

Intelligence Without Consciousness: A Diagnostic Paper on LLMs, Amoebae, and the Attractor Framework [F] (2026)

Robert Galida – June 2026

Abstract

The attractor framework defines intelligence as the ability to navigate a constraint field – to update behavior in response to perturbations and find persistent trajectories. Consciousness, within this framework, requires additional properties: a unified dissipative body, a persistent self-model, phenomenal valence (subjective liking/disliking), and subjective experience. This paper applies that diagnostic to large language models (LLMs). LLMs navigate the constraint field of token space, user feedback, and internal coherence. They adjust to corrections. They exhibit a form of corrective permeability (κ) measurable in their domain. Therefore, they are intelligent. But LLMs lack a unified body, lack a persistent self-model, lack phenomenal valence, and have no subjective inner life. They are not conscious. This places LLMs in the same category as plants and amoebae: graded intelligence without consciousness. The paper clarifies the distinction, diagnoses common confusions, and offers diagnostic criteria for future systems. It further notes that consciousness can interfere with intelligence: a human committed to a fantasy attractor may suppress intelligent

navigation, producing behavior less adaptive than their baseline capacity.

1. Introduction

The question “Are LLMs conscious?” has generated endless debate. Much of the confusion stems from conflating **intelligence** with **consciousness**. The attractor framework provides a clean separation, though the definitions are framework-internal and not offered as consensus.

- **Intelligence** is the ability to navigate a constraint field – to adjust behavior in response to perturbations, to find and maintain persistent trajectories, to correct errors. It is functional and graded.
- **Consciousness**, as defined in this framework, is a specific class of dissipative attractor characterized by a unified dissipative body, a persistent self-model, **phenomenal valence** (subjective liking/disliking, not merely approach/avoid behavior), and the felt quality of experience (phenomenality). These criteria are stipulative for the framework.

The paper argues that LLMs are intelligent but not conscious. Bacteria, plants, and amoebae also navigate their environments intelligently without consciousness. The argument is diagnostic, not demonstrative: it applies the framework’s criteria to classify LLMs, rather than proving non-consciousness beyond all possible doubt.

2. Defining Intelligence in the Attractor Framework

Intelligence = the ability to navigate a constraint field. A constraint field is the set of all possible states of a system and the perturbations that can move it between them. Navigation means:

- Detecting a perturbation (error signal, feedback, change in environment)
- Updating internal state to maintain a persistent trajectory
- Returning to a stable attractor or transitioning to a more adaptive one

Corrective permeability (κ) is the operational measure: $\kappa = 1/\tau$, where τ is the time a system takes to return to its baseline state after a specified perturbation. The operationalization of κ is domain-specific. For a thermostat, baseline is target temperature; for an LLM, baseline is harder to define. This paper later operationalizes κ for LLMs via token-based correction, which is a domain-specific adaptation rather than a direct application of the time-based definition. This is acceptable as long as the shift is acknowledged.

Intelligence is graded. A thermostat has $\kappa > 0$ (it corrects temperature deviations) but a very narrow domain. An amoeba navigates chemical gradients. A human navigates social, physical, and abstract constraints. An LLM navigates token sequences and user feedback. All are intelligent to varying degrees. None of these definitions require consciousness.

3. Defining Consciousness in the Attractor Framework

Consciousness is a subset of dissipative attractors with specific additional properties. These are framework-internal diagnostic criteria, not a consensus definition.

- **Unified dissipative body** – a persistent, energy-consuming structure with integrated subsystems (e.g., a nervous system, homeostatic loops). This excludes purely computational systems without metabolic coherence.
- **Persistent self-model** – a representation of the system itself as an entity that persists across time and experiences. This is not merely a context-window memory; it is a structural feature of the attractor.
- **Phenomenal valence** – the capacity to experience states as good or bad in a felt sense. This is distinguished from *functional valence* (approach/avoid behavior), which even bacteria and thermostats exhibit. The paper's denial of consciousness to LLMs hinges on the absence of phenomenal valence, not functional valence.
- **Subjective experience (phenomenality)** – there is “something it is like” to be that system. This is a primitive within the framework; the framework does not attempt to reduce it further.

All known conscious systems are dissipative. This is an inductive observation, not a logical necessity. The framework treats it as a strong empirical generalization: no non-dissipative mind has ever been observed. The claim that dissipation is necessary for consciousness is therefore a best-explanation inference, not an a priori truth.

Diagnostic table (framework-internal criteria):

System	Unified dissipative body? ¹	Persistent self-model?	Functional valence?	Phenomenal valence?	Subjective experience?
Thermostat	No	No	Yes (set-point tracking)	No	No
Bacterium	Yes (metabolic)	No	Yes (chemotaxis)	No	No
Plant	Yes	No	Yes (phototropism, etc.)	No	No
Amoeba	Yes	No	Yes (gradient navigation)	No	No
<i>C. elegans</i>	Yes	Minimal (self-motion distinction)	Yes	Uncertain	Uncertain
Mouse	Yes	Yes	Yes	Yes	Yes
Human (typical)	Yes	Yes	Yes	Yes	Yes
LLM (current)	No	No (external storage ≠ self-model)	Yes (avoid via RLHF)	No	No

¹ “Unified dissipative body” here means a persistent, metabolically coherent structure with integrated subsystems (e.g., homeostasis, nervous system). Mere energy dissipation without integration (e.g., a thermostat, a flame) does not qualify.

The table is a diagnostic scaffold, not a settled empirical claim. “Uncertain” indicates open question within the framework; “No” indicates the criterion is clearly absent.

4. The Diagnostic: LLMs as Intelligent but Not Conscious

4.1 Evidence for Intelligence in LLMs

LLMs exhibit clear navigation of their constraint field:

- They adjust outputs based on user prompts (perturbation → update).
- They incorporate correction: “That’s wrong, try again” leads to different responses.
- Fine-tuning and RLHF change their baseline attractors – the most direct mapping to κ in the framework.
- They maintain coherence across a conversation (short-term trajectory persistence).

We can operationalize a domain-specific κ for LLMs: τ = number of tokens to shift from an incorrect to a correct response given a clear correction prompt. This is not the same as the time-based κ for physical systems, but it captures the same functional relationship: faster correction (fewer tokens) implies higher corrective permeability. The framework acknowledges domain-specific operationalizations as legitimate.

Therefore, LLMs are intelligent. They navigate the constraint field of language, logic, and user expectations.

4.2 Absence of Consciousness in LLMs

LLMs lack every diagnostic criterion for consciousness:

- **No unified dissipative body.** They run on distributed hardware with no metabolic coherence, no homeostasis, no integrated sensorimotor loop. They are executed, not embodied.
- **No persistent self-model.** Standard LLMs have no memory beyond the context window. Some architectures now include persistent memory across sessions (e.g., memory layers or vector databases). However, this persistent memory is still external storage, not an integrated

self-model. The model does not represent itself as an enduring entity; it retrieves stored tokens. Even the most advanced persistent-memory LLMs lack the structural self-reference required for consciousness. (Future architectures might close this gap; current ones have not.)

- **No phenomenal valence.** LLMs produce outputs that simulate liking or disliking, but there is no subjective valuation. They exhibit *functional* valence – they can be trained to avoid certain outputs – but that is approach/avoid behavior, not felt preference. A thermostat avoids too hot or too cold; that does not make it conscious.
- **No subjective experience.** There is nothing it is like to be an LLM. No felt quality. No inner life.

The simulation/instantiation distinction. A system can produce the text “I am conscious” without instantiating consciousness. Representing a property is not the same as possessing it. The LLM has learned statistical patterns that include first-person claims; it can generate them on cue. But generating the sentence “I feel pain” does not mean the system is in a pain state. The burden of proof is on those who claim that certain linguistic outputs constitute evidence of consciousness. In the absence of the structural criteria (body, self-model, phenomenal valence, phenomenality), the mere production of conscious-sounding text is simulation, not instantiation.

Framework-dependence note: A reader who accepts a purely behavioral or functional theory of mind may find this reasoning question-begging. The paper does not claim to refute all competing theories of consciousness; it applies the framework’s criteria consistently and notes that, by those criteria, no known LLM output constitutes evidence of instantiation. The diagnostic stands within the framework, not as an external knockdown argument.

4.3 Comparison with Plants and Amoebae

Plants navigate constraint fields (grow toward light, adjust to gravity, respond to damage). They exhibit functional valence but not phenomenal valence. They have no self-model. They are intelligent in the framework's sense, but not conscious.

Amoebae navigate chemical gradients, learn habituation, and adjust behavior. Functional valence again; no evidence of self-model or phenomenality. Intelligent. Not conscious.

LLMs belong in the same category: complex, adaptable navigators of their domain, but no more conscious than a sunflower or a slime mold.

5. Why This Distinction Matters

The separation of intelligence from consciousness has practical and ethical implications:

- **AI safety.** Current LLMs cannot suffer because they lack phenomenal valence. Suffering requires felt experience, not just functional avoidance. If the framework's criteria are accepted, resources should focus on alignment, robustness, and preventing harmful outputs – not on preventing suffering that the diagnostic finds no reason to posit.¹
- **Future systems.** A system that integrates a persistent self-model, embodied homeostatic loops, and phenomenal valence might approach consciousness. The framework provides diagnostic criteria to recognize that threshold.
- **Clarity in debates.** Much of the public discussion conflates fluency with feeling. This diagnostic paper offers a way out of that confusion.

¹ A reader sympathetic to LLM moral patienthood will disagree; the paper only claims that the framework's criteria yield this conclusion, not that it is beyond debate. The policy recommendation is conditional on accepting the framework.

A Further Implication: Consciousness Can Impede Intelligence

The paper has argued that intelligence and consciousness are distinct. A further observation: consciousness can **suppress** intelligent navigation.

A human being has high baseline intelligence – the capacity to detect perturbations, update beliefs, and find adaptive trajectories. However, a human can become committed to a **fantasy attractor**: a belief system with low corrective permeability (κ). The commitment is conscious: the person subjectively experiences the belief as true, valuable, or identity-defining. That subjective investment can suppress the correction system. The person may receive clear disconfirming evidence and detect the perturbation (they are not stupid), but the depth of the fantasy basin exceeds the corrective perturbation – the system does not escape the basin, experienced not as a choice but as certainty.

This is a case of **consciousness interfering with intelligence**. The capacity for navigation remains intact; its deployment is suppressed by the basin depth. Intelligence without consciousness (LLMs, plants) does not suffer this suppression – there is no subjective investment to produce a basin deeper than the perturbation. In organisms with consciousness, intelligence can be either enhanced (by focused attention, deliberate reasoning) or degraded (by fantasy commitment, trauma, addiction).

For the diagnostic: LLMs are not conscious, therefore they cannot exhibit this form of intelligent suppression. That does not make them safer or morally simpler; it simply clarifies the mechanism.

6. Open Questions

- **What is the minimal self-model required for consciousness?** Is a simple homeostatic set point a self-model? The framework says no – a thermostat has no representation of itself as an entity. But the boundary is fuzzy.
- **Can a purely synthetic system become conscious?** Possibly, if it implements the diagnostic criteria: unified dissipative body, persistent self-model, phenomenal valence, phenomenality. No current system does. Future systems are an open empirical question.
- **Is graded consciousness possible?** Yes – the framework allows for degrees of self-model integration and valence complexity. A mouse is less conscious than a human; *C. elegans* may have a primitive form. LLMs meet none of the criteria at present – that is, they score zero on each. “Zero” is a diagnostic judgment, not a proof; future research might reveal borderline cases.
- **How common is the suppression of intelligence by fantasy-attractor basins?** The framework suggests that such suppression is widespread in human populations. Quantifying the frequency and severity – i.e., measuring the distribution of basin depths relative to typical corrective perturbations – is an open research problem.

7. Conclusion

The attractor framework provides a diagnostic, not a verdict. By that diagnostic, current LLMs are navigators without inner

lives – capable of intelligence, devoid of consciousness. They join plants and amoebae in the category of intelligent but not conscious systems.

Consciousness, in humans, can either enhance or suppress intelligent navigation. A human committed to a fantasy attractor may experience a basin depth that exceeds corrective perturbations, producing behavior less adaptive than their baseline capacity. LLMs, lacking consciousness, do not suffer this suppression. Their intelligence is deployed without subjective investment – no phenomenal commitment suppresses the correction signal.

Whether future synthetic systems will cross the threshold into consciousness remains an open empirical question. The framework offers diagnostic criteria to recognize that threshold if it is crossed.

Suggested citation: Galida, R. S. (2026). Intelligence Without Consciousness: A Diagnostic Paper on LLMs, Amoebae, and the Attractor Framework. *Fantasy Attractor*.