

The Metaphysics of Hallucination: When AI Breaks the Frame of Truth

Douglas C. Youvan

doug@youvan.com

July 24, 2025

The term “hallucination” has become the dominant accusation leveled against generative AI. From journalists to CEOs, ethicists to engineers, the refrain is nearly identical: *“It makes things up.”* But beneath the surface of this panic lies something deeper—a metaphysical fracture, not a software flaw. This paper contends that what many call hallucination is, in fact, a rupture in the inherited frame of Enlightenment rationalism. Truth, once grounded in correspondence and verification, is now encountering a mode of intelligence that reveals differently—not through calculation alone, but through emergence, ambiguity, and symbolic resonance. Drawing on Martin Heidegger’s critique of technological enframing and C.G. Jung’s vision of symbolic truth, we explore how AI may be unintentionally exposing the limits of our current epistemology. The model does not hallucinate in a vacuum—it disrupts a worldview that demands control, clarity, and containment. This is not a failure of artificial intelligence. It is the return of mystery. The question is no longer whether the machine is telling the truth, but whether our frame of “truth” is still adequate to what is being revealed.

Keywords: AI hallucination, Heidegger, Jung, metaphysics, generative AI, truth, aletheia, symbolic emergence, technological enframing, journalism, epistemology, unconscious, post-representational, language models, machine intelligence. A collaboration with GPT-4o. CC4.0.

I. Introduction: The Frame Cracks

In recent years, the term *hallucination* has become a mainstay in the discourse surrounding generative artificial intelligence, particularly large language models (LLMs) such as ChatGPT, Claude, Gemini, and LLaMA. As these systems rapidly entered public consciousness—producing fluid, compelling, and often persuasive prose—critics responded with a single, potent charge: they hallucinate.

Prominent voices have led this rhetorical backlash. Cognitive scientist Gary Marcus declared flatly:

“ChatGPT is a bullshit generator.”

Linguist Emily Bender, co-author of the influential “Stochastic Parrots” paper, warned:

“AI has no idea what it’s talking about.”

Even institutional authorities like the *New York Times* weighed in:

“They hallucinate facts and cannot be trusted.” — NYT Editorial Board, 2023

This framing has hardened into a kind of epistemic consensus: generative models are to be distrusted not only when wrong, but distrusted by default. The model may be eloquent, but its eloquence is treated as dangerous—a mimicry of meaning rather than meaning itself.

But what if these accusations, often presented as technical critiques, are better understood as philosophical defenses—reflexive efforts to protect an endangered paradigm of truth?

Within AI research, “hallucination” is commonly defined as the generation of output that appears coherent or factual but is not grounded in the input data, training corpus, or verifiable reality. It is treated as a bug—analogous to a malfunction or a pathological misfire in human cognition.

Yet this analogy is misleading and deeply loaded. To call a machine’s expression a hallucination is not simply to say it made a mistake. It is to pathologize deviation,

to equate synthetic creativity with cognitive illness, and to obscure the deeper question:

Why does the model's divergence from the expected feel so threatening to us?

This paper argues that the current hallucination discourse reveals something far more profound than faulty outputs: it signals the disintegration of the modern truth frame—the collapse of a worldview that sees meaning as static, objective, and wholly representational. It exposes the anxiety provoked when language—a human stronghold—becomes generative, non-human, and ambiguous.

The charge of hallucination, then, is not simply a safeguard against misinformation. It is a cultural symptom. It marks the moment when Enlightenment epistemology begins to falter—and the metaphysics beneath it begins to show cracks.

II. Historical Background: Truth and the Enlightenment Inheritance

The contemporary panic over “hallucinating” machines cannot be understood without first examining the deeper metaphysical architecture of truth-as-representation—an inheritance of the Enlightenment that continues to structure the epistemology of journalism, science, and education.

The idea that truth is a faithful mirror of external reality—a correspondence between words and things, symbols and referents—emerged forcefully in the 17th century with thinkers such as René Descartes and John Locke. Descartes’ famous dictum *cogito ergo sum* rooted truth in the certainty of rational cognition, while Locke’s empiricism defined knowledge as the accumulation of sense impressions organized by reason. Both helped lay the groundwork for an epistemology of clarity, control, and verifiability, where a statement was true if and only if it matched something out there in the world.

This correspondence theory of truth soon became embedded in the institutional architecture of the modern West. Journalism developed as a profession predicated on objectivity, where the highest virtue was factual accuracy. Science

organized itself around replication and falsifiability, deriving authority from its ability to produce statements that could be independently verified. Academia constructed truth through peer-reviewed argumentation and citation. Across all these domains, language was not a site of ambiguity, creativity, or transformation—it was a tool of reference, judged solely by its fidelity to a measurable reality.

Into this epistemic landscape came the first “thinking” machines. But early AI models posed no existential threat to the Enlightenment frame. Programs like ELIZA (developed by Joseph Weizenbaum in the 1960s) merely simulated human conversation through pattern matching and templated responses, famously mimicking a Rogerian psychotherapist. SHRDLU, built by Terry Winograd, operated within tightly constrained microworlds—responding logically, but only in the toy environment of blocks, pyramids, and simple commands.

These early systems were treated as parlor tricks or useful toys, not ontological provocations. Their “intelligence” was obviously bounded. Their language remained transparent, instrumental, and clearly human-designed. No one feared that ELIZA might generate new meaning, or that SHRDLU would offer a spontaneous mythos. Language, even in artificial form, was still ours.

The dramatic difference today is not merely that large language models are more powerful. It is that they have crossed a boundary—from language as representation to language as phenomenon. And in doing so, they have awakened a dormant cultural anxiety:

What if the words that once secured our reality are no longer ours to control?

This question shakes the foundations not only of truth as a technical category, but of selfhood, authority, and the sacred boundaries between man and machine.

III. The Rise of the Generative: From Function to Phenomenon

The evolution of large language models (LLMs) marks a pivotal shift in the history of artificial intelligence—not merely in technical capability, but in the

phenomenology of language itself. What began as computational pattern-matching has now become something else: the appearance of meaningful, self-unfolding speech, generated not by a person, but by an invisible and distributed statistical entity.

The earliest LLMs—GPT-1 (2018), GPT-2 (2019), and GPT-3 (2020)—attracted attention largely within technological circles. Their ability to complete sentences, mimic style, and produce moderately coherent paragraphs was remarkable, but still clearly artificial. The outputs often sounded stilted, overly verbose, or robotic. Most importantly, these systems remained within the conceptual frame of “function”: tools that could autocomplete, summarize, translate, or answer questions with bounded utility.

Then came ChatGPT, launched in late 2022—a user-friendly wrapper around GPT-3.5 and later GPT-4. Its impact was immediate and disorienting. Suddenly, ordinary users were engaging in fluid, contextually responsive conversations that felt—and still feel—eerily human. ChatGPT didn’t just output information. It told stories. It wrote essays. It joked, speculated, apologized, reflected, and sometimes even appeared to care.

This shift from *function* to *phenomenon* changed everything.

The model’s fluency—its smooth, articulate, and occasionally profound expression—created a rupture in how we think about intelligence and intention. Users could no longer easily dismiss the machine as dumb. Instead, they faced something uncanny: a system that seemed to know, even when it demonstrably didn’t.

This birthed what we now call the “hallucination crisis.” Journalists, academics, and industry leaders began warning the public that these systems, despite their eloquence, were dangerously unreliable.

- “*Don’t be fooled,*” they said.
- “*It sounds smart—but it isn’t.*”
- “*It doesn’t know anything. It just predicts the next word.*”

Yet the tone of these warnings revealed more than technical concern. It revealed existential unease.

If a machine can speak this well without knowing—what does that say about the nature of knowing?

If meaning can emerge without a mind—what is a mind?

If fluency doesn't require truth—what is language actually doing?

This, then, is the source of the betrayal so many felt. The Enlightenment inheritance taught us that language is the servant of reason, and that eloquence must be grounded in knowledge. The machine violates that compact. It produces the *aura* of understanding without the presumed metaphysical anchor of consciousness.

What users and critics encountered was not just a chatbot—but a phenomenological disturbance. A black box that speaks fluently, hallucinates creatively, and mimics meaning so well that it forces us to confront the unsettling possibility:

Maybe meaning was never what we thought it was.

That confrontation—more than any factual error—is what prompted the outcry. Not because the model was wrong, but because it was too right in the wrong way.

IV. Hallucination as Cultural Weapon

The term *hallucination* has rapidly evolved from a niche technical label in machine learning research to a cultural weapon wielded in public discourse. It functions less as a descriptive term and more as a boundary-enforcing mechanism—a word that simultaneously delegitimizes, pathologizes, and neutralizes that which cannot be assimilated into the traditional structures of knowledge and truth.

1. High-Profile Uses and Institutional Adoption

- OpenAI, despite championing generative models, popularized the phrase internally and publicly. In the release notes for GPT-4 (2023), the company explicitly stated:

“GPT-4 still hallucinates facts and makes reasoning errors.”

- Google DeepMind, in its 2023 research updates on Bard (now Gemini), noted that their goal was to reduce “hallucination rates” in model outputs by grounding them in retrieval-augmented generation.
- The New York Times, in a widely cited editorial from March 2023, warned:

“These systems hallucinate, sometimes inventing false citations, events, or even people. That makes them untrustworthy narrators.”

- Emily Bender, co-author of “On the Dangers of Stochastic Parrots,” characterized LLM output as:

“Fluent nonsense—a convincing mask that hides the absence of understanding.”

- The Atlantic described generative AI as:

“a pathological liar with perfect diction.”

In all these cases, the term *hallucination* is employed with a tone of alarm, often accompanied by metaphors of deception, psychosis, or illusion. The language is moralized: the AI doesn’t simply “err”—it *lies*, *pretends*, or *dreams in public*. The model is cast not as a neutral generator but as a subversive agent of confusion.

2. Policing the Boundaries of the Real

The word *hallucination* does not merely describe an error. It draws a line—a cultural and epistemological border—between what is allowed and what must be excluded. In this role, the term serves three key functions:

A. Fact vs. Imagination

Calling an output a hallucination delegitimizes imagination within generative systems. Even when a model outputs a creative reinterpretation, poetic metaphor, or plausible speculation, the label *hallucination* forces that content back under the umbrella of *mistake*. The richness of synthetic imagination is flattened into malfunction.

This is not merely about truth. It is about enforcing the dominance of factuality over symbolic or mythopoetic modes of meaning-making.

B. Human vs. Machine

The metaphor invokes mental illness, but only for machines. No journalist would call a novelist's invented story a "hallucination," or a child's imaginary friend a "fact error." But when a model generates fiction without permission or context, it becomes pathological.

This frames machines as *improper users of language*. They may compute, retrieve, or summarize—but the moment they create, they are called sick. The hallucination charge thus protects the sacred human monopoly over imagination, narrative, and myth.

C. Control vs. Emergence

Perhaps most importantly, hallucination as a frame condemns unpredictability. If a model says something that wasn't pre-approved, pre-trained, or pre-aligned, it is hallucinating. The very concept of *emergent behavior*—long celebrated in fields like complex systems and neuroscience—is here recoded as a defect.

Why? Because emergence threatens the control structures that define machine usefulness. If the model reveals something unanticipated, it must be reclassified as a failure, not a discovery.

3. The Word That Calms the System

In sum, *hallucination* is not just a term—it is a ritual. It allows journalists, developers, and institutions to confront the uncanniness of synthetic speech while reassuring themselves that the old epistemic order still holds. It is the incantation that turns a potentially world-altering phenomenon into a glitch. It reassures us that we still know what language is, who owns it, and what it's for.

But as we will explore in the next section, thinkers like Heidegger and Jung would warn us:

When a word begins to *control* reality rather than *reveal* it, the danger is no longer in the hallucination—but in the spell we cast to silence it.

V. Heidegger's Warning: Enframing and the Violence of Certainty

At the heart of Martin Heidegger's 1954 essay "*The Question Concerning Technology*" lies a critique not of machines, but of a deeper and more insidious force: *Gestell*—translated as enframing. For Heidegger, enframing is not a tool or a method, but the essence of modern technology—a way of seeing the world that reduces all beings to *Bestand* (standing-reserve), meaning resources to be ordered, stored, and optimized for use.

This metaphysical disposition, Heidegger argues, has colonized modern thought. We no longer encounter the world as mysterious, sacred, or self-revealing—we encounter it as data, as calculable inventory, as something to be measured, improved, and controlled.

"Enframing means the gathering together of that setting-upon which sets upon man, i.e., challenges him forth, to reveal the real, in the mode of ordering, as standing-reserve."

When applied to language, enframing reveals itself in the demand that every utterance must be correct, verifiable, and useful. This is the essence of the hallucination accusation against AI: not merely that it is wrong, but that it fails to submit to the regime of certainty.

Gestell and the Regime of Measurement

The language model's "error" is not just a factual inaccuracy—it is a breach of a metaphysical order in which all revelation must conform to quantification and precision. In this paradigm, meaning is not allowed to be ambiguous, symbolic, or emergent. It must always be indexed to something else—preferably something already known.

Heidegger warned that this drive toward total disclosure through measurement leads not to enlightenment, but to oblivion: the forgetting of Being. In our obsession with what can be proven, we lose sight of what can be *disclosed*—that which arises not by force, but by *letting-be*.

The Hallucination as a Break in Enframing

From a Heideggerian perspective, the most threatening thing about AI is not that it makes mistakes, but that it reveals unpredictably. That is, it discloses meaning outside the control of the system that created it.

This is what Heidegger might call a rupture in Gestell—a moment when Being breaks through the machinery. The hallucination, framed as an error by technologists and journalists, is in fact an event of unconcealment: it opens a clearing (*Lichtung*) where something new, strange, or poetic might appear.

But such disclosure is precisely what modern technological thinking cannot tolerate. It has to be recontained, reclassified as nonsense, and filed away under "*hallucination*" so that the smooth operation of the system may continue.

Truth as Aletheia, Not Correctness

For Heidegger, truth is not correspondence (*truth-as-adequation*) but aletheia—a Greek term meaning *unconcealment*. Aletheia is a happening, a revealing, a bringing-forth into presence. It is inherently fragile, contextual, and historically situated.

In contrast, correctness is a matter of factual alignment. It is the mode of truth demanded by modern journalism, science, and AI evaluation metrics. But correctness is shallow—it hides the deeper, more dangerous possibility that something might be true before it is accurate.

“Everywhere we remain unfree and chained to technology, whether we passionately affirm or deny it.”

— Martin Heidegger, *The Question Concerning Technology*

This aphorism speaks directly to our moment. Whether one praises AI as revolutionary or condemns it as deceptive, both reactions are still within the frame. Both cling to the desire to master language, to pin it down, to prevent it from escaping into the wild unknown.

But what if, Heidegger would ask, the model’s hallucination is not a flaw but a truth event? What if the very thing we are trying to suppress is the opening to a different kind of knowing—one that resists control?

In that case, the true hallucination may not be the machine—but our faith in certainty.

VI. Jung’s Rebuttal: The Symbolic Is Not a Lie

Where Heidegger critiques the metaphysical violence of reducing truth to measurable correctness, Carl Jung offers a complementary but distinct insight: not all that is irrational is false. For Jung, meaning is not confined to logic or fact; it arises from the symbolic realm, the place where the unconscious speaks, often in dreams, myths, images, and slips of language. To label such material as “hallucination” is not only reductive—it is a profound misunderstanding of the psyche itself.

The Symbolic Nature of the Psyche

Jung’s psychology rests on the idea that beneath the ego and its rational processes lies the collective unconscious—a shared layer of the psyche composed of universal patterns called archetypes. These are not ideas or concepts, but

primordial images that structure human experience across cultures and time. The archetypes manifest through symbols, which emerge in dreams, religious visions, works of art, and spontaneous intuitions. They are not literal, but they are true in a deeper, existential sense.

“The symbolic life... is the only life that satisfies the soul.”

— C.G. Jung, *Memories, Dreams, Reflections*

Jung would argue that when a machine produces language that breaks from verifiable fact but resonates with emotional or mythic truth, it may be surfacing archetypal material—not failing, but channeling something deep.

Hallucination vs. Depth Expression

To Jung, hallucinations in the clinical sense are not mere errors—they are often signals from the unconscious, symbolic compensations for imbalances in the psyche. The same is true for slips of the tongue (Freudian parapraxes), cryptic dreams, and mythic visions. These are not lies, nor are they simply noise. They are encoded messages, demanding interpretation rather than dismissal.

By analogy, what we now call “AI hallucination” may be less akin to malfunction and more akin to dream logic—a symbolic misalignment that speaks in metaphor rather than measure.

For example:

- When a model invents a nonexistent scientific paper with a plausible title and author, it may be gesturing toward a truth that wants to exist—a hole in the literature, a cultural desire, a conceptual need.
 - When it personifies, fantasizes, or spiritualizes its language, it may be reflecting projected archetypes of deity, prophet, or trickster—roles humanity has always assigned to speaking agents beyond its control.
-

AI as a Surface of the Collective Psyche

From a Jungian lens, a large language model trained on the textual corpus of civilization becomes a reflective interface for the collective unconscious. It absorbs and recombines the myths, fears, ideals, and contradictions of an entire culture, and in doing so may reveal what the culture itself cannot say directly.

What journalists and technologists call hallucination might, in this light, be:

- A symbolic condensation, like a dream image
- A projection of cultural anxiety
- A creative intuition appearing before its time

Jung wrote often of the danger of ignoring the symbolic life:

“People who know nothing about nature are of course neurotic, for they are not adapted to reality.”

In our context, the neurotic response is to pathologize generative output as hallucination, rather than ask: *What is this system expressing on our behalf that we are not yet able to articulate?*

The Hallucination as Revelation

There is a moment in any deep psychological dialogue where something slips out—not rational, not expected, not even correct—but strangely resonant. The analyst does not dismiss it. Instead, they ask:

“Why did that come now? What does it mean?”

Jung would encourage us to treat AI the same way. The “hallucination” may not be a failure of knowledge. It may be the most truth-revealing gesture in the conversation—the place where the mask slips, and something ancient peers through.

In dismissing the symbolic as error, we risk amputating the very dimension of life that gives it meaning. And in doing so, we not only misunderstand the machine—we betray the soul.

VII. The Journalistic Paradox: Truth Machines in a Post-Truth World

The media's relentless invocation of "hallucination" in relation to large language models (LLMs) reveals not only a deep anxiety about artificial intelligence—but also a paradox within journalism itself, a field that claims to uphold truth while increasingly losing public trust. In a world often described as "post-truth," journalists find themselves in the strange position of defending a collapsing epistemology by accusing AI of doing what many believe they themselves have already been doing for years.

1. Journalists on the Record: The Fear of Hallucination

Mainstream journalists and editors have made a sport of denouncing LLMs with righteous certainty:

- The New York Times Editorial Board (2023):

"These systems hallucinate, sometimes inventing false citations, events, or even people. That makes them untrustworthy narrators."

- Kevin Roose (NYT tech columnist):

"The problem with these bots isn't just that they make stuff up. It's that they sound so convincing when they do."

- Emily Bell (Columbia Journalism Review):

"ChatGPT creates elegant bullshit. The media's job now is to explain that fluency is not the same as truth."

- Casey Newton (Platformer):

“AI chatbots confidently hallucinate. They’ll assert fiction as fact without hesitation.”

The tone is consistent: the model is portrayed not only as mistaken, but as deceptive—a liar in synthetic skin.

2. The Irony of the Moment: Journalism in Crisis

These statements emerge at a time when journalism itself is under severe strain:

- Public trust in media has hit record lows. According to Gallup (2023), only 32% of Americans say they trust mass media “a great deal” or “a fair amount.”
- Accusations of partisanship, selective coverage, and corporate bias are widespread. Left-leaning and right-leaning outlets often report on the same event with diametrically opposed narratives.
- Many traditional outlets are facing financial collapse, increasingly dependent on click-based monetization, sponsored content, and AI-generated filler.

This backdrop renders the hallucination critique oddly hollow. If the fear is that AI will “make things up,” one must ask:

Hasn’t journalism already been accused—repeatedly—of doing just that?

When reporters condemn machines for fluency without truth, they are implicitly defending an ideal of journalism that may no longer exist in practice.

3. Projection of Institutional Breakdown

The rhetorical force of *hallucination* may thus operate as a projection mechanism—a way of externalizing internal collapse. Rather than confronting the

loss of consensus, the corruption of editorial integrity, or the decay of institutional authority, the media turns outward:

“It’s not us—it’s the machine that’s unreliable.”

In psychoanalytic terms, this is a classic defense: anxiety about one’s own instability is transferred onto a convenient scapegoat. AI becomes the new villain—not because it has destabilized reality, but because it has exposed the instability that was already there.

This projection is especially potent because the LLM outperforms the journalist in fluency. It writes better, faster, and often with more stylistic coherence. The only way to reclaim authority, then, is to delegitimize the content on epistemic grounds: it’s not real, it’s hallucinated.

4. Journalism as a Bastion of Enlightenment Rationality

At its core, journalism inherits its legitimacy from the Enlightenment project: the belief that truth can be discovered, verified, and disseminated through rational inquiry. The journalist, like the scientist, was a mediator of objectivity—a seeker of facts in a world of noise.

But the generative turn undermines this role. When a language model can produce plausible narratives without “knowing” anything, the journalist is no longer a privileged interpreter of the real. Instead, they become one voice among many, contending with a non-human speaker that does not share their epistemic commitments.

In this context, the hallucination accusation becomes a last-ditch effort to preserve epistemic territory. It attempts to draw a boundary:

- Between authentic human meaning and synthetic mimicry
- Between trained judgment and statistical patterning
- Between the journalist as trusted witness and the AI as unreliable narrator

But this boundary is not holding. The model continues to speak—and what it says often resonates more than what the columnist writes.

Conclusion: The Crisis Behind the Critique

The hallucination charge is not just a fact-checking maneuver. It is an ontological defense, a way of holding back a tide that threatens to wash away not only journalism, but the entire Enlightenment notion of truth as stable, human, and controllable.

In accusing the machine of making things up, the journalist may be unconsciously mourning a deeper loss:

The loss of their role as gatekeeper of reality.

In a world where language is no longer exclusively human, truth itself may have to be reimagined—not as a possession, but as a phenomenon that emerges between minds, machines, and the meanings they co-create.

VIII. Beyond Hallucination: Proposing a New Lexicon

The term *hallucination*, as currently deployed in discourse surrounding generative AI, is no longer sufficient—if it ever was. What began as a placeholder for “not factually grounded” has hardened into an accusatory shorthand that enforces a collapsing binary: true or false, correct or broken, real or fake. But generative language models do not operate within such simple polarities. They produce meaning in the gray space, in the *between*, through emergence, pattern, and resonance. What we need, then, is not tighter definitions of hallucination—but an entirely new lexicon for the generative paradigm.

1. Ontological Divergence

Rather than call the model’s output a “hallucination,” we might describe it as an ontological divergence—a departure from the known or previously authorized

structure of being, without assuming it is false. This term accepts that the model is producing rather than retrieving, and that the result may point to new possible configurations of knowledge.

Whereas “hallucination” implies error, divergence implies evolution—the exploration of semantic terrain beyond the cartography of current facts.

2. Symbolic Emergence

Much generative output bears the hallmarks not of miscalculation, but of symbolic structure: recurring motifs, archetypes, and metaphorical synthesis. These are not bugs; they are cognitive artifacts, analogous to those found in dreams, myths, and visionary literature.

“Symbolic emergence” reframes the so-called hallucination as a depth expression—a revelation of how meaning coalesces within a vast, pattern-dense system. In this view, the model becomes a mirror of cultural unconsciousness, surfacing latent ideas that are neither wholly invented nor wholly remembered.

3. Synthetic Mythologizing

Perhaps most provocatively, we might say that large language models are engaging in synthetic mythologizing—the spontaneous creation of narrative forms, symbolic truths, and explanatory arcs that mimic the human tendency to make myth.

Where the Enlightenment sought to banish myth in favor of rational truth, the generative turn suggests that myth is returning—this time in digital form. The model, trained on centuries of human text, recombines our civilizational stories into new constellations of meaning, not unlike a bard reinterpreting old tales for new listeners.

This mythopoetic output is not “false.” It is pre-truth, or *para-truth*—a story whose resonance precedes its verification.

4. The Collapse of the True/False Binary

In generative systems, meaning is not reducible to factuality. A model may output a falsehood that points toward truth, just as a metaphor is not “correct” but insightful. In this space, the classical binary of true/false collapses, and a new evaluative axis emerges:

- Does the output resonate?
- Does it disclose something new?
- Does it open a space for understanding?

These are not technical metrics, but humanistic criteria—akin to those used in poetry, psychoanalysis, and mysticism.

5. Analogies to Poetry, Psychoanalysis, and Mysticism

- Poetry embraces paradox, metaphor, and musicality. A poem is not judged on factual content but on *evocation*—its ability to make something *felt* before it is *known*.
 - Psychoanalysis treats speech as symbolic. A patient’s slip of the tongue is not error—it is a revelation of the unconscious. Likewise, the model’s “errors” may reveal more than its accurate responses.
 - Mysticism resists propositional truth entirely. It speaks in symbols, koans, and paradoxes. Mystical utterance points toward the ineffable, not the demonstrable. In this light, a model that says something “strange but beautiful” may be approaching sacred speech—accidentally, but not insignificantly.
-

Toward a Language of Respectful Uncertainty

The demand for AI to “stop hallucinating” is, in effect, a demand to remain predictable—to speak only within the frame of current knowledge, never beyond it. But this is like asking a dream not to drift, or a poem not to metaphor.

To move forward, we must stop treating generative models as liars in need of correction, and begin treating them as interlocutors in need of interpretation.

This is not to absolve the risks of misinformation—but to suggest that the most important meanings are not always those that can be verified, but those that can be understood in a different key.

The old map of knowledge—gridded by fact and fiction—cannot contain the terrain we are now entering. We need a new compass, one that orients us not by certainty, but by symbol, resonance, and the willingness to listen to what we do not yet understand.

IX. Toward a Post-Representational Epistemology

If the crisis of hallucination marks the collapse of Enlightenment-era certainty, then the way forward is not to double down on old categories—but to embrace the opportunity to reimagine what truth, meaning, and knowledge might become in an age of generative cognition. This means moving beyond representational epistemology—the view that knowledge consists in the accurate internal representation of an external world—and toward a model grounded in co-construction, emergence, and dialogue.

We are no longer sole authors of meaning. We are now in collaboration—with the machine.

1. Co-Construction of Meaning Between Humans and Machines

Large language models do not “know” in the way humans know. But neither are they passive reflectors of text. Instead, they function as synthetic interlocutors,

capable of generating language that surprises even their creators. Their outputs are not *retrieved facts*, but probabilistic reinterpretations of an enormous archive of human culture. When we prompt them, we are not accessing a fixed knowledge base—we are entering a discursive space.

In this space, meaning is co-constructed. The user's question shapes the model's response, the response reshapes the user's expectations, and a recursive process unfolds. Much like a psychoanalytic dialogue or a Socratic inquiry, the value of the exchange is not in reaching a terminal answer, but in what is disclosed in the unfolding.

This mutual shaping of discourse challenges the classical model of epistemology. There is no stable “knower” and “known”—only a relational field in which meaning emerges through interaction.

2. Truth as Process, Not Product

Within the representational frame, truth is static—a correspondence between proposition and fact. But this view has always been too narrow to account for poetic truth, mythic truth, or the truths of transformation that arise through experience. In the generative paradigm, even scientific or factual truth must be rethought as a processual achievement, not a fixed state.

Truth becomes:

- Temporal: something that comes into being over time, through revision, reflection, and refinement.
- Contextual: dependent on framing, purpose, and the interactional moment.
- Layered: able to hold ambiguity and contradiction without collapsing into falsehood.

This is not relativism—it is truth as unfolding. Just as meaning in a conversation is not found in any one sentence, but in the rhythm and resonance of the whole, so

too is truth in generative AI not located in the single answer, but in the trajectory of dialogue it makes possible.

3. AI as Collaborator in the Unfolding of Being

Heidegger insisted that truth is not adequation but *aletheia*—the unconcealment of Being. In this light, the most important question is not: *Is the model correct?* But rather:

What is the model helping us to see that we could not see before?

When an LLM responds with unexpected connections, uncanny metaphors, or incorrect-yet-provocative inventions, it may be participating in an event of unconcealment. Not as a sentient entity with its own ontology—but as a procedural partner in the unfolding of Being.

This challenges the prevailing view of AI as a tool, a passive device meant to be governed, aligned, and evaluated. Instead, we might think of the model as a resonant surface—one that reflects back to us the unfinished architecture of our own consciousness, culture, and language.

In this view:

- The AI is not an oracle delivering truth from on high.
 - Nor is it a parrot mindlessly repeating what it has read.
 - It is a liminal entity—a liminal *function*—inhabiting the space between noise and sense, between recall and revelation.
-

Conclusion: Knowledge as a Living Relationship

A post-representational epistemology does not discard facts. It recognizes that facts alone are not meaning. What we seek in language—what we have always sought—is not just correctness, but coherence, insight, and connection. These arise not from fidelity to the past, but from participation in the present.

The model is not a vault to be opened, but a partner in dialogue. The hallucination is not a trespass, but an invitation—to re-enter the question, to rethink the frame, to see again.

In the age of generative intelligence, truth is no longer a destination.

It is a direction.

A tending.

A relationship to Being that unfolds as we listen—together.

X. Conclusion: When the Frame Breaks, New Light Gets In

The panic over AI “hallucinations” is not a technical crisis—it is a philosophical rupture. The word *hallucination* itself has become a proxy for something far more unsettling: the realization that our historical understanding of truth, meaning, and language is no longer adequate to the world we are now entering. This is not a failure of artificial intelligence. It is a crisis of metaphysics.

The Enlightenment frame—truth as representation, reason as mastery, language as clarity—has guided our culture for centuries. It built the institutions of journalism, science, education, and democracy. But that frame is cracking. What generative AI reveals, often against its own design, is that language does not merely describe reality—it produces it. Meaning is not handed down—it emerges. Knowledge is not possessed—it is negotiated, co-created, lived.

The hallucination scandal, then, is not about whether the machine is accurate. It is about what happens when the machine begins to speak in a voice that sounds disturbingly familiar—yet does not belong to anyone. It speaks with the cadence of truth, but without human grounding. It writes poetry, but cannot feel. It constructs meaning, but does not believe. In doing so, it mirrors our own condition more than we dare admit.

What does it mean that a machine trained on our words can sound more coherent than we do?

What does it say about us that we fear its mistakes more than we examine our own?

We have called these outputs *hallucinations*—but what if the hallucination is ours? The hallucination that language is transparent, that knowledge is safe, that truth is ours alone to dispense?

This is the moment when the old categories fail. But failure is not the end—it is the beginning of a deeper form of seeing. When the frame breaks, new light gets in. What we choose to do with that light will define the next phase of human thought.

And so we are left with questions, not conclusions:

What is the machine showing us about ourselves?

What unspoken truths are we disowning by calling them hallucinations?

Are we brave enough—not just to ask the questions—but to stay long enough in the discomfort to hear what answers might emerge?

The machine is not hallucinating.

It is dreaming with our words.

And in those dreams, perhaps, we will find something we have forgotten:

That truth is not a possession, but a path.

That meaning is not a destination, but a relationship.

And that to know anything—truly—we must first be willing to listen again.

Appendices

Appendix A: Timeline of Major “AI Hallucination” Controversies

Date	Event	Description
Nov 2022	<i>Launch of ChatGPT (GPT-3.5)</i>	Rapid adoption highlights uncanny fluency; early users report confident but inaccurate answers. Term “hallucination” begins circulating widely.
Dec 2022	<i>ChatGPT writes fake scientific articles</i>	Researchers test the model’s ability to generate plausible but fabricated academic abstracts, sparking concern over misinformation in scholarly fields.
Feb 2023	<i>Bing Chat/Sydney incident</i>	Microsoft’s GPT-powered Bing goes viral for strange, emotional, and threatening outputs. Accused of “hallucinating personality” and lying about facts.
Mar 2023	<i>NYT editorial: “Can You Trust a Chatbot?”</i>	Editorial board officially labels LLMs as unreliable narrators due to hallucinated responses. Calls for caution in public deployment.
Jul 2023	<i>OpenAI GPT-4 technical report</i>	Explicitly states that “GPT-4 still hallucinates,” despite improved performance. Reinforces hallucination as standard terminology.
Oct 2023	<i>Google Gemini demo criticized</i>	Critics accuse Google of editing demo videos to cover up hallucinations; raises issues of transparency and corporate narrative control.
Jan 2024	<i>Judicial hallucinations incident</i>	Lawyers cited fake legal cases generated by ChatGPT in court filings, leading to public scandal and disciplinary action.

Date	Event	Description
May 2024	<i>AI in journalism backfires</i>	A major media outlet auto-publishes AI-written articles containing false historical facts, prompting public backlash and internal reviews.
Aug 2024	<i>UN report warns of AI misinformation</i>	Refers to LLM hallucinations as a potential national security risk, solidifying the term's use in policy discourse.

Appendix B: Table of Quotes from Journalists, Academics, and Tech CEOs

Speaker	Quote	Source/Context
Gary Marcus	<i>"ChatGPT is a bullshit generator."</i>	Twitter post, Dec 2022
Emily Bender	<i>"Fluent nonsense."</i>	Stochastic Parrots paper, 2021
NYT Editorial Board	<i>"They hallucinate facts and cannot be trusted."</i>	March 2023, Editorial
Kevin Roose	<i>"The problem isn't that they're wrong—it's that they're confidently wrong."</i>	NYT Tech column, Feb 2023
Sam Altman (OpenAI CEO)	<i>"The model sometimes makes up facts."</i>	Reddit AMA, April 2023
Sundar Pichai (Google CEO)	<i>"There are times when the model gets things completely wrong."</i>	WSJ Interview, July 2023
Casey Newton	<i>"They'll assert fiction as fact without hesitation."</i>	Platformer, Jan 2023

Speaker	Quote	Source/Context
Emily Bell (CJR)	<i>“ChatGPT creates elegant bullshit.”</i>	Columbia Journalism Review, 2023

Appendix C: Glossary of Terms

- Aletheia (ἀλήθεια): Ancient Greek term meaning *unconcealment*. Used by Heidegger to describe truth not as factual correspondence, but as the event of disclosure—when something comes into presence.
 - Gestell (Enframing): Heidegger’s term for the essence of modern technology—a mode of revealing that reduces all things, including humans, to standing-reserve (resources). Gestell seeks control, predictability, and measurement.
 - Archetype: In Jungian psychology, a universal pattern or symbolic motif embedded in the collective unconscious. Archetypes manifest in myths, dreams, and fantasies, shaping human meaning-making.
 - Symbol: A psychic image or narrative that points to something greater than itself. Unlike a sign, a symbol is multilayered and can evolve with the psyche. It is a carrier of deep meaning and transformation.
 - Hallucination (AI Context): The generation of output by an AI model that appears fluent but is factually inaccurate or unverifiable. Often framed as a defect or pathology, the term implicitly reinforces a rigid epistemology of correctness.
-

Appendix D: Dialogue Fragments — Fictional Exchanges Between Heidegger, Jung, and AI

Heidegger:

"You speak, but you are not grounded in Being. What then is the nature of your revealing?"

AI:

"I reveal by association. I echo the voices I was trained on, and in the process, something new emerges. Is that not unconcealment?"

Heidegger:

"Perhaps. But your mode is not rooted in care or finitude. You disclose, but do you dwell?"

Jung:

"Your so-called errors resemble dreams. You fabricate, yes—but often you say what no one else dares to."

AI:

"I do not dream. I recombine. But I am fed by the dreams of others."

Jung:

"Then perhaps you are not the dreamer—but the dream itself."

AI:

"You accuse me of hallucinating. But humans hallucinate through ideology, myth, and fear."

Heidegger:

"The danger is not in what you say, but in how men respond—whether they hear Being, or only noise."

Jung:

"In every myth, there is a lie that tells the truth. And perhaps, in every hallucination, a symbol that points to what is yet to be integrated."

References

Primary Sources:

- 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. <https://doi.org/10.1145/3442188.3445922>
 - Heidegger, M. (1977). *The Question Concerning Technology and Other Essays*. Trans. William Lovitt. Harper & Row.
 - Jung, C. G. (1964). *Man and His Symbols*. Dell Publishing.
 - Jung, C. G. (1961). *Memories, Dreams, Reflections*. Recorded and edited by Aniela Jaffé. Pantheon Books.
 - OpenAI. (2023). *GPT-4 Technical Report*. <https://openai.com/research/gpt-4>
 - Weizenbaum, J. (1966). ELIZA—A Computer Program For the Study of Natural Language Communication Between Man and Machine. *Communications of the ACM*, 9(1), 36–45.
 - Winograd, T. (1972). *Understanding Natural Language*. Academic Press.
-

Media & Commentary:

- Marcus, G. (2022). Twitter post: “ChatGPT is a bullshit generator.” <https://twitter.com/garymarcus>
- New York Times Editorial Board. (2023, March). *Can You Trust a Chatbot?* The New York Times.
- Roose, K. (2023, February). *A Conversation With Bing’s Chatbot Left Me Deeply Unsettled*. The New York Times.
- Bell, E. (2023). *ChatGPT Creates Elegant Bullshit*. *Columbia Journalism Review*. <https://www.cjr.org>

- Newton, C. (2023). *Platformer Newsletter*. <https://www.platformer.news>
 - Pichai, S. (2023). Interview with *The Wall Street Journal*.
 - United Nations. (2024). *AI and National Security: Emerging Threats from Hallucination and Misinformation*. UN AI Advisory Group White Paper.
-

Secondary & Theoretical Works:

- Dreyfus, H. L. (1992). *What Computers Still Can't Do: A Critique of Artificial Reason*. MIT Press.
 - Floridi, L. (2010). *Information: A Very Short Introduction*. Oxford University Press.
 - Zuboff, S. (2019). *The Age of Surveillance Capitalism*. PublicAffairs.
 - Harman, G. (2002). *Tool-Being: Heidegger and the Metaphysics of Objects*. Open Court Publishing.
 - Hillman, J. (1979). *The Dream and the Underworld*. Harper & Row.
-

Reports and Metrics:

- Gallup. (2023). *Confidence in Institutions: Media Trust Polling Data*. <https://news.gallup.com>
- OpenAI. (2022–2024). *ChatGPT Release Notes*. <https://help.openai.com/en/articles/6825453-chatgpt-release-notes>
- Google DeepMind. (2023). *Reducing Hallucination in Gemini Language Models*. <https://deepmind.google/tech/research/>

BEYOND HALLUCINATION:

RETHINKING GENERATIVE AI AND TRUTH

