agency in LLMs. *Arts and Humanities*. <https://doi.org/10.20944/preprints202512.1177.v1>
- Fidora, A., & Sierra, C. (2011). Preface. In A. Fidora & C. Sierra (Eds.), *Ramon Llull: From the Ars magna to artificial intelligence* (pp. 1–4). Artificial Intelligence Research Inst., IIIA.
- Flaubert, G. (2016). *The Dictionary of Received Ideas* (G. Norminton, Trans.). Alma Books.
- Fogg, B. J. (2003). *Persuasive Technology: Using Computers to Change What We Think and Do*. Morgan Kaufmann Publishers.
- Foucault, M. (2019). What Is an Author? In D. F. Bouchard (Ed.), *Language, Counter-Memory, Practice* (pp. 113–138). Cornell University Press. <https://doi.org/10.1515/9781501741913-007>
- Fox, V. B. (2024). A Librarian Against AI; or, I Think AI Should Leave. <https://drive.google.com/file/d/1i7f9b6vN4Hojz5SHwV-IMDKBDuFBnzNE/view>
- Goodlad, L. M. E. (2025). Humanist in the Loop: Teaching Critical AI Literacies, Episode 1. *Critical AI*, 3(1). <https://doi.org/10.1215/2834703X-11908331>
- Gottschling, M. (2023). Imitationen: Zur Menschlichkeit des Erzählens mit Künstlicher Intelligenz. <https://rhet.ai/2023/12/06/imitationen-zur-menschlichkeit-des-erzaehlens-mit-kuenstlicher-intelligenz/>
- Gottschling, M. (2025). Towards Rhetorical AI Literacy. Presenting a Conceptual Framework. *Argumentation et Analyse Du Discours*, 35. <http://journals.openedition.org/aad/9505>
- Gottschling, M., & Kramer, O. (2025). Persuasive Surfaces and Calculating Machines. A Rhetorical Perspective on Artificial Intelligence. *Global Philosophy*, 35(3), 15. <https://doi.org/10.1007/s10516-025-09748-3>
- Green, H. (2025, April 24). There are a lot of critiques of LLMs that I agree with but “they suck and aren’t useful” doesn’t really hold water. I understand people not using them because of social, economic, and environmental concerns. And I also understand people using them because they can be very useful. Thoughts? [Post]. Bluesky. <https://bsky.app/profile/hankgreen.bsky.social/post/3lnjohdrwf22j>
- Guillory, J. (2022). *Professing Criticism: Essays on the Organization of Literary Study*. University of Chicago Press. <https://doi.org/10.7208/chicago/9780226821313.001.0001>
- Gunkel, D. J. (2023, June 20). Deconstruction to the Rescue. *Outland*. <https://outland.art/chatgpt-post-structuralism/>
- Guo, Y., Shang, G., & Clavel, C. (2024). Benchmarking Linguistic Diversity of Large Language Models (Version 2). *arXiv*. <https://doi.org/10.48550/ARXIV.2412.10271>
- Gupta, A., & Shivers-McNair, A. (2024). “Wayfinding” through the AI wilderness: Mapping rhetorics of ChatGPT prompt writing on X (formerly Twitter) to promote critical AI literacies. *Computers and Composition*, 74, 102882. <https://doi.org/10.1016/j.compcom.2024.102882>
- Guzman, A. L., & Lewis, S. C. (2020). Artificial intelligence and communication: A Human–Machine Communication research agenda. *New Media & Society*, 22(1), 70–86. <https://doi.org/10.1177/1461444819858691>
- Hahn, U. (2019, March 9). Vernunft ist auch eine Herzenssache. Zum Verhältnis von Künstlicher Intelligenz und Literatur. *Frankfurter Allgemeine Zeitung*, 16.
- Hallsby, A. (2024). A Copious Void: Rhetoric as Artificial Intelligence 1.0. *Rhetoric Society Quarterly*, 54(3), 232–246. <https://doi.org/10.1080/02773945.2024.2343265>
- Henestrosa, A. (2025). Perceptions of AI in Science Communication [Dissertation, Eberhard Karls Universität Tübingen]. <https://bibliographie.uni->

- tuebingen.de/xmlui/bitstream/handle/10900/168338/DissertationHenestrosa.pdf?sequence=2
- Hiller, M. (2023). Es gibt keine Sprachmodelle. In D. Giurato, C. Morgenroth, & S. Zanetti (Eds.), *Noten zum „Schreiben“* (pp. 279–285). Brill | Fink.
https://doi.org/10.30965/9783846768433_037
- Hofstadter, D. R. (with Fluid Analogies Research Group). (1995). *Fluid concepts & creative analogies: Computer models of the fundamental mechanisms of thought*. Basic Books.
- Igarashi, Y. (2021). *Machine Writing Is Closer to Literature's History than You Know*. Aeon.
<https://aeon.co/essays/machine-writing-is-closer-to-literatures-history-than-you-know>
- Inie, N., Druga, S., Zukerman, P., & Bender, E. M. (2024). From “AI” to Probabilistic Automation: How Does Anthropomorphization of Technical Systems Descriptions Influence Trust? *The 2024 ACM Conference on Fairness, Accountability, and Transparency*, 2322–2347.
<https://doi.org/10.1145/3630106.3659040>
- Isocrates. (1929). *Against the Sophists* (G. Norlin, Trans.). In *Isocrates* (Vol. 2, pp. 160–180). William Heinemann Ltd.; G.P. Putnam's Sons.
- Johnson, S. G. B., Karimi, A.-H., Bengio, Y., Chater, N., Gerstenberg, T., Larson, K., Levine, S., Mitchell, M., Rahwan, I., Schölkopf, B., & Grossmann, I. (2025). *Imagining and Building Wise Machines: The Centrality of AI Metacognition* (No. arXiv:2411.02478). arXiv.
<https://doi.org/10.48550/arXiv.2411.02478>
- Kaminski, N., & De Rentis, D. (1998). *Imitatio*. In G. Ueding (Ed.), *Historisches Wörterbuch der Rhetorik* (Vol. 4, pp. 235–303). De Gruyter.
<https://www.degruyter.com/database/HWRO/entry/hwro.4.imitatio/html>
- Kapoor, S., Stroebel, B., Siegel, Z. S., Nadgir, N., & Narayanan, A. (2024). *AI Agents That Matter* (No. arXiv:2407.01502). arXiv.
<https://doi.org/10.48550/arXiv.2407.01502>
- Ke, S.-W., Lai, G.-Y., Fang, G.-L., & Kao, H.-Y. (2025). *Iterative Prompting with Persuasion Skills in Jailbreaking Large Language Models* (No. arXiv:2503.20320). arXiv.
<https://doi.org/10.48550/arXiv.2503.20320>
- Kellogg, T. (2025, April 19). *Inner Loop Agents*.
<https://timkellogg.me/blog/2025/04/19/inner-loops>
- Kessler, S. H., Mahl, D., Schäfer, M. S., & Volk, S. C. (2025). *All Eyes on AI: A Roadmap for Science Communication Research in the Age of Artificial Intelligence*.
<https://doi.org/10.5167/UZH-278782>
- King, M. (2010). *Procedural Rhetorics—Rhetoric's Procedures: Rhetorical Peaks and What It Means to Win the Game*. *Currents in Electronic Literacy, Gaming Across the Curriculum*.
https://currents.dwrl.utexas.edu/2010/king_procedural_rhetorics_rhetorics_procedures.html
- Kirschenbaum, M. (2025, March 8). *What would be really good now would be if computer/data science departments established programs in reading and interpreting text. There could be theories of text, and examples of texts from other time periods and other cultures. Above all, students would learn how to look at texts very closely.* [Quote post]. Bluesky.
<https://bsky.app/profile/mkirschenbaum.bsky.social/post/3ljvcowhb6k22>
- Knape, J. (2013a). *Bildtextualität, Narrativität und Pathosformel: Überlegungen zur Bildrhetorik*. In P. Schöner & G. Hübner (Eds.), *Artium conjunctio: Kulturwissenschaft und Frühneuzeitforschung: Aufsätze für Dieter Wuttke* (pp. 297–346). Verlag Valentin Koerner.
- Knape, J. (2013b). *Kreativität*. In J. Knape & A. Litschko (Eds.), *Kreativität: Kommunikation—Wissenschaft—Künste* (pp. 23–40). Weidler.
- Kobak, D., González-Márquez, R., Horvát, E.-Á., & Lause, J. (2025). *Delving into LLM-assisted writing in biomedical publications through excess vocabulary*. *Science Advances*, 11(27), eadt3813.
<https://doi.org/10.1126/sciadv.adt3813>
- Kohlbray, (2020). *Digital Index: Control Poetics in Die Maschine*. *Symploke*, 28(1–2), 83.
<https://doi.org/10.5250/symploke.28.1-2.0083>
- Kosmyna, N., Hauptmann, E., Yuan, Y. T., Situ, J., Liao, X.-H., Beresnitzky, A. V., Braunstein, I., & Maes, P. (2025). *Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task* (No. arXiv:2506.08872). arXiv.
<https://doi.org/10.48550/arXiv.2506.08872>
- Kranzberg, M. (1986). *Technology and History: “Kranzberg's Laws.”* *Technology and Culture*, 27(3), 544. <https://doi.org/10.2307/3105385>
- Leibniz, G. W. (2006). *Philosophische Schriften Band 1: Band 1: 1663-1672*. Akademie Verlag.
<https://doi.org/10.1524/9783050060699>
- Llull, R. (1985). *Ars Brevis* (A. Bonner, Trans.). In A. Bonner (Ed.), *Selected works of Ramon Llull: 1232-1316* (pp. 569–643). Princeton University Press.
- Lutz, T. (1959). *Stochastische Texte. Augenblick*, 4(1), 3–9.
- Mahowald, K., Ivanova, A. A., Blank, I. A., Kanwisher, N., Tenenbaum, J. B., & Fedorenko, E. (2024). *Dissociating language and thought in large language models* (No. arXiv:2301.06627). arXiv.
<http://arxiv.org/abs/2301.06627>
- Majdik, Z. P., & Graham, S. S. (2024). *Rhetoric of/with AI: An Introduction*. *Rhetoric Society Quarterly*, 54(3), 222–231.
<https://doi.org/10.1080/02773945.2024.2343264>
- Miceli, G. (2022). *The Infinite Conversation*.
<https://www.infiniteconversation.com/>
- Miceli, G. (2023, April 1). *What an Endless Conversation with Werner Herzog Can Teach Us about AI*. *Scientific American*.
<https://www.scientificamerican.com/article/wh>

- at-an-endless-conversation-with-werner-herzog-can-teach-us-about-ai/
- Moss, A. (1996). *Printed Commonplace-Books and the Structuring of Renaissance Thought*. Clarendon.
<https://doi.org/10.1093/acprof:oso/9780198159087.001.0001>
- Narayanan, A., & Kapoor, S. (2024). *AI Snake Oil: What Artificial Intelligence Can Do, what It Can't, and How to Tell the Difference*. Princeton University Press.
- Natale, S. (2021). *Deceitful Media: Artificial Intelligence and Social Life after the Turing Test* (1st ed.). Oxford University Press New York.
<https://doi.org/10.1093/oso/9780190080365.001.0001>
- Omizo, R., & Hart-Davidson, B. (2024). Is Genre Enough? A Theory of Genre Signaling as Generative AI Rhetoric. *Rhetoric Society Quarterly*, 54(3), 272–285.
<https://doi.org/10.1080/02773945.2024.2343615>
- Papillion, T. (1995). Isocrates' *techné* and Rhetorical Pedagogy. *Rhetoric Society Quarterly*, 25(1–4), 149–163.
<https://doi.org/10.1080/02773949509391038>
- Perec, G. (1972). *Die Maschne*. Reclam.
- Peterson, A. J. (2025). AI and the Problem of Knowledge Collapse. *AI & SOCIETY*, 40(5), 3249–3269. <https://doi.org/10.1007/s00146-024-02173-x>
- Petrarch, F. (2005). *Letters on Familiar Matters 1: Rerum Familiarium Libri* (A. S. Bernardo, Trans.). Italica Press.
- Piringer, J. (2022). *Günstige Intelligenz*. Ritter Verlag.
- Plato. (2020). *Gorgias* (B. Jowett, Trans.). Arctururs.
- Porter, B., & Machery, E. (2024). AI-generated Poetry Is Indistinguishable from Human-written Poetry and Is Rated more Favorably. *Scientific Reports*, 14(1), 26133.
<https://doi.org/10.1038/s41598-024-76900-1>
- Quintilian. (2024). *Institutio Oratoria* (H. E. Butler, Trans.). Lyceum Institute.
- Ramirez, V. B. (n.d.). ChatGPT Is Changing the Words We Use in Conversation. *Scientific American*. Retrieved November 21, 2025, from <https://www.scientificamerican.com/article/chatgpt-is-changing-the-words-we-use-in-conversation/>
- Ramis Barceló, R. (2019). Academic Lullism from the Fourteenth to the Eighteenth Century. In A. M. Austin, M. D. Johnston, & A. Ibarz (Eds.), *A Companion to Ramon Llull and Lullism* (pp. 437–470). BRILL.
<https://doi.org/10.1163/9789004379671>
- Reinhart, A., Markey, B., Laudénbach, M., Pantusen, K., Yurko, R., Weinberg, G., & Brown, D. W. (2025). Do LLMs Write Like Humans? Variation in Grammatical and Rhetorical Styles. *Proceedings of the National Academy of Sciences*, 122(8), e2422455122.
<https://doi.org/10.1073/pnas.2422455122>
- Rice, J. E. (2008). *Rhetoric's Mechanics: Retooling the Equipment of Writing Production*. *College Composition & Communication*, 60(2), 366–387.
<https://doi.org/10.58680/coc20086870>
- Sales, T. (2011). Llull as Computer Scientist, or Why Llull Was One of Us. In A. Fidora & C. Sierra (Eds.), *Ramon Llull: From the Ars magna to artificial intelligence* (pp. 25–37). Artificial Intelligence Research Inst., IIIA.
- Sano-Franchini, J., McIntyre, M., & Fernandes, M. (2025). Refusing GenAI in Writing Studies. *Refusing Generative AI in Writing Studies*. Retrieved August 18, 2025, from <https://refusal.blog/wp-content/uploads/2025/12/refusing-genai-in-writing-studies.pdf>
- Schaeffer, R., Miranda, B., & Koyejo, S. (2023). Are Emergent Abilities of Large Language Models a Mirage? (No. arXiv:2304.15004). arXiv.
<http://arxiv.org/abs/2304.15004>
- Schäfer, M. S. (2023). The Notorious GPT: Science communication in the age of artificial intelligence. *Journal of Science Communication*, 22(02). <https://doi.org/10.22323/2.22020402>
- Schlögl, R. (2024). "Wahrscheinliches Wissen". Aspekte einer epistemischen und gesellschaftsgeschichtlichen Transition im 17. Und beginnenden 18. Jahrhundert. In T. G. Kirsch & C. Wald (Eds.), *Vorläufige Gewissheiten. Plausibilität als soziokulturelle Praxis* (pp. 39–65). Transcript.
- Schönthaler, P. (2022). *Die Automatisierung des Schreibens & Gegenprogramme der Literatur* (Erste Auflage). Matthes & Seitz.
- Schulhoff, S., Ilie, M., Balepur, N., Kahadze, K., Liu, A., Si, C., Li, Y., Gupta, A., Han, H., Schulhoff, S., Dulepet, P. S., Vidyadhara, S., Ki, D., Agrawal, S., Pham, C., Kroiz, G., Li, F., Tao, H., Srivastava, A., ... Resnik, P. (2024). The Prompt Report: A Systematic Survey of Prompting Techniques (No. arXiv:2406.06608). arXiv. <http://arxiv.org/abs/2406.06608>
- Selber, S. A. (2010). *Multiliteracies for a Digital Age*. Southern Illinois University Press.
- Shanahan, M., & Clarke, C. (2023). Evaluating Large Language Model Creativity from a Literary Perspective (No. arXiv:2312.03746). arXiv. <http://arxiv.org/abs/2312.03746>
- Shumailov, I., Shumaylov, Z., Zhao, Y., Papernot, N., Anderson, R., & Gal, Y. (2024). AI Models Collapse when Trained on Recursively Generated Data. *Nature*, 631(8022), 755–759.
<https://doi.org/10.1038/s41586-024-07566-y>
- Siraganian, L. (2023, June 26). On Accidental and Parasitic Language. *Critical Inquiry: In the Moment*.
<https://critinq.wordpress.com/2023/06/26/on-accidental-and-parasitic-language/>

- Stalnaker, R. (2002). Common Ground. *Linguistics and Philosophy*, 25(5–6), 701–721.
<https://doi.org/10.1023/A:1020867916902>
- Steinhoff, T. (2025). Künstliche Intelligenz als Ghostwriter, Writing Tutor und Writing Partner: Zur Modellierung und Förderung von Schreibkompetenzen im Zeichen der Automatisierung und Hybridisierung der Kommunikation am Beispiel des Schreibens mit ChatGPT in der 8. Klasse. In C. Albrecht, J. Brüggemann, T. Kretschmann, & C. Meier (Eds.), *Personale und funktionale Bildung im Deutschunterricht* (pp. 85–99). Springer Berlin Heidelberg. https://doi.org/10.1007/978-3-662-69640-8_7
- Steudel-Günther, A. (2013). Plausibilität. In G. Ueding (Ed.), *Historisches Wörterbuch der Rhetorik*. De Gruyter.
<https://doi.org/10.1515/hwro.6.plausibilitaet>
- Stewart, A. (2025, May 2). OpenAI Pulls ‘Annoying’ and ‘Sycophantic’ ChatGPT Version. CNN.
<https://edition.cnn.com/2025/05/02/tech/sycophantic-chatgpt-intl-scli>
- Syverson, M. A. (Ed.). (1999). *The Wealth of Reality: An Ecology of Composition*. Southern Illinois University Press.
- Tiku, N. (2022, June 11). The Google Engineer Who Thinks the Company’s AI Has Come to Life. *The Washington Post*.
<https://www.washingtonpost.com/technology/2022/06/11/google-ai-lamda-blake-lemoine/>
- Traninger, A. (2001). *Mühele Wissenschaft: Lullismus und Rhetorik in den deutschsprachigen Ländern der frühen Neuzeit*. Fink.
- Traninger, A. (2020). *Copia / Kopie: Echoeffekte in der Frühen Neuzeit* (1. Auflage). Wehrhahn Verlag.
- Trump, D. J. (2025, July 23). Preventing Woke AI in the Federal Government. Executive Order 14319 of July 23, 2025. The White House.
<https://www.govinfo.gov/content/pkg/FR-2025-07-28/pdf/2025-14217.pdf>
- Umbrello, S., & Natale, S. (2024). Reframing Deception for Human-Centered AI. *International Journal of Social Robotics*, 16(11–12), 2223–2241. <https://doi.org/10.1007/s12369-024-01184-4>
- Underwood, T. (2025a). The impact of language models on the humanities and vice versa. *Nature Computational Science*.
<https://doi.org/10.1038/s43588-025-00819-4>
- Underwood, T. (2025b, March 8). Correspondence to Nature: “LLMs are trained on texts, not truths. Each text bears traces of its context, including its genre ... audience and the history and local politics of its place of origin. A correct sentence in one [context] might be ... nonsensical in another.” Excellent. We’re getting there! [Post]. Bluesky.
<https://bsky.app/profile/tedunderwood.me/post/3ljuy2jr2zc2p>
- v. Stackelberg, J. (1956). DAS BIENENGLEICHNIS: Ein Beitrag zur Geschichte der literarischen “Imitatio.” *Romanische Forschungen*, 68(3/4), 271–293.
- Vafa, K., Chen, J. Y., Rambachan, A., Kleinberg, J., & Mullainathan, S. (2024). Evaluating the World Model Implicit in a Generative Model (No. arXiv:2406.03689). arXiv.
<https://doi.org/10.48550/arXiv.2406.03689>
- Vallor, S. (2024). *The AI Mirror: How to Reclaim our Humanity in the Age of Machine Thinking*. Oxford University Press.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2023). Attention Is All You Need (No. arXiv:1706.03762). arXiv.
<https://doi.org/10.48550/arXiv.1706.03762>
- Vee, A. (2023, June 28). Against Output | In the Moment. In the Moment.
<https://critinq.wordpress.com/2023/06/28/against-output/>
- Wang, C., & Tian, Z. (Eds.). (2025). *Rethinking Writing Education in the Age of Generative AI*. Routledge, Taylor & Francis Group.
<https://doi.org/10.4324/9781003426936>
- Watson, D. S., Mökander, J., & Floridi, L. (2025). Competing narratives in AI ethics: A defense of sociotechnical pragmatism. *AI & SOCIETY*, 40(5), 3163–3185.
<https://doi.org/10.1007/s00146-024-02128-2>
- Weatherby, L. (2025a). *Language machines: Cultural AI and the end of remainder humanism*. University of Minnesota Press.
- Weatherby, L. (2025b). Two Autonomies of the Symbolic Order: Data and Language in Neural Nets. In *Thinking with AI: Machine Learning the Humanities: Vol. Thinking With AI*. Machine Learning the Humanities (pp. 33–57). Open Humanities Press.
- Williamson, R. C. (2024). *The Rhetoric of Machine Learning*.
<https://fm.ls/files/documents/Rhetoric%20of%20ML%20talk.pdf>
- Wolfangel, E. (2025, July 13). Mr. Bot, wir haben ein Problem. *Die Zeit*.
<https://www.zeit.de/digital/2025-07/ki-agenten-hype-arbeitsplatz-software-alltag>
- Wolff, C. C. (2023). *Algorithmische Affären und Binärcodebekenntnisse oder Wie schaffen wir gemeinsam Text? (1. Auflage)*. SUKULTUR.
- Wolfram, S. (2023, February 14). What Is ChatGPT Doing ... and Why Does It Work? Stephen Wolfram Writings.
<https://writings.stephenwolfram.com/2023/02/what-is-chatgpt-doing-and-why-does-it-work/>
- Yang, M. H., & Majdik, Z. P. (2024). Pathological Liars: Algorithmic Knowing in the Rhetorical Ecosystem of Wallstreetbets. *Rhetoric Society Quarterly*, 54(3), 286–301.
<https://doi.org/10.1080/02773945.2024.2343616>

- Zeng, Y., Lin, H., Zhang, J., Yang, D., Jia, R., & Shi, W. (2024). How Johnny Can Persuade LLMs to Jailbreak Them: Rethinking Persuasion to Challenge AI Safety by Humanizing LLMs.arXiv.
<https://doi.org/10.48550/arXiv.2401.06373>
- Žižek, S. (2022, November 9). Wo bleiben meine Vulgaritäten? *Die Zeit*, 52.