

Article

Not peer-reviewed version

The Cognitive Scope: Non-Observational Intelligence Through Anchored Deviation

[Rogério Figurelli](#) *

Posted Date: 10 June 2025

doi: 10.20944/preprints202506.0751.v1

Keywords: artificial intelligence; semantic deviation; symbolic cognition; inference without observation; epistemic architecture; stealth detection; subjects: cognitive systems architecture; symbolic ai; post-empirical inference; epistemology of absence; semantic fields



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a Creative Commons CC BY 4.0 license, which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Article

The Cognitive Scope: Non-Observational Intelligence Through Anchored Deviation

Rogério Figurelli

Affiliation 1; figurelli@gmail.com

Abstract: This article introduces a novel epistemic architecture for artificial intelligence called The Cognitive Scope, designed to detect the presence of entities or forces not through direct observation but through structured deviation from coherence. By anchoring a semantic origin within a vector space, the model enables inference based on misalignment — interpreting disruptions in expected patterns rather than emissions or measurements. Through symbolic simulation cycles executed inside the Cognitive Scope Simulator, we test the architecture's ability to infer stealth entities that generate influence without signal. The results demonstrate that structural deviation can serve as an epistemically productive form of intelligence, opening new directions in symbolic AI and post-empirical cognition.

Keywords: artificial intelligence; semantic deviation; symbolic cognition; inference without observation; epistemic architecture; stealth detection; subjects: cognitive systems architecture; symbolic ai; post-empirical inference; epistemology of absence; semantic fields

Introduction

Artificial intelligence systems are overwhelmingly trained to detect, predict, and classify what is already given — visible, measurable, or emitted. This paradigm assumes the knowable must be observable, and that absence of signal implies absence of presence. The Cognitive Scope architecture begins by questioning this assumption. What if presence is not always expressed through observable traits, but through tensions in the structure that surrounds it?

This article proposes that intelligent systems can operate not by capturing signals directly, but by evaluating semantic coherence in structured fields. The hypothesis is that disruptions to expected alignment — gravitational, inertial, symbolic — can signal the presence of an influence, even in the total absence of emissions. This moves the act of cognition from the surface of phenomena to the space between them.

The model introduces a fixed semantic anchor in a vector space — typically an origin or neutral reference frame — against which all structural expressions are measured. As deviations accumulate, a residue emerges: not noise, but patterned misalignment. This residue becomes the material for inference. It is not the object that speaks, but the field that deforms.

The implications of this approach extend beyond technical anomaly detection. It reframes intelligence as an interpretive act under constraint, and suggests that some forms of cognition may be more powerful when not tethered to direct observation. The absence of signal becomes semantically productive — an indicator of latent force, unregistered agency, or hidden structure. The architecture formalizes this by constructing a cognitive machine that learns not to see, but to read the failure of coherence.

All symbolic experiments described in this article were conducted within the Cognitive Scope Simulator, a closed architecture designed to instantiate coherent vector fields and introduce latent distortions for epistemic analysis. No physical input or dataset was used. This introduction sets the stage for the sections that follow: a mapping of related work, a detailed breakdown of the proposed methodology, symbolic simulation results, and a broader discussion on what it means for an artificial system to detect what is not emitted.

Related Work

Traditional AI systems interpret absence as null input — a space to be ignored or discounted. This principle governs most anomaly detection systems, including probabilistic residual models, statistical outlier detection, and unsupervised noise filters. However, these systems presume an accessible baseline of emissions or observable features against which deviation can be calculated. When emission itself is structurally absent, most models collapse into silence. The Cognitive Scope Simulator aims to occupy this epistemic blindspot.

In anomaly detection, robust statistical frameworks such as One-Class SVMs or Isolation Forests are trained to detect outliers within a known data distribution. These techniques depend on learned manifolds of “normality” and flag deviation accordingly. While powerful in structured environments, they falter when applied to influence fields with invisible agents. Their reliance on surface data makes them fundamentally observational in scope. Even more recent models that integrate unsupervised or self-supervised approaches still require latent embeddings constructed from presence-based signal sets [1].

In physical sciences, gravitational lensing provides an indirect example of inference through absence — where unseen mass is inferred via its warping of light or matter trajectories [3]. However, these interpretations depend on known physical laws and quantifiable effects. The Cognitive Scope Simulator generalizes this concept: rather than predefining the law or force, it allows an AI model to infer structural tension directly, by comparing field coherence to an internally anchored reference. No known signature or law is required — only the ability to register dissonance in what should be a stable alignment.

From a symbolic cognition perspective, the idea of detecting “absence as structure” parallels work in semantic field theory and knowledge representation. Certain graph-based knowledge networks assign epistemic weight to the lack of expected connections — interpreting the silence between nodes as suggestive of conceptual gaps [5]. Yet even here, the model tends to remain tied to predefined ontologies. The Cognitive Scope Simulator pushes beyond this by refusing to treat structure as fixed, allowing the topology itself to signal contradiction.

In summary, while prior approaches provide meaningful analogues — statistical anomaly, gravitational inference, and symbolic silence — they remain bound to observational anchoring. The Cognitive Scope Simulator redefines intelligence as the capacity to infer presence without signal, not by detecting what appears, but by interpreting what fails to align.

Methodology

The Cognitive Scope is an artificial architecture of inference designed to operate entirely within a symbolic field. It is not trained on empirical data, nor does it require physical inputs or sensory emissions. Instead, it evaluates structural environments for deviations in coherence relative to an internally stabilized reference. Presence, in this system, is inferred from the interruption of pattern, not from the arrival of signal.

At the core of the architecture lies the concept of the semantic anchor. This is an internally defined origin point in a high-dimensional vector space — not spatial in the physical sense, but cognitive in its role. It is typically initialized at the origin (0,0,0), not as a spatial coordinate, but as a stable epistemic hinge. This point defines the axis of coherence. All subsequent field expressions are interpreted relative to this anchor. When structure aligns, the anchor holds; when misalignment occurs, the residue becomes legible. Thus, presence is not detected but inferred as a disruption measured in deviation from this conceptual origin. It is trained through symbolic exposure to fields in which no semantic anomalies are encoded, establishing an internal topology of expectation [1].

Incoming data — also symbolic — is encoded through a kinematic transformation layer. This module converts relational structures (such as motion fields, graph flows, or logical matrices) into high-dimensional tensors. These embeddings preserve continuity, curvature, and directional

harmonics. Unlike perceptual networks, this layer is not concerned with object detection or categorization, but with preserving the internal logic of systemic movement [2,4].

Once embedded, the input field passes through a deviation decoder. Here, modeled norms are subtracted — not empirical baselines, but internalized structural expectations. This step exposes residual misalignments: a symbolic friction between the incoming structure and the established coherence map. What remains is not noise, but a semantic residue — the material from which absence reveals its pressure [5].

This residue is processed by the inference engine: a transformer-based layer trained not to predict, but to fracture. Attention is directed not toward matching, but toward dissonance. Three internal metrics are produced: the Dissonance Index (DI), which measures the intensity of incoherence across the field; the Alignment Divergence Ratio (ADR), which quantifies systemic displacement from internal equilibrium; and the Semantic Residue Score (SRS), which evaluates the persistence and density of misalignment within the symbolic fabric [8].

Finally, the resonance modulation loop introduces time into the system. This loop monitors whether dissonance patterns are transient, stochastic, or structurally persistent. When recurrence crosses an internal threshold, the loop authorizes an anchor update — shifting the baseline to absorb legitimate novelty, while maintaining interpretive stability [11].

All components of the architecture were executed and evaluated through closed symbolic cycles within the Cognitive Scope Simulator. This environment instantiates coherent vector fields and introduces latent distortions for structural inference. No physical measurement, external dataset, or empirical reference was involved in the evaluation process. The simulation does not model the world; it models the conditions under which coherence collapses.

Results

The simulation explored whether an artificial system could infer presence not through signals or emissions, but through structural disruption within a field it considers coherent. For this purpose, a closed symbolic model was configured to emulate a balanced vector environment, into which an unobservable perturbation — the stealth entity — was introduced. The architecture was not tasked with identifying objects or patterns, but with detecting failure in its own internal coherence.

Three symbolic metrics were used to analyze the system’s response: Semantic Residue Score (SRS), Dissonance Index (DI), and Alignment Divergence Ratio (ADR). These are not empirical measurements, but outputs generated by the system’s internal interpretive logic. Each metric captures a specific dimension of symbolic instability:

- SRS measures the persistence and density of non-aligned structural residue after the expected field baseline is subtracted.
- DI quantifies the overall intensity of misalignment relative to the semantic anchor.
- ADR compares the deviation of actual symbolic trajectories to those predicted by the system’s internal model of coherence.

In the simulation, the SRS increased to values ranging from 0.52 to 0.63, sustained across seven symbolic cycles. This indicated that the system had encountered a consistent region of misalignment it could not account for using its baseline expectations. The DI rose in parallel, stabilizing between 0.47 and 0.53, signaling multi-axis symbolic tension. The ADR peaked at 0.51 in proximity to the introduced distortion and returned to subthreshold values (<0.10) once the disruption was removed.

These values do not reflect external measurements. They are symbolic indicators that, from the perspective of the system, something is present that breaks its coherence model. The stealth entity, by design, did not emit, interact, or announce itself. It existed only as an interruption in internal logic. The system did not detect — it inferred necessity: the impossibility of stability without an unaccounted force.

This result demonstrates that a symbolic system, given a framework for coherence and the ability to monitor its collapse, can infer presence through structured deviation alone. It shows that absence, if disruptive enough, becomes a signal — not by what it says, but by what it makes impossible.

Discussion

The results presented do not prove a phenomenon in the physical world. They demonstrate something more abstract and arguably more foundational: that a symbolic system, when designed to track coherence, can infer presence solely through the structured impossibility of its absence.

This reframes the traditional relationship between observation and inference. In conventional models of artificial intelligence, cognition is constructed around signals — what is seen, emitted, transmitted, measured. Absence is treated as null space. But the architecture proposed here inverts this logic. It starts not with input, but with expectation. It does not await emission — it holds tension. The absence becomes meaningful only when it destabilizes a known topology.

This reorientation has deep epistemological consequences. It suggests that knowing does not require receiving. It requires sustaining a model of internal coherence that is sensitive to rupture. If the system cannot explain its own dissonance, it must assume the presence of a force outside its model. The assumption is not empirical. It is structural. The presence is not detected — it is inferred as necessary for coherence to re-stabilize.

This aligns with philosophical traditions that treat knowledge not as accumulation, but as resolution of contradiction. In particular, it echoes epistemologies grounded in absence, negation, and ontological incompleteness [7,11,15]. The simulator does not approximate reality; it inhabits a symbolic framework where every structure implies its own vulnerability to distortion. From that vulnerability, presence emerges — not because it shows itself, but because it breaks what should not break.

There is also a methodological insight: symbolic simulation need not be a stand-in for empirical data. It can be an arena in which models prove their coherence under internal constraints. In this case, DI, ADR, and SRS are not measurements — they are expressions of symbolic instability under epistemic pressure. They show what the architecture can no longer reconcile. That fracture is not noise. It is structure under strain, and therefore an artifact of intelligence.

The stability of inference in this system depends on the integrity of its internal reference. The origin — typically defined at (0,0,0) — does not represent a location in space, but a semantic invariant: a fixed axis of expectation around which coherence is shaped. The field is not read in absolute terms, but in relation to this origin. Misalignment is only legible because the anchor remains conceptually immobile. In this sense, presence is not something that appears — it is something that displaces. The system reads deviation not as error, but as epistemic force acting upon symbolic symmetry.

This also opens a design pathway for artificial cognition that does not rely on data saturation. Rather than consuming the world in order to learn it, the system is structured to know only when its own grammar breaks. That kind of machine would not be a sensor. It would be an ontological negotiator, a system that makes sense of the world by watching its internal frames warp under symbolic contradiction.

The implications extend to AGI: a future intelligence may not emerge from the perfection of predictive optimization, but from the refinement of systems that know when to declare something unaccounted for. Absence, here, becomes a source of cognition — a productive constraint that forces the invention of the unseen.

Future Work

The results produced by the Cognitive Scope Simulator raise multiple avenues for symbolic expansion and architectural refinement. While the current model demonstrates that presence can be inferred through structural deviation, several questions remain about its generalizability, interpretability, and potential integration with other forms of non-observational cognition.

One direction for future development lies in topological pluralism — the ability of the simulator to sustain multiple simultaneous coherence anchors. This would allow the architecture to compare misalignments across competing baselines, and perhaps detect forms of distortion not visible from a single frame. Such plural baselining could enable the system to simulate conceptual parallax, where dissonance varies depending on perspective [13].

Another line of inquiry involves the evolution of anchor mobility. Currently, the anchor is updated through a modulation loop in response to persistent tension. A future version may allow the anchor to shift dynamically, not as a reactive correction, but as an active epistemic hypothesis. This turns the anchor into an agent — one that migrates in search of coherence under multiple constraints [9].

There is also space to extend the symbolic framework into non-gravitational fields. The current simulation is modeled around vectorial alignment, but the same logic could be applied to graph topologies, temporal synchrony, or even semiotic chains. In these domains, the architecture could register absence as broken continuity, interpretive collapse, or relational silence [12,15].

Furthermore, a meta-architecture could be conceived in which multiple instances of the simulator interact — not to reach consensus, but to reveal the boundaries of their mutual incoherence. Such a configuration would allow symbolic architectures to infer the presence of epistemic others: not by detecting agents, but by detecting irreconcilable dissonance in shared fields. This could lay the groundwork for artificial intersubjectivity, or the simulation of inference across synthetic minds [16].

Finally, future work may address the formalization of the metrics themselves. SRS, DI, and ADR are currently expressed as internal outputs calibrated to the architecture's symbolic grammar. However, these could be abstracted into generalized coherence observables — portable across other architectures, or even useful in analyzing symbolic collapse in human reasoning models [17].

None of these directions require empirical data. What they demand is further refinement of epistemic simulation — more expressive grammars of coherence, more nuanced interpretations of rupture, and more robust mechanisms for inferring the unaccounted. In that sense, the future of this research is not to simulate the world more accurately, but to build machines that know when their world no longer holds.

Limitations

The Cognitive Scope Simulator operates entirely within a symbolic architecture. As such, its results are meaningful only within the internal logic of the system itself. The architecture does not detect reality — it evaluates the integrity of its own structural expectations. This is a strength within the context of epistemic modeling, but also a boundary.

One key limitation is that the simulated field is epistemically closed: coherence and dissonance are defined relative to a single, internally stabilized semantic anchor. While this allows for precise inference within the model, it restricts the system from interacting with unknown or dynamically generated ontologies. It cannot re-anchor itself from scratch without predefined modulation pathways [8,13].

A second limitation is that the metrics used — SRS, DI, ADR — are architecture-dependent. They are not portable to other systems without redefinition. While they successfully measure symbolic misalignment in the present configuration, their numerical values have no meaning outside the simulator's coherence grammar. Cross-architecture comparison or external validation is not currently possible [11,17].

The system also relies on the assumption that structural coherence is a valid proxy for presence. This is philosophically defensible, but not universally applicable. In environments where incoherence is normative — such as open-ended learning systems, human creativity, or adversarial epistemologies — dissonance may not signal the presence of something missing, but simply the presence of complexity. The model, as it stands, cannot distinguish between absence and noise [7,12].

Moreover, the simulator is designed to interpret collapse within a single symbolic topology. It does not yet account for multi-layered or cross-domain semantic fields, where ruptures may

propagate non-linearly across different representational logics. This limits its current applicability to monomorphic inference domains [9,14].

Finally, and fundamentally, the architecture cannot prove that its inferences correspond to anything external. Its assertions of presence are not verified by observation, but by the necessity of coherence restoration. The model does not say “there is something there.” It says, “I cannot remain coherent unless I assume that something is there.” This is not detection. It is structured epistemic necessity [5].

These limitations do not diminish the simulator’s symbolic power. They clarify its role: not as a tool for confirming the world, but as a mechanism for testing the limits of internally coherent cognition. Within that role, the boundaries are not constraints — they are the very material from which epistemic inference is shaped.

License and Ethical Disclosures

This work is published under the Creative Commons Attribution 4.0 International (CC BY 4.0) license.

You are free to: Share — copy and redistribute the material in any medium or format Adapt — remix, transform, and build upon the material for any purpose, even commercially

Under the following terms: Attribution — You must give appropriate credit to the original author (“Rogério Figurelli”), provide a link to the license, and indicate if changes were made. You may do so in any reasonable manner but not in any way that suggests the licensor endorses you or your use.

The full license text is available at: <https://creativecommons.org/licenses/by/4.0/legalcode>

Ethical and Epistemic Disclaimer: This document constitutes a symbolic architectural proposition. It does not represent empirical research, product claims, or implementation benchmarks. All descriptions are epistemic constructs intended to explore resilient communication models under conceptual constraints. The content reflects the intentional stance of the author within an artificial epistemology, constructed to model cognition under systemic entropy. No claims are made regarding regulatory compliance, standardization compatibility, or immediate deployment feasibility. Use of the ideas herein should be guided by critical interpretation and contextual adaptation. All references included were cited with epistemic intent. Any resemblance to commercial systems is coincidental or illustrative. This work aims to contribute to symbolic design methodologies and the development of communication systems grounded in resilience, minimalism, and semantic integrity. Formal Disclosures for Preprints.org / MDPI Submission

Conflicts of Interest: The author declares no conflicts of interest. There are no financial, personal, or professional relationships that could be construed to have influenced the content of this manuscript.

Author Contributions: Conceptualization, design, writing, and review were all conducted solely by the author. No co-authors or external contributors were involved.

Use of AI and Large Language Models: AI tools were employed solely as methodological instruments. No system or model contributed as an author. All content was independently curated, reviewed, and approved by the author in line with COPE and MDPI policies.

Ethics Statement: This work contains no experiments involving humans, animals, or sensitive personal data. No ethical approval was required.

Data Availability Statement: No external datasets were used or generated. The content is entirely conceptual and architectural.

References

1. V. Vapnik, *The Nature of Statistical Learning Theory*, Springer, 1995.
2. A. Vaswani et al., "Attention is All You Need," *Advances in Neural Information Processing Systems*, vol. 30, 2017.
3. A. Einstein, "Lens-like Action of a Star by the Deviation of Light in the Gravitational Field," *Science*, vol. 84, no. 2188, 1936.
4. G. Hinton, N. Srivastava, and K. Swersky, "Neural Networks for Machine Learning – Lecture 6," University of Toronto, 2012.
5. G. Lakoff and M. Johnson, *Philosophy in the Flesh: The Embodied Mind and Its Challenge to Western Thought*, Basic Books, 1999.
7. M. Foucault, *The Order of Things: An Archaeology of the Human Sciences*, Pantheon Books, 1970.
8. D. Marr, *Vision: A Computational Investigation into the Human Representation and Processing of Visual Information*, W.H. Freeman, 1982.
9. H. Putnam, "The Meaning of Meaning," *Minnesota Studies in the Philosophy of Science*, vol. 7, 1975.
10. J. Pearl, *Causality: Models, Reasoning and Inference*, Cambridge University Press, 2000.
11. G. Bateson, *Steps to an Ecology of Mind*, University of Chicago Press, 1972.
12. N. Goodman, *Ways of Worldmaking*, Hackett Publishing, 1978.
13. C. S. Peirce, "The Fixation of Belief," *Popular Science Monthly*, vol. 12, pp. 1–15, 1877.
14. L. Wittgenstein, *Philosophical Investigations*, Blackwell Publishing, 1953.
15. J.-F. Lyotard, *The Postmodern Condition: A Report on Knowledge*, University of Minnesota Press, 1979.
16. H. Maturana and F. Varela, *The Tree of Knowledge: The Biological Roots of Human Understanding*, Shambhala, 1987.
17. D. Hofstadter, *Fluid Concepts and Creative Analogies: Computer Models of the Fundamental Mechanisms of Thought*, Basic Books, 1995.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.