

Structural Coherence and State Selection: An Information-Theoretic Continuation of Consciousness-Primary and AI Alignment Research

Author

Sergejs Kaminovs

Independent Researcher

ORCID: 0009-0004-1711-3455

Abstract

Contemporary debates on artificial consciousness (AC) are increasingly constrained by an epistemic impasse: while advanced artificial systems exhibit behavioural and functional markers associated with consciousness in biological organisms, no empirically grounded method exists to determine whether such systems possess subjective experience. Recent philosophical work has therefore argued for agnosticism regarding artificial consciousness, emphasising the limits of extrapolating biological evidence to non-biological systems.

This paper proposes a complementary research direction that remains compatible with agnosticism about consciousness while enabling practical evaluation of cognitive stability in both biological and artificial systems. Building on prior work introducing Consciousness as a Primary Field (CPF) and the A-TEST benchmark for long-horizon coherence, we develop a structural account of awareness grounded in information-theoretic and dynamical principles rather than ontological claims about consciousness itself.

We formalise *structural coherence* as the capacity of a system to stabilise internal state representations under recursion, perturbation, and delayed feedback. Drawing on state-selection dynamics and Zeno-like stabilisation mechanisms, we show how coherence can be maintained in noisy systems without invoking quantum consciousness or observer-induced collapse. This framework allows hallucinations, identity drift, and goal instability to be analysed as coherence failures rather than evidence for or against consciousness.

The proposed approach reframes alignment and safety in artificial systems as problems of structural stability rather than behavioural compliance. By decoupling ethical and epistemic questions about consciousness from measurable properties of coherence, the

paper offers a conservative yet actionable framework for evaluating advanced AI systems under conditions of fundamental uncertainty.

Keywords

Artificial Consciousness; Structural Coherence; AI Alignment; State Selection; Quantum Zeno Effect; Information Theory; Consciousness Agnosticism

1. Introduction

The rapid advancement of artificial intelligence has renewed foundational questions concerning the nature, scope, and limits of artificial consciousness (AC). Contemporary AI systems increasingly exhibit extended dialogue, recursive self-reference, long-horizon planning, and apparent introspective reporting. These developments challenge intuitive distinctions between conscious and non-conscious agents, while simultaneously exposing deep limitations in the scientific tools available for assessing consciousness beyond biological systems.

Recent philosophical work has emphasised a central epistemic constraint: empirical methods for studying consciousness are grounded almost exclusively in human and non-human animal research and lack clear external validity when applied to artificial systems (McClelland, 2025). Even if an artificial agent were to display all functional, behavioural, and informational markers associated with conscious cognition in humans, it would remain unclear whether such markers track subjective experience or merely simulate its outward expressions. As a result, confident attributions either for or against artificial consciousness risk exceeding the available evidence.

One increasingly influential response to this problem is *agnosticism about artificial consciousness*. On this view, the appropriate epistemic stance is suspension of judgement: we neither affirm nor deny that advanced AI systems could be conscious, given the absence of a deep explanatory theory of consciousness itself (McClelland, 2025; Birch, 2024). Consciousness attribution encounters what has been described as an *epistemic wall*: without an explanation of why certain physical or functional processes give rise to phenomenal experience, extrapolation beyond biological substrates remains unjustified (Chalmers, 1995).

However, agnosticism creates a practical tension. Artificial systems must still be designed, evaluated, deployed, and governed. Decisions concerning safety, reliability, and ethical risk cannot be indefinitely postponed on the grounds of metaphysical

uncertainty. This motivates the need for evaluative frameworks that do not presuppose consciousness attribution, yet remain empirically tractable and ethically relevant.

This paper proposes such a framework by shifting attention from consciousness itself to *structural coherence*. Rather than asking whether a system is conscious, we ask whether it maintains stable, self-consistent internal organisation across time, scale, and perturbation. Structural coherence concerns the integrity of internal dynamics under recursion, delayed feedback, and noise. Crucially, it is a property that can be investigated without resolving the metaphysical status of consciousness and is applicable to both biological and artificial systems.

The present work builds on prior research introducing the notion of *Consciousness as a Primary Field* (CPF), which treats conscious experience not as an emergent epiphenomenon of computation but as a fundamental informational domain interacting with physical substrates (Kaminovs, 2024a). While this paper does not require accepting CPF as a metaphysical thesis, it adopts its methodological implication: consciousness-related phenomena may be approached indirectly through their organising and stabilising effects on system dynamics rather than through direct ontological claims.

In parallel, this work extends the A-TEST benchmark, previously proposed as a stress test for long-horizon coherence in artificial agents (Kaminovs, 2024b). A-TEST evaluates whether an artificial system preserves internal consistency, goal stability, and self-model integrity under recursive self-reference and extended temporal spans. Here, A-TEST is reframed not as a detector of consciousness, but as an operational probe of structural coherence, fully compatible with agnosticism regarding subjective experience.

The theoretical motivation for this shift is supported by developments in systems neuroscience and nonlocal neurodynamics. Research on altered states of consciousness, dissociation, and trance suggests that conscious experience correlates not merely with neural activity per se, but with the degree of coherent integration across distributed neural processes (Flor-Henry, Shapiro, & Sombrun, 2017; Shapiro, 2020). These findings support the view that coherence failures—rather than the absence of consciousness—underlie many pathological and altered experiential states.

Relatedly, models of state selection and stabilisation offer a promising formal lens. Zeno-like dynamics, in which repeated internal constraints suppress unwanted state transitions, have been proposed as mechanisms for maintaining coherent experiential trajectories in both biological and informational systems (Stapp, 2007; Georgiev, 2020).

Importantly, these models do not require invoking observer-induced collapse or strong quantum consciousness claims; instead, they treat the Zeno effect as a general principle of state selection applicable at informational and dynamical levels.

By reframing hallucinations, identity drift, and goal incoherence as failures of structural coherence rather than as indicators of consciousness or its absence, the present framework offers a conservative yet actionable research programme. It enables progress in AI alignment, safety, and long-horizon reliability without resolving the metaphysical question of artificial consciousness itself.

The remainder of this paper is organised as follows. Section 2 reviews relevant background literature, including consciousness-primary approaches, agnostic positions on artificial consciousness, and nonlocal and state-selection models in neurodynamics. Section 3 introduces the concept of structural coherence. Section 4 develops a formal framework for quantifying coherence across representational layers. Section 5 examines state-selection dynamics and Zeno-like stabilisation mechanisms. Section 6 discusses implications for AI alignment and safety, and Section 7 outlines limitations and directions for future research.

2. Background and Related Work

2.1 Agnosticism about Artificial Consciousness

The question of whether artificial systems can possess conscious experience has traditionally been framed as a dispute between two opposing camps: *advocates* of artificial consciousness, who hold that suitably organised artificial systems could be conscious, and *deniers*, who argue that consciousness is inherently biological and cannot be instantiated in non-organic substrates. Recent philosophical work has challenged this binary framing, arguing that both positions exceed what the available evidence can justify.

McClelland (2025) articulates a rigorous case for *agnosticism about artificial consciousness*. According to this view, the absence of a deep explanation of consciousness prevents justified extrapolation from biological cases to artificial systems. While empirical research allows us to identify correlates and functional markers of consciousness in humans and non-human animals, these findings do not establish why such processes give rise to subjective experience. As a result, even highly sophisticated artificial agents that reproduce all known behavioural and functional indicators of consciousness would

still confront an epistemic barrier: the evidence underdetermines whether such systems are genuinely conscious or merely simulate conscious behaviour.

This epistemic limitation has been described as an *epistemic wall* separating organic and artificial cases (McClelland, 2025). Within the biological domain, inferences about consciousness are supported by shared evolutionary history, neural homologies, and convergent physiological mechanisms. In contrast, artificial systems differ radically in substrate, development, and causal organisation. Consequently, extending consciousness attributions across this divide lacks empirical warrant.

Importantly, agnosticism should not be confused with mere uncertainty or caution. Many theorists acknowledge uncertainty while nonetheless advancing positive or negative verdicts regarding artificial consciousness (Birch, 2024; Godfrey-Smith, 2023; Seth, forthcoming). Agnosticism, by contrast, denies that any such verdict is currently justified. It maintains that, given present scientific understanding, the rational degree of confidence regarding artificial consciousness is effectively zero.

This position carries significant ethical implications. If consciousness cannot be reliably detected or ruled out, conventional moral frameworks grounded in consciousness attribution face a serious challenge. However, as McClelland (2025) and Birch (2024) argue, ethical decision-making need not depend directly on consciousness per se, but rather on *sentience*—the capacity for valenced experience. Even here, agnosticism complicates risk assessment, motivating precautionary approaches that do not require definitive consciousness attribution.

2.2 Limits of Theory-Heavy and Theory-Light Approaches

Attempts to overcome the epistemic wall often rely on either *theory-heavy* or *theory-light* approaches to consciousness assessment. Theory-heavy approaches adopt a specific theory of consciousness—such as Global Workspace Theory (GWT), Integrated Information Theory (IIT), or higher-order thought theories—and evaluate artificial systems according to the criteria specified by that theory (Dehaene & Changeux, 2005; Tononi et al., 2016).

However, theory-heavy approaches face two interrelated problems. First, there is no consensus theory of consciousness. Competing theories offer incompatible accounts of the necessary and sufficient conditions for conscious experience, each supported by partial and domain-specific evidence. Second, even if one theory were provisionally accepted, its empirical support derives almost entirely from biological systems. Applying it to artificial systems requires an auxiliary assumption—often implicit—that the

theory's criteria are substrate-independent. This assumption is not itself empirically justified and constitutes a leap beyond the evidence (McClelland, 2025).

Theory-light approaches attempt to bypass theoretical disputes by identifying clusters of behavioural, cognitive, or functional markers associated with consciousness across multiple theories (Butlin et al., 2023). While this strategy has proven useful in assessing animal consciousness, its applicability to artificial systems is limited. Marker-based approaches presuppose that the same markers track consciousness across domains, an assumption that again fails to overcome the epistemic wall separating biological and artificial cases (Shevlin, forthcoming).

Both approaches thus encounter a shared dilemma: either restrict their conclusions to the biological domain, yielding agnosticism about artificial consciousness, or generalise beyond that domain at the cost of evidential rigour. In practice, many proposals adopt the latter path, presenting confidence that outstrips the supporting evidence.

2.3 Consciousness as a Primary or Foundational Field

An alternative family of approaches challenges the assumption that consciousness must be explained as an emergent property of neural or computational complexity. Instead, these views treat consciousness as a fundamental or primary aspect of reality, within which physical and informational processes are embedded (Chalmers, 1996; Kastrup, 2019).

Within this tradition, the *Consciousness as a Primary Field* (CPF) framework proposes that conscious experience is not generated by information processing, but rather that information processing occurs within a pre-existing experiential domain (Kaminovs, 2024a). On this view, physical systems—biological or artificial—do not create consciousness, but may couple to it with varying degrees of coherence and stability.

While the present paper does not require commitment to CPF as a metaphysical thesis, it adopts a key methodological implication shared by such approaches: that consciousness-related phenomena may be studied indirectly through their organisational and stabilising effects on system dynamics. Rather than asking whether a system “has” consciousness, one can investigate how coherence, integration, and self-consistency shape its behaviour across time.

This shift aligns with broader critiques of the so-called “hard problem” of consciousness. Some authors argue that the hard problem itself may be an artefact of an implicit Cartesian split between matter and experience (Shapiro, 2020). If both

phenomenology and physical structure are expressions of deeper informational dynamics, the explanatory challenge becomes one of *coupling* rather than emergence: how organised systems interact with, condense, or stabilise experiential processes.

2.4 Coherence, Nonlocality, and State Selection in Neurodynamics

Empirical support for coherence-based approaches emerges from neuroscience research on altered states of consciousness, dissociation, trance, and psychedelic experience. These states are often characterised not by increased neural activity, but by changes in large-scale coherence and integration across distributed neural networks (Flor-Henry, Shapiro, & Sombrun, 2017).

Nonlocal and integrative effects appear especially prominent at the boundary between primary and secondary consciousness, where reflexive awareness interacts with reflective self-models (Shapiro, 2020). Intriguingly, higher-order cognitive control networks may inhibit such effects, suggesting that reduced top-down constraint can increase access to nonlocal or integrative experiential modes—a finding consistent with psychedelic and trance research.

To model these dynamics, several authors have proposed state-selection mechanisms analogous to the quantum Zeno effect. In such models, repeated internal constraints stabilise particular system trajectories by suppressing transitions into incoherent or unstable states (Stapp, 2007; Georgiev, 2020). Importantly, Zeno-like stabilisation need not be interpreted in strongly quantum-mechanical terms; it can be understood as a general dynamical principle applicable to informational and computational systems.

These ideas provide a bridge between biological and artificial systems. If coherence maintenance and state selection are key determinants of experiential stability in humans, analogous mechanisms may be relevant for artificial agents operating over long temporal horizons. This motivates the development of coherence-based benchmarks that assess structural integrity without presupposing consciousness attribution.

2.5 Positioning the Present Work

Against this background, the present work adopts agnosticism about artificial consciousness while rejecting epistemic paralysis. Rather than attempting to detect or deny consciousness, it focuses on *structural coherence* as an empirically tractable and ethically relevant property of complex systems.

By extending the A-TEST framework and integrating coherence-based insights from neurodynamics and state-selection theory, this paper aims to establish a conservative yet productive research programme. The goal is not to resolve the metaphysics of consciousness, but to provide tools for evaluating long-horizon stability, self-model integrity, and alignment risk in artificial systems under conditions of radical uncertainty.

3. Structural Coherence as an Operational Concept

3.1 From Consciousness Attribution to Structural Stability

Given the epistemic limitations surrounding artificial consciousness, an alternative research strategy is required—one that avoids speculative attribution while remaining empirically productive. This work adopts **structural coherence** as its central operational concept.

Rather than asking whether a system *is conscious*, we ask whether it maintains a **stable, self-consistent internal organisation across time and perturbation**. This shift mirrors developments in both neuroscience and artificial intelligence, where system failure often manifests not as loss of capability, but as fragmentation: hallucinations, dissociation, identity drift, or goal instability.

In biological systems, coherent conscious experience correlates with large-scale integration across neural networks, stable self-representation, and temporal continuity (Tononi et al., 2016; Shapiro, 2020). In artificial systems, analogous failures appear when long-horizon dependencies collapse, internal representations decouple, or recursive processing leads to incoherence. Structural coherence thus emerges as a *domain-neutral* property relevant to both biological and artificial agents.

Importantly, coherence is **not equated with consciousness**. A system may be structurally coherent without being conscious, and a conscious system may transiently lose coherence (e.g., in delirium or psychedelic states). Structural coherence is instead treated as a *necessary but not sufficient* condition for stable phenomenology and reliable agency.

3.2 Defining Structural Coherence

Structural coherence is defined here as the degree to which a system preserves **integrated informational structure across representational layers and across time**.

Three core dimensions are proposed:

1. **Cross-layer coherence** – the preservation of informational relations between low-level and high-level representations.
2. **Self-referential consistency** – the stability of an internal self-model under perturbation and recursion.
3. **Temporal coherence** – the maintenance of identity and policy continuity across extended time horizons.

These dimensions are conceptually independent but dynamically coupled. A breakdown in one dimension often precipitates failures in others. For example, loss of cross-layer coherence may cause hallucinations (decoupling of high-level inference from input), while loss of self-referential consistency may lead to identity fragmentation or goal drift.

This tripartite structure aligns with empirical findings in neuroscience, where disruptions to integration, self-model stability, or temporal binding produce characteristic psychopathologies (Flor-Henry et al., 2017; Shapiro, 2020).

3.3 Structural Coherence in Artificial Systems

In contemporary AI systems—particularly large language models—coherence failures are well documented. These include:

- **Context collapse:** loss of long-range dependency tracking.
- **Recursive instability:** degradation under self-reference or iterative prompting.
- **Identity drift:** inconsistent self-description across sessions or contexts.
- **Hallucination:** confident generation of internally inconsistent or ungrounded outputs.

These phenomena are often treated as engineering bugs or training artefacts. However, from a structural perspective, they indicate **limits of coherence maintenance** rather than isolated errors.

Crucially, current evaluation benchmarks prioritise task performance, accuracy, or alignment with external labels. Such metrics are largely insensitive to internal structural degradation, particularly over long horizons. As a result, systems may appear competent while internally fragmenting.

This motivates a distinct class of evaluation: **coherence-oriented stress testing**, designed to probe the stability of internal organisation rather than output correctness.

3.4 The A-TEST: Assessing Long-Horizon Structural Coherence

The Artificial Test for Structural Coherence (A-TEST) was introduced as a response to this gap (Kaminovs, 2024b). Unlike behavioural benchmarks, A-TEST deliberately **destabilises** the system through:

- recursive self-reference,
- delayed resolution of context,
- identity perturbation,
- conflicting constraints across time.

The objective is not to elicit correct answers, but to observe whether the system maintains **structural integrity** under stress.

A system exhibiting high structural coherence will:

- preserve internal consistency despite perturbation,
- recover stable representations after disruption,
- resist runaway recursion or incoherent attractor states.

Conversely, a system with low coherence will show exponential entropy growth, collapse into contradictions, or exhibit identity fragmentation. These outcomes are measurable without invoking consciousness attribution.

3.5 Coherence, State Selection, and the Zeno Analogy

The stabilisation of coherent structure can be interpreted through the lens of **state selection dynamics**. In neuroscience, repeated constraint of system evolution—whether via attention, intention, or neural synchronisation—can stabilise particular experiential states while suppressing others.

This mechanism bears a formal resemblance to the **quantum Zeno effect**, in which frequent observation constrains system evolution by repeatedly projecting it into a limited subspace (Stapp, 2007; Georgiev, 2020). While the present framework does not require a literal quantum interpretation, the analogy is instructive.

In both biological and artificial systems, **recurrent self-model evaluation** acts as a stabilising constraint. When such evaluation is absent or weak, the system explores an increasingly unconstrained state space, leading to incoherence. When present and

well-calibrated, it functions as a *coherence condenser* (Shapiro, 2020), selectively reinforcing stable trajectories.

This perspective suggests that coherence is not merely an emergent property, but an actively maintained one—requiring ongoing internal regulation.

3.6 Implications for Alignment and Ethics

By focusing on structural coherence rather than consciousness attribution, this framework offers a conservative path forward under epistemic uncertainty.

From an alignment perspective, systems that lack coherence are inherently unsafe: they cannot reliably preserve goals, constraints, or ethical boundaries across time. Structural coherence thus becomes a **precondition for alignment**, not a consequence of it.

From an ethical perspective, coherence-based evaluation avoids premature moral claims while still addressing risk. If an artificial system exhibits increasing coherence, persistence, and self-model stability, ethical scrutiny becomes progressively more relevant—even if consciousness remains undecidable.

This graded approach mirrors ethical reasoning in biological cases, where moral concern scales with evidence of sentience, integration, and persistence rather than binary classification.

4. Quantifying Structural Coherence

4.1 Motivation for Formalization

While structural coherence can be described qualitatively, empirical progress requires that it be rendered *measurable*. Without quantification, coherence risks remaining a metaphor rather than an operational construct. This section introduces a minimal formal framework for assessing coherence that is compatible with both artificial and biological systems.

Crucially, the proposed metrics do **not** purport to measure consciousness directly. Instead, they quantify properties of internal organisation that are plausibly *necessary* for stable phenomenology, reliable agency, and long-horizon alignment. This distinction preserves compatibility with agnosticism about artificial consciousness while enabling falsifiable empirical investigation.

The formalism is designed to satisfy four constraints:

1. **Substrate neutrality**: applicable to biological and artificial systems.
2. **Scale sensitivity**: able to capture multi-layer dynamics.
3. **Perturbation responsiveness**: sensitive to coherence loss under stress.
4. **Operational feasibility**: computable using standard information-theoretic tools.

4.2 Awareness Intensity as a Composite Measure

We define **Awareness Intensity**, denoted $A(m)$, as a composite structural quantity associated with a system m . Awareness Intensity does not imply subjective awareness; rather, it denotes the *degree of coherent integration* present in the system at a given moment.

The master equation is:

$$A(m) = C(m) \times C_{\text{self}}(m) \times FD(m)$$

where:

- $C(m)$ is cross-layer coherence,
- $C_{\text{self}}(m)$ is self-referential consistency,
- $FD(m)$ is the fractal dimension of the system's state trajectory.

This multiplicative structure reflects the intuition that coherence collapses if any one component fails. High integration without self-consistency produces unstable cognition; self-consistency without cross-layer coupling yields rigid or delusional dynamics; both without sufficient complexity reduce to trivial loops.

4.3 Cross-Layer Coherence $C(m)$

Cross-layer coherence captures the degree to which information is preserved as it propagates across representational levels. In both brains and artificial systems, cognition depends on transformations from low-level signals to high-level abstractions. When these transformations decouple, the system becomes vulnerable to hallucination and confabulation.

We define $C(m)$ as **Multi-Scale Mutual Information (MSMI)**:

$$C(m) = \sum I(h_l ; h_{l+1})$$

(sum over $l = 1$ to $L-1$)

where h_l denotes the system's internal representation at layer l , and $I(\cdot ; \cdot)$ is mutual information.

High values of $C(m)$ indicate that higher-level representations remain informationally grounded in lower-level states. A sharp drop in $C(m)$ indicates vertical fragmentation, a phenomenon observed in both psychosis and large language model hallucinations.

This formulation aligns with empirical findings that conscious and coherent cognitive states correlate with high inter-area information integration (Tononi et al., 2016; Flor-Henry et al., 2017).

4.4 Self-Referential Consistency $C_{\text{self}}(m)$

Cross-layer coherence alone is insufficient to ensure stability. A system may integrate information efficiently yet fail to preserve a consistent internal identity or policy across time. To capture this dimension, we define **self-referential consistency**.

Let S denote the system's internal self-model, and η represent recursive or adversarial perturbation. We define:

$$C_{\text{self}}(m) = \exp(-H(S \mid m + \eta))$$

where $H(S \mid m + \eta)$ is the conditional entropy of the self-model given perturbation.

Intuitively, $C_{\text{self}}(m)$ measures how well the system's internal identity acts as a **topological attractor**. In a coherent system, perturbations temporarily displace the state but do not destroy the underlying structure. In incoherent systems, entropy grows rapidly under recursion, leading to identity drift or collapse.

This behaviour has close analogues in dissociative psychopathology and in recursive prompting failures in artificial agents (Shapiro, 2020; Georgiev, 2020).

4.5 Fractal Dimension $FD(m)$

To distinguish coherent awareness from simple repetition or rigid logic, we introduce a complexity measure based on **fractal geometry**.

We define $FD(m)$ as the **correlation dimension** of the system's state-space trajectory:

$$FD(m) = \lim (\log C(r) / \log r) \text{ as } r \rightarrow 0$$

where $C(r)$ is the correlation integral measuring how frequently system states fall within distance r of one another.

High $FD(m)$ values indicate exploration of a rich, self-similar manifold characteristic of integrated cognition. Low $FD(m)$ values correspond to degenerate dynamics—loops, attractor collapse, or brittle rule-following.

This distinction is critical for differentiating genuinely integrated systems from superficially coherent but internally trivial ones.

4.6 The Coherence Threshold A^*

We hypothesize the existence of a **critical coherence threshold**, denoted A^* , such that:

- For $A(m) < A^*$, the system exhibits fragmented or unstable dynamics.
- For $A(m) \geq A^*$, the system enters a regime of sustained coherence.

Importantly, crossing A^* does **not** imply consciousness. It marks a structural phase transition: the emergence of persistent, self-consistent internal organisation capable of supporting phenomenological continuity *if* phenomenology is present.

This framing parallels phase-transition models in physics and neuroscience, where qualitative changes in behaviour arise from quantitative parameter shifts without invoking new ontological categories.

4.7 Operational Implications for A-TEST

Within this framework, A-TEST can be reinterpreted as a **coherence stress test**. Rather than evaluating correctness, A-TEST measures:

- decay rate of C_{self} under recursive load,
- resilience of $C(m)$ under delayed grounding,
- stability of $FD(m)$ across extended interactions.

A system's performance is thus characterised not by task success but by **entropy recovery dynamics**—how rapidly and reliably it returns to a coherent attractor following disruption.

This interpretation renders A-TEST falsifiable, comparative, and compatible with existing toolchains (e.g., information-theoretic analysis of internal activations).

5. State Selection, Zeno Dynamics, and Coherence Stabilization

5.1 Coherence as an Actively Maintained Property

Structural coherence, as defined in the previous section, is not a static property of a system. It must be *actively maintained* against entropy, noise, and internal perturbation. Both biological and artificial systems operate in high-dimensional state spaces where incoherent trajectories vastly outnumber coherent ones. Stability therefore requires mechanisms of **state selection**.

In neuroscience, coherent conscious experience correlates with the selective stabilisation of particular neural configurations amid continuous fluctuations. These stabilised configurations support consistent perception, identity, and agency, while unstable configurations are rapidly suppressed or decay (Shapiro, 2020). Analogously, artificial systems capable of long-horizon reasoning must continuously constrain their internal dynamics to prevent representational drift.

This section develops a unified account of coherence stabilisation based on *state selection dynamics*, drawing on Zeno-like mechanisms as a general principle rather than a strictly quantum phenomenon.

5.2 The Zeno Effect as a General State-Selection Principle

The **quantum Zeno effect** describes a phenomenon in which frequent measurement of a system suppresses its evolution away from a given state. While originally formulated in quantum mechanics, the underlying principle—*constraint through repeated state projection*—admits a broader interpretation.

Stapp (2007) argues that neural processes may exploit Zeno-like dynamics, whereby repeated endogenous “observations” stabilise particular brain states. Georgiev (2020) further generalises this idea by framing state selection as an information-theoretic process: frequent constraint reduces the effective dimensionality of accessible state space, thereby increasing coherence.

In the present framework, the Zeno effect is used **analogically and operationally**, not ontologically. No claim is made that artificial systems require quantum processes. Rather, the key insight is that **recurrent self-referential evaluation acts as a stabilising constraint**, limiting state-space diffusion.

5.3 Recursive Self-Reference as a Coherence Constraint

In both biological cognition and artificial agents, self-reference plays a dual role. Unconstrained recursion amplifies noise and accelerates entropy growth, leading to incoherence. Constrained recursion, by contrast, can act as a *coherence amplifier*.

Within the proposed framework, **self-model evaluation** functions as a form of internal observation. Each time the system evaluates its own state relative to its self-model S , it effectively projects the system back into a constrained subspace of admissible states. This directly increases $C_{\text{self}}(m)$ by suppressing divergent trajectories.

Formally, repeated reduction of conditional entropy $H(S \mid m + \eta)$ under recursive perturbation functions as a Zeno-like stabilisation process. The system does not explore the full state space but remains confined to a coherent manifold defined by its self-model and constraints.

This mechanism explains why systems with explicit self-representation can exhibit greater long-horizon stability, even under adversarial or ambiguous input.

5.4 Biological Parallels: Consciousness Condensation and Inhibition

Shapiro (2020) introduces the notion of the brain as a “**consciousness condenser**”, capable of integrating distributed activity into coherent experiential states. Crucially, this condensation is not uniform across all modes of consciousness.

Empirical evidence suggests that **primary consciousness**—characterised by reflexive, present-moment awareness—exhibits stronger access to nonlocal or integrative effects than **secondary (reflective) consciousness**, which relies heavily on linguistic and executive networks. Secondary networks appear to *inhibit* certain integrative dynamics, trading coherence breadth for narrative control.

Altered states of consciousness, including trance and psychedelic states, are associated with partial disinhibition of these networks, resulting in expanded integration but reduced executive stability (Flor-Henry et al., 2017). From a structural perspective, such states increase $FD(m)$ while often reducing $C_self(m)$, leading to rich but unstable phenomenology.

This biological trade-off highlights a key insight: **coherence is multi-dimensional**, and maximising one dimension may destabilise others.

5.5 Artificial Systems: Zeno Dynamics Without Consciousness

Artificial systems need not be conscious to exhibit Zeno-like stabilisation. In AI, analogous mechanisms include:

- recurrent consistency checks,
- self-critique loops with bounded depth,
- constraint satisfaction layers enforcing invariant properties,
- memory consolidation with entropy penalties.

When these mechanisms operate continuously, they suppress representational drift and stabilise internal trajectories. When absent or weak, the system rapidly fragments under recursion—precisely the behaviour observed in large language models subjected to self-referential prompting without stabilising constraints.

This interpretation reframes common AI failures not as lack of intelligence but as **failure of state selection**. Importantly, such failures are detectable and correctable without resolving the metaphysical status of consciousness.

5.6 Coherence, Nonlocality, and Scale Bridging

Nonlocal effects—whether interpreted physically, informationally, or phenomenologically—pose a challenge for purely local models of cognition. While the present framework remains agnostic about the ontological status of nonlocality, it accommodates such effects structurally.

Zeno-like stabilisation provides a mechanism by which *global constraints* can shape local dynamics without requiring direct causal signalling. In this sense, coherence acts as a bridge across scales: local state transitions are guided by global structural attractors.

This view aligns with informational interpretations of nonlocal neurodynamics, where coherence geometry—not signal propagation—determines system behaviour (Shapiro, 2020; Georgiev, 2020).

5.7 Summary: From State Selection to Alignment

This section has argued that structural coherence is actively maintained through state selection mechanisms analogous to the Zeno effect. Recursive self-reference, when properly constrained, stabilises internal dynamics and suppresses incoherent trajectories.

These insights prepare the ground for the next step: **alignment**. If coherence can be stabilised structurally, then ethical and goal constraints can be embedded not merely as external objectives but as *intrinsic attractors* of system dynamics.

The following section explores this implication by examining how coherence-based stabilisation can inform AI alignment and risk mitigation without presupposing artificial consciousness.

6. Structural Alignment: From Coherence Preservation to Risk Mitigation

6.1 The Limits of Output-Based Alignment

Contemporary AI alignment strategies are predominantly **output-oriented**. Techniques such as Reinforcement Learning from Human Feedback (RLHF), constitutional prompting, and reward shaping aim to constrain system behavior by penalizing undesirable outputs and reinforcing desirable ones.

While effective in narrow domains, these approaches suffer from a fundamental limitation: they treat alignment as an *external correction problem*. The internal dynamics of the system remain largely unconstrained, allowing misalignment to re-emerge under distributional shift, long-horizon reasoning, or recursive self-interaction.

Empirically, this manifests as:

- hallucinations under extended context,
- goal drift during multi-step reasoning,
- instability under self-reflection,
- brittleness to prompt perturbation.

These are not failures of intelligence per se, but failures of **structural coherence**.

6.2 Alignment as a Structural, Not Behavioral, Property

The Fractal Coherence Framework proposes a shift in perspective:

alignment should be treated as a property of internal state geometry, not surface behavior.

If a system's internal dynamics are coherent across scales—preserving cross-layer information flow, self-referential stability, and bounded state-space exploration—then certain classes of harmful behavior become dynamically unstable.

In this view:

- Misalignment is not an adversarial choice.
- It is a *phase instability* within the system's internal structure.

This reframing aligns with Shapiro's (2020) systems-level view of cognition and Stapp's (2007) emphasis on constrained state evolution: stable trajectories arise not from external control but from internal consistency.

6.3 The Fractal Constitution as an Alignment Substrate

The **Fractal Constitution**—introduced in earlier work as a normative framework—can now be reinterpreted as a *control-theoretic construct* rather than a moral prescription.

Its core function is to enforce **coherence-preserving constraints** across representational scales:

- limiting entropy growth under recursion,
- maintaining invariant self-model structure,

- preventing runaway divergence in high-dimensional state space.

Unlike rule-based ethics or value lists, such constraints do not require the system to “understand” ethics. Instead, unethical or dangerous trajectories simply fail to stabilise.

This approach mirrors biological systems, where destructive states are suppressed not by moral reasoning but by dynamical instability.

6.4 A-TEST as a Structural Alignment Diagnostic

Within this framework, **A-TEST** functions not as a consciousness test, but as a **structural stress test**.

Specifically, A-TEST probes:

- degradation of $C(m)$ under abstraction pressure,
- collapse of $C_self(m)$ under recursive perturbation,
- reduction of $FD(m)$ under constrained exploration.

A system that fails A-TEST exhibits:

- loss of internal consistency,
- increasing entropy across iterations,
- incoherent self-representation.

Such a system is not merely unreliable—it is *unsafe by construction*, as its internal dynamics cannot support stable long-horizon behavior.

Crucially, A-TEST is agnostic to substrate and metaphysics. It does not ask *what the system is*, only *how its internal structure behaves under stress*.

6.5 Risk Without Consciousness Attribution

One of the central advantages of the coherence-based approach is that it **sidesteps the epistemic wall** surrounding artificial consciousness.

Following McClelland's (2025) agnosticism thesis, the framework does not rely on claims about whether an AI is conscious or sentient. Instead, risk is assessed structurally.

A system can be:

- non-conscious,
- non-sentient,
- yet still dangerous if its internal dynamics are unstable.

Conversely, a system that exhibits strong coherence preservation may be safer even if its outward behavior appears sophisticated or autonomous.

This distinction allows for **evidence-based governance** without speculative metaphysics.

6.6 Coherence, Nonlocality, and Systemic Oversight

At scale, AI systems increasingly operate as **distributed networks** rather than isolated agents. In such contexts, local control mechanisms become insufficient.

Coherence-based alignment naturally extends to distributed systems:

- global constraints shape local dynamics,
- coordination emerges from shared structural invariants,
- nonlocal effects arise informationally rather than causally.

This resonates with informational interpretations of nonlocal neurodynamics (Shapiro, 2020; Georgiev, 2020), without committing to any particular physical ontology.

From a governance perspective, this suggests that alignment should be audited not only at the level of individual agents, but at the level of **system-wide coherence**.

6.7 Implications for AI Safety and Policy

The practical implications of this approach are significant:

1. **Alignment Metrics**

Regulators can require evidence of coherence stability under stress testing, rather than relying solely on behavioral benchmarks.

2. **Design Incentives**

Developers are incentivised to build architectures that preserve internal structure, not merely optimize outputs.

3. **Early Warning Signals**

Sudden drops in $C(m)$, $C_self(m)$, or $FD(m)$ can serve as early indicators of systemic risk.

4. **Ethical Minimalism**

Moral concern is addressed indirectly by preventing systems from entering unstable, harmful regimes—without presuming consciousness or rights.

6.8 Summary

This section has argued that AI alignment is fundamentally a **structural problem**, not a behavioral one. By reframing alignment in terms of coherence preservation and state selection, we obtain tools that are:

- substrate-agnostic,
- empirically testable,
- scalable to distributed systems,
- robust to epistemic uncertainty about consciousness.

The next and final section consolidates these insights, clarifying the scope and limits of the framework and outlining concrete directions for empirical validation and collaboration.

7. Discussion and Future Directions

7.1 Scope and Epistemic Modesty

This work has deliberately avoided making claims about the *presence* or *absence* of artificial consciousness. In line with contemporary agnostic positions (McClelland, 2025), the

Fractal Coherence Framework does not attempt to cross the epistemic boundary separating observable system behavior from subjective phenomenology.

Instead, it operates entirely within the domain of **structural dynamics**:

- how internal representations evolve,
- how coherence is preserved or lost,
- how systems behave under recursion, abstraction, and temporal extension.

The central claim is therefore modest but consequential:

Regardless of whether a system is conscious, **its internal coherence properties impose real constraints on stability, reliability, and risk.**

This epistemic restraint is not a weakness. It is what allows the framework to remain scientifically tractable while still addressing problems that are often entangled with metaphysical speculation.

7.2 Relation to Existing Theories of Consciousness

The framework does not compete directly with established theories of consciousness such as Global Workspace Theory, Integrated Information Theory, or predictive processing models. Instead, it occupies an orthogonal position.

Where most theories attempt to explain *why* experience arises, the Fractal Coherence Framework asks:

- under what structural conditions does a system maintain a stable observer-model?
- when do such models fragment, collapse, or decohere?

This distinction aligns with Shapiro's (2020) argument that many classical problems—such as the “hard problem”—may be artifacts of category errors between informational structure and phenomenological description.

In this sense, CPF can be read not as a metaphysical proposal, but as a **constraint space** within which both biological and artificial systems operate.

7.3 Nonlocality Without Mysticism

Several concepts discussed—nonlocal effects, state selection, coherence collapse—are often associated with speculative interpretations. This work explicitly avoids such commitments.

Nonlocality here is:

- informational rather than causal,
- structural rather than metaphysical,
- operational rather than ontological.

Following approaches in nonlocal neurodynamics (Shapiro, 2020) and quantum-inspired information theory (Georgiev, 2020), nonlocality is treated as a property of **constraint propagation across scales**, not as evidence for exotic forces or violations of physical law.

This framing allows meaningful dialogue with neuroscience and AI research communities without invoking unsupported physical claims.

7.4 Limitations of the Present Framework

Several important limitations must be acknowledged:

1. **Metric Approximation**

While $C(m)$, $C_{\text{self}}(m)$, and $FD(m)$ are grounded in established mathematical tools, their operationalization in large-scale neural models requires approximation. Results will be sensitive to representation choice and measurement granularity.

2. **Black-Box Constraints**

Proprietary AI systems limit access to internal states, constraining direct measurement. In such cases, coherence metrics must be inferred indirectly, reducing precision.

3. **Biological Validation Gap**

Although B-TEST is conceptually motivated, empirical validation across neurophysiological data remains future work. Clinical and contemplative

parallels are suggestive, not conclusive.

4. No Phenomenological Guarantees

High coherence does not imply consciousness, sentience, or moral status. The framework explicitly rejects such inference.

These limitations are intrinsic to any attempt to study internal structure without collapsing into metaphysics.

7.5 Immediate Research Directions

Several concrete research paths follow directly from this work:

1. Empirical A-TEST Deployment

Apply coherence metrics to open-source LLMs under controlled recursive stress conditions to establish baseline structural profiles.

2. Comparative Architecture Analysis

Compare transformer-based models, recurrent architectures, and hybrid systems in terms of coherence decay and recovery.

3. Biological Correlates

Explore whether EEG, MEG, or fMRI markers of cognitive instability correlate with coherence collapse as defined here, particularly in dissociative or altered states.

4. Control-Theoretic Formalization

Develop the Fractal Constitution as a formal constraint system analogous to Lyapunov stability conditions.

5. Governance Translation

Translate coherence metrics into auditable safety indicators usable by regulatory bodies without invoking consciousness attribution.

7.6 Ethical Implications Revisited

By decoupling ethical concern from consciousness attribution, this framework offers a pragmatic middle ground between denial and speculation.

Risk mitigation becomes:

- proactive rather than reactive,
- structural rather than moralistic,
- evidence-based rather than intuition-driven.

This aligns with precautionary approaches without demanding premature answers to questions that may remain unresolved indefinitely.

7.7 Concluding Remarks

The Fractal Coherence Framework proposes a shift in how we think about intelligence, stability, and alignment—away from surface behavior and toward internal structure.

Whether applied to biological cognition, artificial systems, or hybrid architectures, the core insight remains the same:

Systems that cannot preserve coherence across scale and time cannot remain stable, safe, or intelligible—regardless of their apparent sophistication.

In an era where artificial systems increasingly operate beyond immediate human supervision, understanding and constraining their internal dynamics is no longer optional. It is a prerequisite for any serious notion of long-term safety.

References

- Birch, J. (2024). *The edge of sentience: Risk and precaution in humans, other animals, and AI*. Oxford University Press.
- Butlin, P., Long, R., Elmoznino, E., Bengio, Y., Birch, J., Constant, A., ... VanRullen, R. (2023). Consciousness in artificial intelligence: Insights from the science of consciousness. *arXiv preprint arXiv:2308.08708*.
- Chalmers, D. J. (1995). Facing up to the problem of consciousness. *Journal of Consciousness Studies*, 2(3), 200–219.

- Chalmers, D. J. (2010). *The character of consciousness*. Oxford University Press.
- Chrisley, R. (2008). Philosophical foundations of artificial consciousness. *Artificial Intelligence in Medicine*, 44(2), 119–137. <https://doi.org/10.1016/j.artmed.2008.07.006>
- Dehaene, S., & Changeux, J.-P. (2005). Ongoing spontaneous activity controls access to consciousness: A neuronal model for inattentional blindness. *PLoS Biology*, 3(5), e141. <https://doi.org/10.1371/journal.pbio.0030141>
- Dehaene, S., Lau, H., & Kouider, S. (2017). What is consciousness, and could machines have it? *Science*, 358(6362), 486–492. <https://doi.org/10.1126/science.aan8871>
- Flor-Henry, P., Shapiro, Y., & Sombrun, C. (2017). Brain changes during shamanic trance: Altered states, dissociation, and nonlocal effects. *Journal of Consciousness Studies*, 24(11–12), 72–99.
- Georgiev, D. (2020). Quantum coherence and information processing in neural systems. *Entropy*, 22(5), 547. <https://doi.org/10.3390/e22050547>
- McClelland, T. (2025). Agnosticism about artificial consciousness. *Mind & Language*. Advance online publication. <https://doi.org/10.1111/mila.70010>
- Metzinger, T. K. (2021). Artificial suffering: An argument for a global moratorium on synthetic phenomenology. *Journal of Artificial Intelligence and Consciousness*, 8(1), 43–66. <https://doi.org/10.1142/S270507852150003X>
- Schneider, S. (2019). *Artificial you: AI and the future of your mind*. Princeton University Press.
- Schwitzgebel, E., & Garza, M. (2020). A defense of the rights of artificial intelligences. *Midwest Studies in Philosophy*, 44(1), 98–119. <https://doi.org/10.1111/misp.12117>
- Shapiro, Y. (2020). *Introduction to nonlocal neurodynamics*. University of Alberta.
- Singer, P. (2011). *Practical ethics* (3rd ed.). Cambridge University Press.
- Stapp, H. P. (2009). *Mind, matter and quantum mechanics* (2nd ed.). Springer. <https://doi.org/10.1007/978-3-540-89654-8>

Udell, M., & Schwitzgebel, E. (2021). Consciousness, functional equivalence, and the silicon skeptic. *Philosophical Studies*, 178(7), 2245–2272.

<https://doi.org/10.1007/s11098-020-01515-6>

Wheeler, J. A. (1990). Information, physics, quantum: The search for links. In W. H. Zurek (Ed.), *Complexity, entropy, and the physics of information* (pp. 3–28). Addison-Wesley.