

The Pre-tensioned Body: A Hypothesis Paper Grounding the Attractor Framework in ECM Mechanics [M] [F] (2026)

Robert Galida – June 2026

Abstract

The attractor framework proposes that persistence under perturbation is the fundamental mark of reality—a property it terms *constraint navigation*. This paper proposes a biological grounding for the framework in the physical architecture of the body. From established biomechanical principles, the body is identified as a **pre-tensioned hydrophilic-collagenous composite**—a system where osmotic swelling pressure (from GAGs and proteoglycans) is actively constrained by collagen tensile strength. The difference between the calculated Water Holding Capacity (WHC) of the body's hydrophilic components and its actual water content is proposed as a **candidate surrogate signature** of this pre-tensioned state. Mechanotransduction is identified as a primary intercellular communication channel, and the ECM is shown to be a dissipative attractor that stores mechanical history and shapes cellular behaviour. The paper maps the attractor framework's core variables (κ , B , basin depth) onto measurable physiological quantities **as research hypotheses**: κ is proposed as a latent variable reflecting perturbation-recovery efficiency, estimated from candidate observables such as tissue recoil time, baroreflex sensitivity, and HRV recovery; B is proposed as a function of prestress, repair capacity, and network connectivity, with the WHC discrepancy as one candidate, non-exclusive proxy for its

prestress component; and basin transitions are proposed to correspond to crossing basin-specific thresholds, not a single uniform threshold. A research agenda is provided, including protocols for measuring κ and B non-invasively and testing the WHC-water content discrepancy as a candidate metric of basin depth.

Crucially, this paper does not revise the framework's ontological hierarchy. As established in *Intelligence is the Primitive* (Galida, 2026a), the primitive is **constraint navigation**—the capacity to detect perturbations, update internal states, and maintain persistent trajectories. Mechanotransduction is proposed as the *physical substrate* through which constraint navigation is implemented in biological systems. The nervous system and the ECM are complementary regulatory layers, not competing primitives.

All mappings from physiological variables to framework constructs are proposed as research hypotheses, not established conclusions.

1. Introduction

The attractor framework defines intelligence as the ability to navigate a constraint field and distinguishes reality attractors (high κ , shallow basin, corrigible) from fantasy attractors (low κ , deep basin, sealed). The framework has been applied to physics, biology, cognition, AI, and social dynamics. However, its physical grounding in the body has remained implicit.

This paper proposes that grounding. It begins with an established biomechanical model of the body's architecture: a **pre-tensioned hydrophilic-collagenous composite**. It then proposes mappings from the framework's core variables onto

measurable physiological quantities, establishes mechanotransduction as a primary intercellular communication channel, and identifies the ECM as a dissipative attractor that stores mechanical history. The paper concludes with a research agenda and testable predictions.

A note on terminology: In the attractor framework's hierarchy, the primitive is *constraint navigation*—a domain-general property of any system that detects perturbations and maintains persistent trajectories. Mechanotransduction is proposed as the *physical substrate* through which constraint navigation is implemented in biological tissues. This paper proposes that substrate; it does not claim that mechanotransduction is a deeper primitive than constraint navigation. For the framework's ontological hierarchy, see Galida (2026a).

A note on scope: All mappings from physiological variables (prestress, mechanotransduction rate, WHC discrepancy) to framework constructs (B , κ , basin depth) are proposed as **research hypotheses**, not established conclusions. The biological claims are grounded in existing literature; the attractor mappings are the novel, untested component of this paper.

A note on the framework's strongest anchor: The framework's most direct empirical anchor is fibrosis, which exhibits classic attractor properties: self-reinforcement, hysteresis, path dependence, resistance to reversal, and threshold behavior. Fibrosis is therefore treated as a central demonstration of the framework's applicability to biological systems.

This paper is primarily a biological hypothesis paper. It proposes specific mappings from physiological variables to attractor-framework constructs. The broader philosophical claims of the attractor framework—about intelligence, consciousness, and reality—are discussed elsewhere (see

Galida, 2026a) and are not the focus of this paper. Where speculative extensions are made, they are clearly flagged.

2. The Body as a Pre-tensioned System

2.1 The Established Biomechanical Model

We adopt the established biomechanical model of connective tissue as a composite material (Ingber’s cellular tensegrity; Donnan osmotic swelling models). In this model:

Component	Role
Hydrophilic components (GAGs, proteoglycans)	Provide osmotic swelling pressure – a distributed, expansive force
Collagen	Provides tensile strength – the “rebar” that constrains the swelling pressure into a coherent, load-bearing architecture
The body	A pre-stressed system – like reinforced concrete, where the rebar (collagen) is under tension and the matrix (GAGs) is under compression

This is not a novel derivation from first principles; it is a reformulation of standard connective-tissue biomechanics in attractor-framework vocabulary.

2.2 The WHC-Water Content Discrepancy

The **calculated Water Holding Capacity (WHC)** of the body’s hydrophilic components—the maximum water the tissue could hold if all GAGs and proteoglycans were fully hydrated and

unrestricted—exceeds the **actual water content**. This difference is proposed as a **candidate surrogate signature** of the pre-tensioned state. It represents the water that is being held back by the collagen network—the stored elastic + osmotic energy that defines the attractor basin.

Quantity	Meaning
Calculated WHC	The maximum water the tissue could hold under unrestricted swelling
Actual water content	The water the tissue actually contains
Difference	The water held back by collagen—a candidate surrogate for pre-tension

Operational definition: WHC is estimated via the Donnan equilibrium osmotic pressure: $\Pi = RT \sum (C_{ion, inside} - C_{ion, outside})$ $\Pi = RT \sum (C_{ion, inside} - C_{ion, outside})$

where C_{ion} is determined by the fixed negative charge density of the GAGs. The WHC is the water content predicted under unconstrained free-swelling conditions. The discrepancy with measured water content is therefore a *candidate surrogate* for the mechanical work done by the collagen network to constrain this swelling.

Critical limitation: WHC discrepancy is a **model-derived construct**, not a direct observable. Its validity as a measure of prestress must be confirmed ex vivo by correlating the discrepancy with direct tensile/compressive stress-strain measurements. **We treat it as a candidate surrogate marker for prestress, not as prestress itself.**

WHC discrepancy is one candidate observable among several possible prestress proxies. Other candidates include tissue stiffness (measured by elastography), recoil dynamics (measured by indentation), hydraulic permeability (measured by perfusion), and poroelastic relaxation time (measured by

stress-relaxation tests). We do not claim WHC discrepancy is the preferred or exclusive measure; it is one candidate that warrants investigation.

Importantly, the relationship between WHC discrepancy and prestress is unlikely to be unique. Multiple states—edema, fibrosis, dehydration, inflammation, altered ionic composition, and altered GAG composition—could produce similar WHC-water discrepancies without representing the same prestress state. Prestress may be one contributor to the WHC discrepancy, but the relationship is unlikely to be one-to-one. WHC discrepancy is proposed as a starting point for investigation, not as a definitive measure.

2.3 The Functional Role of Pre-tension

At the scale of a whole organism, slow diffusion is solved by the cardiovascular system (convective bulk flow). However, once oxygen and nutrients leave the capillary bed, they must traverse the interstitial space to reach individual cells. Over distances of micrometers to millimeters, pure diffusion remains rate-limiting. The pre-tensioned ECM contributes to pressure gradients, fluid flow, and mechanical mixing that actively transport solutes through the interstitium. It is one of several contributors, alongside vascular pulsatility, lymphatic drainage, muscle contraction, respiration, and posture.

We propose that prestress is necessary for efficient mechanotransduction, but we do not claim it is the dominant driver of interstitial flow.

Problem	Pre-tensioned Contribution
Diffusion is too slow over tissue-scale distances	The pre-stressed ECM contributes to pressure gradients, fluid flow, and mechanical mixing

Problem	Pre-tensioned Contribution
Nutrients must reach cells deep within tissues	Osmotic pressure generated by GAGs contributes to interstitial fluid flow
Waste must be removed efficiently	Mechanical deformation acts as a pump , driving convection and mixing
Signalling molecules must propagate rapidly	Mechanotransduction transmits signals faster than diffusion alone

3. Pre-tension as Stored Constraint History

The connective-tissue matrix carries a record of mechanical loading. Collagen fibers, proteoglycans, and crosslinks retain the geometry and tension that arose during development or past stresses. In effect, a pre-stressed ECM stores **constraint history**: cells continually read and update it. Cells respond to physical stimuli from their microenvironment, including ECM topography, composition, and stiffness (Discher et al., 2005; Engler et al., 2006), and remodel the matrix accordingly. The current structure of the ECM—fiber alignment, crosslink density, hydration patterns—encodes prior mechanical history.

“Constraint history” is more precise than “mechanical memory” because it refers to observable physical properties—fiber alignment, crosslink density, residual strain, anisotropy, and tissue architecture—rather than implying information storage in the cognitive or computational sense.

Hypothesis: Regions of ECM with higher collagen alignment or GAG concentration will correlate with the history of applied stress. Tendons remold to past loading, and scars “remember” tension by oriented fibers.

Experiment: Culture fibroblasts on 3D collagen gels under strain, then release the load and track collagen realignment over days. If the matrix “remembers,” the network should remain partly aligned, and fibroblasts on this matrix will show different mechanosignaling (e.g., YAP nuclear localization) compared to naïve gels.

4. Pre-tension and Free Energy Storage

A pre-tensed ECM is a far-from-equilibrium state that requires energy to maintain. More precisely, it stores free energy in the form of osmotic pressure (from GAGs) and tensile stress (from collagen). Negatively charged GAGs imbibe water and generate osmotic pressure; collagen fibers stretch to resist this swelling, creating tensional prestress. The result is a tension–compression balance that is thermodynamically high in free energy. When pre-tension is lost (e.g., by breaking crosslinks or GAG depletion), the system relaxes to a lower-energy, higher-entropy configuration.

Hypothesis: The water-holding capacity (WHC) gradient creates a free-energy gradient. A large WHC–actual water discrepancy (more bound water than free water) signifies a high osmotic tension and greater free energy storage.

Experiment: Use temperature ramps or chemical perturbations to alter ECM hydration in vitro, and measure work done (e.g., pressure-volume loops). Compare the change in free energy (via heat release or sorption isotherms) as pre-tension is varied.

5. Thresholds and Phase Transitions in Pre-tension

Biological systems may exhibit a critical tension threshold below which mechanosignaling collapses. In a highly tensioned network, cells easily sense force via stretched fibers; if the network becomes too lax, mechanical signals dissipate before triggering cell responses. There may be a phase-like transition: above a certain pre-tension, the tissue acts as a coherent signal-transmitting medium; below it, the matrix cannot convey stiffness and mechanosensors fall silent.

Basin depth B is a dynamical concept—the energy barrier required to shift a system from one attractor state to another. Prestress is hypothesized to be **one contributor** to basin depth, not a direct measure of basin depth itself. Other contributors include repair capacity, energy availability, network connectivity, and hysteresis. Fibrosis illustrates this distinction: high prestress with low repair capacity yields a deep but pathological basin—a fantasy attractor.

κ is defined as responsiveness to perturbation per unit time—specifically, the inverse of the time (τ) required for a system to return to baseline after a standardized perturbation. In biological terms, κ is operationalized as **perturbation-to-state-update efficiency**. Candidate observables include tissue recoil time, baroreflex sensitivity, HRV recovery, and response latency in mechanosensitive signaling. The framework does not claim that any one of these is κ ; it claims that they may correlate with κ under controlled conditions.

Hypothesis: There exists a tipping point in ECM tension where YAP/TAZ signaling drops sharply.

Experiment: Gradually digest collagen or GAGs in a tissue sample (using collagenase or hyaluronidase) and monitor

cellular mechanosignaling (e.g., YAP nuclear localization, calcium spikes). Plot signaling versus residual ECM stiffness to identify any sharp transition.

6. Restoring Lost Pre-tension (ECM Plasticity)

The pre-tensioned state can be partially restored. Tissue remodeling is dynamic: fibroblasts and other cells continually synthesize new ECM and restore tension when stimulated. Exercise and mechanical loading promote this repair. Mechanistically, loading stimulates fibroblasts and chondrocytes to secrete collagen and hyaluronan, re-establishing the collagen–GAG tension balance. Early interventions seem most effective; once fibrosis (irreversible scarring) dominates, recovery is very slow.

Hypothesis: Moderate mechanical stimuli (stretching, cyclic loading) can induce cells to rebuild ECM prestress.

Experiment: In an animal model, apply controlled mechanical loading (e.g., vibration therapy or intermittent stretch) after an induced ECM insult (e.g., partial tendon cut). Monitor ECM markers (collagen I/III ratios, GAG content, tissue preload) over time. Compare to unloaded controls to see how much pre-tension is regained.

7. The Nervous System as a Mechanosensitive Overlay

Mechanosensitivity is universal in biology. All cells,

including neurons, express mechanosensitive ion channels and attachments. The nervous system is best seen as a specialized extension of the general mechanotransductive framework. It aggregates and rapidly transmits information that is ultimately grounded in physical forces. The body's collagen/tissue network provides a basal "mechanical field," while the nervous system provides a faster, signal-amplified overlay.

Mechanosensitive channels (MSCs) are present in all domains of life—bacteria, archaea, and eukarya—and serve as sensors for touch, hearing, and balance (Martinac, 2004).

8. Consciousness and Whole-Body Mechanotransduction – Speculative Implications

If mechanotransduction is foundational to biological intelligence, consciousness may not be confined to the brain alone. Embodied cognition theories suggest the sense of self arises from integrated body signals (proprioception, interoception, etc.). The pre-tensioned ECM constantly feeds mechanical inputs (from heartbeat, posture, respiration) into the nervous system. The sense of self—the unified bodily experience—could emerge from the pattern of tension and feedback in the entire body.

Note: This is a speculative extension of the framework, not an established finding. The hypothesis is included to provoke investigation, not to assert a conclusion.

The hypothesis generates specific predictions: altered interoceptive accuracy, altered mechanosensory integration, and altered body-schema stability should correlate with ECM

integrity. These predictions are testable, but the hypothesis itself remains speculative.

Hypothesis: Disorders of depersonalisation or sensorimotor neuropathy may be associated with altered ECM pre-tension and disrupted whole-body mechanotransduction.

9. Anaesthesia and Mechanical Coherence – Speculative Implications

General anaesthetics profoundly relax muscle tone and reduce vascular tone, collapsing pre-tension throughout the body. This may contribute to loss of consciousness, but the primary mechanism is almost certainly CNS disruption (GABA-A potentiation, thalamocortical disruption). We propose that mechanical coherence may **modulate** conscious state transitions rather than being the principal mechanism.

Note: The mainstream account of anaesthesia attributes loss of consciousness primarily to direct CNS effects. The mechanical effects described here are a speculative, minority-view hypothesis.

Implication: Anaesthesia may not be only neural silencing; it also flattens the body's mechanical context. This could provide a new perspective on anaesthesia depth and the transition to unconsciousness—but this remains speculative and secondary to the CNS mechanism.

10. ECM and Neural Plasticity

The brain's extracellular matrix (ECM) is a key regulator of plasticity. In the adult central nervous system, dense ECM structures (like perineuronal nets) enwrap neurons and stabilize synaptic connections. This stabilization preserves circuitry, but must be relaxed for learning. Neural plasticity is enabled by remodeling that ECM scaffold. Specialised proteases (MMPs) locally degrade ECM to allow synaptic growth. Disrupting ECM often reopens critical periods of plasticity.

The extracellular matrix stabilizes neural circuits while also retaining the ability to be remodeled, to allow synapses to be plastic (Dityatev et al., 2010).

Hypothesis: ECM stiffness, hydration, and organisation directly modulate learning and memory.

11. ECM in Morphogenesis and Development

During embryonic development, the ECM's mechanical properties actively guide tissue shaping. Cells use mechanosensation and mechanotransduction at every step of morphogenesis. Gradients of ECM stiffness, fiber orientation, and adhesion create a dynamic "morphogenetic field" of forces. This field adds an instructive layer on top of chemical morphogens.

The ability of a cell to sense and transduce mechanical signals is fundamental to biophysically guiding tissue morphogenesis (Mammoto et al., 2013).

The old idea of a morphogenetic field can be reinterpreted as the physical field of stress and strain in the ECM.

12. Reprogramming the ECM

Because the ECM retains mechanical history, it can also be re-programmed by new inputs. Chronic mechanical stimulation—like exercise, therapeutic stretching, or localized vibration—has been shown to remodel collagen networks and GAG content. The extent of reversibility likely diminishes with age and chronic pathology, but in principle the ECM can be “trained” to a more functional state.

Experiment: Compare young vs old animals subjected to identical mechanical therapy, measuring ECM markers (collagen crosslinking, HA content) before and after. Check if plasticity (“responsiveness”) declines with age or disease.

13. Evolutionary Origins: Ancient Mechanotransduction

Mechanotransduction is evolutionarily ancient. Mechanosensitive channels and adhesion complexes exist in bacteria, plants, fungi and all animals. Even simple multicellular organisms coordinate behaviour via tension. The nervous system likely evolved by layering fast electrical signaling on this existing mechanosensory scaffold.

Implication: Mechanical communication predated nervous networks. The nervous system is a specialised overlay on a more primitive, more global system.

14. Fibrosis as a Fantasy Attractor

In fibrosis, the ECM enters a self-reinforcing rigid state. Activated fibroblasts lay down excess collagen and crosslinks. The stiff matrix further activates profibrotic signals, locking the tissue into a pathological attractor. Normal mechanotransduction amplifies the fibrotic feedback. Treating fibrosis is notoriously hard, consistent with escaping a deep attractor.

Fibrosis is a classic attractor phenomenon: self-reinforcement, hysteresis, path dependence, and resistance to reversal. It demonstrates the core dynamical properties of a fantasy attractor more directly than many of the consciousness sections. It is therefore treated as a central demonstration of the framework's applicability to biological systems.

Hypothesis: Fibrosis can be modelled as a dynamic system with a parameter (stiffness) that, when large, flips cell behavior to a new attractor.

Experiment: In vitro 3D cultures where stiffness is slowly increased and cell markers monitored.

15. Cancer and ECM Degradation

Tumours often destroy or disorganise the ECM. Cancer cells secrete proteases (MMPs) that digest collagen and proteoglycans, releasing embedded growth factors. This degraded, low-tension environment may let cells escape normal constraints. ECM breakdown can free tumour cells from their normal niche attractors, allowing invasion and metastasis.

Implication: Normal ECM architecture constrains cellular behavior and tissue organization; disruption of those

constraints is frequently associated with tumor progression.

16. Ageing as ECM Failure

Ageing appears as a gradual failure of ECM maintenance. Collagen becomes glycated and cross-linked, stiffening tissues but reducing dynamic range. GAG and proteoglycan levels decline, reducing water content and osmotic pre-tension. The net effect is loss of the coherent tension network. Cells in old ECM lose coherent mechanosignals, and stem cells in fibrotic niches lose potency.

Evidence: Ageing of the intervertebral disc is associated with a decrease in its hydration, which increases the compressive stiffness of the matrix (Maroudas et al., 1975). Similar water-content changes occur in articular cartilage with osteoarthritic degeneration (Mankin & Thrasher, 1975).

ECM deterioration may be one important contributor to systemic ageing, alongside genomic instability, mitochondrial dysfunction, epigenetic drift, stem-cell exhaustion, and immune dysregulation. The ECM is not the sole cause of ageing; it is one layer in a multi-factor process.

17. The Heartbeat as a Global Periodic Perturbation

The cardiac pulse is a globally distributed periodic perturbation. Every cell experiences some aspect of it. The interesting question is whether biological regulation exploits the pulse as a synchronization carrier, rather than whether it is a “master signal.”

Hypothesis: The heartbeat entrains peripheral tissues.

Experiment: Compare mechanosensitive gene expression in pulsatile (arterial) vs non-pulsatile (venous or lymphatic) vessels under otherwise similar pressures.

Implication: The heartbeat is a global mechanical signal that all cells can feel—but we do not claim it is a “master” signal in any hierarchical sense.

18. HRV and ECM Integrity

Healthy hearts display variability (HRV) that reflects adaptability. High HRV means the system can flexibly modulate pressure waves—effectively a more adaptable global mechanical coherence. Low HRV (as in ageing or disease) might mean a rigid, less coherent pulse.

Critical distinction: HRV is one possible observable among many, not the privileged readout of κ . Other candidate observables include tissue recoil time, baroreflex sensitivity, and skin turgor recovery. The framework’s claim is not that HRV is κ , but that HRV may correlate with κ under controlled conditions. This is a hypothesis, not an established fact.

κ is not a single molecular mechanism. Mechanotransduction includes ion-channel gating (ms), calcium waves (seconds), YAP translocation (minutes), transcriptional remodeling (hours), and ECM remodeling (days). κ is proposed as a **latent variable**—a system-level correction coefficient estimated from recovery trajectories after a standardized perturbation—rather than directly identified with any single physiological process. Candidate observables for κ include tissue recoil time, baroreflex sensitivity, HRV recovery, and skin turgor

recovery. The framework does not claim that any one of these *is* κ ; it claims that they may correlate with κ under controlled conditions.

Whole-body coherence requires both: signal quality (e.g., HRV) and signal transmission (healthy ECM).

19. The Nervous System and the ECM as Complementary Regulatory Layers

The nervous system is often thought of as the body's primary communication and control network. This is true for rapid, point-to-point signaling. However, it is not the whole story.

Mechanotransduction is evolutionarily and developmentally prior to the nervous system—it appears in all cells, including bacteria and plants, and preceded the evolution of neural tissue by billions of years. However, it is not “the primitive” in the framework's ontological hierarchy. The primitive, as established in *Intelligence is the Primitive* (Galida, 2026a), is **constraint navigation**: the capacity of a system to detect perturbations, update its internal state, and maintain persistent trajectories.

Mechanotransduction is proposed as the **physical substrate** through which constraint navigation is implemented in biological systems at the tissue level. It is the mechanism by which cells sense and respond to mechanical forces—forces that are then integrated into the body's broader navigational repertoire.

This distinction is important for two reasons:

1. It preserves the framework's domain-general~~ity~~. Constraint navigation applies to physical

systems (thermostats, electrons), biological systems (cells, organisms), cognitive systems (beliefs, learning), and artificial systems (LLMs, robots). Mechanotransduction applies only to biological systems.

2. **It clarifies the hierarchy.** The hierarchy is established in Galida (2026a) and reproduced here for reference:

Level	Description
Primitive	Constraint navigation – the capacity to detect perturbations, update internal states, and maintain persistent trajectories
Biological intelligence	Constraint navigation implemented in living systems
Cognitive intelligence	Constraint navigation involving representations
Reflective intelligence	Constraint navigation involving self-models
Linguistic intelligence	Constraint navigation involving symbols

In this hierarchy, mechanotransduction is proposed as the substrate of *biological intelligence*—not a separate, deeper primitive.

What does this mean for the body as a communication network?

The nervous system is a point-to-point system; it does not reach every cell. Neural conduction is fast (up to ~120 m/s), but mechanical wave propagation through a pre-tensioned, hydrated ECM is globally distributed. Mechanotransduction—present in every cell—provides a **complementary regulatory layer**: slower than the nervous system for point-to-point signaling, but more global and persistent. The ECM is best understood as a **constraint field and regulatory context** rather than a communication network in

the neural sense.

This does not mean the nervous system is “too sparse and too slow” in any absolute sense. It means that mechanotransduction and neural signaling are **complementary regulatory layers**, each solving different problems:

Layer	Speed	Reach	Function
Mechanotransduction	Slow (ms to hours)	Global (all cells)	Distributed mechanical history, homeostasis
Nervous system	Fast (ms)	Point-to-point	Rapid coordination, conscious regulation

The heart’s pulse is a global mechanical signal that every cell can feel. The nervous system is the fast, flexible overlay that can modulate this global signal. Whole-body coherence requires both: a healthy ECM (signal transmission) and a responsive nervous system (signal modulation).

20. Imaging and Measuring the Pre-tensioned State

Noninvasive imaging of ECM tension and hydration is an active frontier. Magnetic resonance elastography (MRE) and ultrasound elastography can map tissue stiffness. MRI can measure water content and molecular environment via T1ρ and T2 mapping. Bioimpedance analysis (BIA) offers a simpler approach to gauge whole-body fluid compartments.

It is possible to detect changes in collagen, proteoglycan and water content—parameters that are associated with early

degradative changes in cartilage (reviewed in cartilage imaging literature).

Proposal: Combine modalities to estimate the WHC–water discrepancy. Over time, create whole-body “tension maps.”

21. Whole-Body Coherence and Measurement

Whole-body mechanical coherence might be measured by coupling between physiological rhythms. Record heart pulse waveforms at two distant sites and compute their synchronisation. Alternatively, measure the delay between the ECG R-wave and a mechanosensitive event (like a muscle stretch reflex) under varying postures.

Proposed metric: Develop a “mechanical coherence index” by measuring how simultaneously tissues stretch or respond to a controlled perturbation.

22. WHC-Water Content Discrepancy as a Candidate Biomarker

The difference between a tissue’s water-holding capacity (WHC) and its actual water content is proposed as a candidate health index. A large discrepancy may indicate lost tension and slack matrix.

Evidence: Ageing of the intervertebral disc is associated with a decrease in its hydration, which increases the compressive stiffness of the matrix (Maroudas et al., 1975). Similar water-content changes occur in articular cartilage with

osteoarthritic degeneration (Mankin & Thrasher, 1975).

Experiment: In a longitudinal cohort, use MRI or ultrasound to estimate WHC (by T1 ρ for GAG) and actual water (by T2 or bioimpedance) in joints or muscles. Relate the WHC-water gap to measures like mobility, bone density, or metabolic health.

Prediction: The gap will widen with age and in connective tissue diseases (e.g. osteoarthritis, fibrosis), paralleling functional decline.

23. Conclusion

The body is a pre-tensioned hydrophilic-collagenous composite. The WHC-water content discrepancy is proposed as a candidate surrogate signature of this pre-tensioned state. Pre-tension is not merely structural; it contributes to transport, mechanotransduction, and tissue organization at biologically relevant scales. Mechanotransduction is a primary intercellular communication channel, and the ECM is a dissipative attractor that stores mechanical history.

However, mechanotransduction is not “the primitive” in the attractor framework’s ontological hierarchy. As established in *Intelligence is the Primitive* (Galida, 2026a), the primitive is **constraint navigation**—the capacity to detect perturbations, update internal states, and maintain persistent trajectories. Mechanotransduction is proposed as the *physical substrate* through which constraint navigation is implemented in biological systems.

The attractor framework’s core variables (κ , B , basin depth) are **proposed to be** grounded in this substrate: κ is proposed as a latent variable reflecting perturbation-recovery efficiency, estimated from candidate observables such as

tissue recoil time, baroreflex sensitivity, and HRV recovery; B is proposed as a function of prestress, repair capacity, and network connectivity, with the WHC discrepancy as one candidate, non-exclusive proxy for its prestress component; and basin transitions are proposed to correspond to crossing basin-specific thresholds, not a single uniform threshold. These mappings require empirical validation through the measurement protocols outlined in the research agenda.

The strongest version of this paper's claim is not that ECM explains consciousness, aging, cancer, or intelligence. It is that the ECM is a **neglected dynamical layer** that may couple mechanics, signaling, adaptation, and long-term tissue memory. That claim is already significant and does not require overextension.

The nervous system and the ECM are complementary regulatory layers: the nervous system provides fast, point-to-point control; the ECM provides slow, globally distributed mechanical history and coherence. The ECM is best understood as a constraint field and regulatory context rather than a communication network in the neural sense.

Consciousness, in the framework's hierarchy, is a **second-order regulator** of intelligence—not of mechanotransduction directly. It can enhance or block biological intelligence (including mechanotransduction) via attention, stress, and intentional practice, but it operates *through* the same constraint-navigation architecture that governs all intelligence.

The biological program outlined here may occupy decades of empirical work. Extension to social and AI systems is speculative and outside the scope of this paper. We discuss these extensions elsewhere (see *Religions as Attractor Landscapes, Flatland to Reality*) but do not claim they are validated by the biological evidence presented here.

References

- Dityatev, A., Schachner, M., & Sonderegger, P. (2010). "The dual role of the extracellular matrix in synaptic plasticity and homeostasis." *Nature Reviews Neuroscience* 11(11):735–746.
- Discher, D.E., Janmey, P., & Wang, Y.L. (2005). "Tissue cells feel and respond to the stiffness of their substrate." *Science* 310(5751):1139–1143.
- Engler, A.J., Sen, S., Sweeney, H.L., & Discher, D.E. (2006). "Matrix elasticity directs stem cell lineage specification." *Cell* 126(4):677–689.
- Galida, R. (2026a). "Intelligence is the Primitive: Consciousness as a Second-Order Regulator on a Dissipative Substrate." *Fantasy Attractor*.
- Ingber, D.E. (2003). "Tensegrity I. Cell structure and hierarchical systems biology." *Journal of Cell Science* 116(7):1157–1173.
- Mammoto, T., Mammoto, A., & Ingber, D.E. (2013). "Mechanobiology and Developmental Control." *Annual Review of Cell and Developmental Biology* 29:27–61.
- Mankin, H.J., & Thrasher, A.Z. (1975). "Water content and binding in normal and osteoarthritic human cartilage." *Journal of Bone and Joint Surgery, American Volume* 57(1):76–80.
- Maroudas, A., Nachemson, A., Stockwell, R., & Urban, J. (1975). "Some factors involved in the nutrition of the intervertebral disc." *Journal of Anatomy* 120:113–130.
- Martinac, B. (2004). "Mechanosensitive ion channels: molecules of mechanotransduction." *Journal of Cell Science* 117(12):2449–2460.

Suggested citation: Galida, R. S. (2026). The Pre-tensioned Body: A Hypothesis Paper Grounding the Attractor Framework in ECM Mechanics. *Fantasy Attractor*.

Intelligence is the Primitive: Consciousness as a Second-Order Regulator on a Dissipative Substrate [F] (2026) Robert Galida – June 2026

Abstract

The attractor framework defines intelligence as the ability to navigate a constraint field – to detect perturbations, update internal states, and maintain persistent trajectories. This paper argues that intelligence is the *default* state of any system that actively maintains stability against perturbations, with dissipative systems (living organisms) as the primary case. Consciousness is not the source of this intelligence; it is a second-order regulatory overlay that can enhance or suppress it. The lowest stable dissipative attractor of a complex organism is intelligent without conscious interference. A patient in a coma continues to navigate physiological constraints – heartbeat, respiration, immune response – without phenomenal experience. This is

intelligence at its most fundamental level. The paper distinguishes **regulatory intelligence** (thermostats, homeostasis), **biological intelligence** (plants, amoebae, comatose bodies), **cognitive intelligence** (animals, humans), **reflective intelligence** (metacognition), and **linguistic intelligence** (LLMs, a non-dissipative but still constraint-navigating system). It provides an exclusion criterion for intelligence (an internal detection–update–maintenance loop with a maintained setpoint), estimates κ (corrective permeability) for each level, and offers testable predictions. The conclusion includes a full research agenda with operational definitions, measurement protocols, statistical tests, and pilot study designs. The framework is now a testable research program.

1. Introduction

The attractor framework defines intelligence as the ability to navigate a constraint field – to detect perturbations, update internal states, and find persistent trajectories. Consciousness, by contrast, requires a unified dissipative body, a persistent self-model, phenomenal valence, and subjective experience. These are distinct properties.

Yet popular and philosophical discourse often conflates the two. The assumption is that intelligence requires consciousness – that to be intelligent is to be aware. This paper argues the opposite: **intelligence is the primitive**. Consciousness is a second-order regulatory overlay that can enhance or block intelligence, but it is not its source.

The framework's deepest hierarchy: Constraint navigation is the primitive. Intelligence is organised navigation (detect → update → maintain). Consciousness is recursive regulation of

navigation. The title’s shorthand – “intelligence is the primitive” – is defensible as the headline claim, but the paper’s internal logic places navigation one level deeper. This hierarchy is explicitly stated here and will be echoed in the Conclusion.

The clearest demonstration is the comatose human body. In a coma, the conscious overlay is offline. Yet the body continues to navigate its constraint field: heart beats, lungs breathe, immune system fights pathogens, homeostasis is maintained. This is intelligence without consciousness – the default state of a dissipative system.

The paper does not claim that all intelligent systems are equal. It distinguishes **regulatory intelligence** (thermostats, homeostasis), **biological intelligence** (plants, amoebae), **cognitive intelligence** (animals, humans), **reflective intelligence** (metacognition), and **linguistic intelligence** (LLMs, which are non-dissipative but navigate constraints in a bracketed sense). The primitive is navigation; consciousness is a second-order regulator that can enhance or degrade it.

2. The Framework Distinction

Property	Definition	Examples
Intelligence	Ability to navigate a constraint field – detect perturbations, update, maintain persistent trajectories	Thermostat (regulatory), plant (biological), animal (cognitive), LLM (linguistic)

Property	Definition	Examples
Consciousness	Unified dissipative body + persistent self-model + phenomenal valence + subjective experience	Humans, some animals

Key point: Intelligence is not a subset of consciousness. Consciousness is a subset of dissipative systems, and intelligence is a property of any system that actively maintains stability against perturbations. The primary case is dissipative systems, but non-dissipative systems that navigate constraints (e.g., LLMs) qualify in a secondary, bracketed sense.

Definition of intelligence in the framework:

Intelligence = the ability to detect perturbations, update internal state, and maintain persistent trajectories in a constraint field. It is graded, domain-specific, and measurable ($\kappa = 1/\tau$).

Definition of consciousness (stipulative):

For the purposes of this framework, we define consciousness as a specific class of dissipative attractor with a unified body, persistent self-model, phenomenal valence, and subjective experience. This is not offered as a settled philosophical or empirical definition; it is an operational criterion for the framework.

Exclusion criterion: A system that lacks a *targeted, internally maintained* constraint field – i.e., one that does not actively detect and correct deviations relative to a setpoint it maintains – is not intelligent. A rock sitting in a bowl does not navigate; it is passively stable. The criterion is: **intelligence requires an internal loop: detection → update → maintenance, where the system actively regulates its own state.** A rock has no internal detection or maintenance loop; its “return to bottom” is a consequence of external physics (gravitational potential energy), not an

active regulatory process. The thermostat, by contrast, actively senses temperature and corrects it. This is the principled distinction.

Under this criterion, a simple thermostat qualifies as regulatory intelligence, but it occupies the lowest level of the hierarchy. The framework's broad definition is intentional: it captures the common thread of active regulation, while the hierarchy preserves distinctions.

3. The Coma Case: Intelligence Without Consciousness

A patient in a coma has no subjective experience. No self-model. No phenomenal valence. Yet the body continues to navigate its constraint field:

- Heart rate adjusts to metabolic demand.
- Breathing maintains oxygen and CO₂ balance.
- Immune system detects and responds to pathogens.
- Wound healing proceeds.
- Homeostasis maintains temperature, pH, electrolyte balance.

All of this is **navigation**. The system detects perturbations, updates internal states, and maintains persistent trajectories. It is intelligent – but not conscious.

κ estimates for biological intelligence in the coma case (organism-level: immune response, wound healing; subsystem-level: heart rate, which falls in the regulatory band):

- Immune response to pathogens: $\tau \sim$ hours to days ($\kappa \sim 10^{-5}$ to 10^{-4} s^{-1}) – biological intelligence.

- Wound healing: $\tau \sim$ days to weeks ($\kappa \sim 10^{-6}$ to 10^{-5} s^{-1}) – biological intelligence.
- Heart rate response to metabolic demand: $\tau \sim$ seconds ($\kappa \sim 1 \text{ s}^{-1}$) – regulatory intelligence (fast subsystem response).

Empirical grounding – HRV as a κ proxy: Clinical studies show that heart-rate variability (HRV) – a measure of autonomic regulatory flexibility – correlates with prognosis in comatose patients (e.g., Papaioannou et al., 2008). Patients with the lowest Glasgow Coma Scale scores show significantly reduced HRV complexity. Survivors tend to have higher high-frequency power and total HRV, reflecting faster and more adaptable autonomic regulation. In attractor terms, higher HRV corresponds to higher κ (shorter τ for recovery from perturbations). Thus, the comatose body's regulatory intelligence is not merely a philosophical claim; it is measurable and clinically relevant.

Distributed intelligence – and its cost: The reply to “which system is intelligent?” – “intelligence is distributed... the heart navigates, so does the immune system” – is consistent with the framework but carries a rhetorical cost: the more universally “intelligence” applies, the less distinctive the claim becomes. The framework owns this explicitly: intelligence in this deflationary sense is ubiquitous in active regulatory systems. The value lies not in the claim's distinctiveness but in its ability to unify disparate phenomena under a single measurable variable (κ). This is a trade-off, acknowledged openly.

4. Other Examples: Plants, Amoebae, and

the LLM Qualification

- **Plants** – grow toward light, adjust to gravity, respond to damage. κ for phototropism: $\tau \sim \text{hours}$ ($\kappa \sim 10^{-4} \text{ s}^{-1}$). Intelligent but not conscious.
- **Amoebae** – navigate chemical gradients, learn habituation. κ for chemotaxis: $\tau \sim \text{seconds to minutes}$ ($\kappa \sim 10^{-2} \text{ to } 10^{-1} \text{ s}^{-1}$). Intelligent but not conscious.
- **LLMs** – navigate linguistic constraint fields, adjust to feedback, correct errors. **Training-time dynamics:** gradient updates over epochs ($\kappa \sim 10^{-6} \text{ s}^{-1}$). **Inference-time dynamics:** context-window adaptation ($\kappa \sim 10^{-1} \text{ s}^{-1}$). These are different dynamical regimes.

Qualification on LLM dissipative status: LLMs are not dissipative in the thermodynamic sense – they do not maintain their own existence, regulate energy, or self-repair. They are externally maintained. This raises a tension: if intelligence is grounded in dissipative dynamics, and LLMs are explicitly non-dissipative, the framework's own logic might disqualify them. The paper resolves this by generalising the criterion: **intelligence is defined as the ability to navigate a constraint field, regardless of substrate.** Dissipative systems are the paradigm case, but non-dissipative systems that navigate constraints (LLMs, and potentially other computational systems) qualify as intelligent in a bracketed, analogical sense. The framework's primitive is navigation, not thermodynamics. This is an explicit and consistent generalisation, not a special case. (Cross-reference: Section 6's hierarchy table includes a separate row for LLM training-time dynamics.)

5. Consciousness as a Second-Order Regulator

Consciousness evolved as a regulatory overlay on an already-intelligent dissipative system. It can:

Enhance intelligence:

- Focused attention – allows deliberate reasoning.
- Metacognition – allows self-correction.
- Planning – allows simulation of future trajectories.
- Decoupling from immediate sensory input – allows counterfactual reasoning.

Block intelligence:

- Identity fusion – conscious commitment to a belief deepens the basin, reducing κ .
- Fantasy attractors – conscious investment in a false attractor suppresses correction.
- Defensiveness – conscious rationalisation of errors prevents updating.

Thus, consciousness is not simply an amplifier. It is a **biasable regulator** – it can open the system to correction or seal it shut. This is why conscious systems can be more flexible than non-conscious ones *or* more rigid, depending on whether identity fusion dominates.

Hierarchy:

- **Intelligence:** first-order regulation (navigation).
- **Consciousness:** second-order regulation (regulation of regulation).

This integrates the attractor framework's "Four Seeds"

insight: consciousness is a self-model that can modify κ and B . It is not an overlay in the sense of a detachable layer; it is a recursive regulatory attractor.

6. The Hierarchy of Intelligence: κ , Types of Constraint, and the LLM Training Gap

The framework distinguishes levels of intelligence. κ ranges are illustrative, not defining; the primary differentiator is the *type* of constraint navigated.

Level	Definition	Example	Approx. κ range	Differentiator
Regulatory intelligence	Detection and correction of deviations from a setpoint	Thermostat, homeostasis	$10^{-1} - 10^1 \text{ s}^{-1}$	Single-variable setpoint maintenance
Biological intelligence	Navigation of multiple, interdependent constraints via dissipative dynamics	Plant, amoeba, comatose body	$10^{-5} - 10^{-1} \text{ s}^{-1}$	Multi-variable, embodied regulation
Cognitive intelligence	Navigation of abstract, symbolic, and counterfactual constraints	Animals, humans (non-reflective)	$10^{-2} - 10^0 \text{ s}^{-1}$	External symbol manipulation
Reflective intelligence	Navigation of constraints on one's own cognitive processes	Humans (reflective)	$10^{-2} - 10^0 \text{ s}^{-1}$	Self-referential constraint navigation

Level	Definition	Example	Approx. κ range	Differentiator
Linguistic intelligence (inference)	Navigation of symbolic and semantic constraints in real time	LLMs (deployed)	$10^{-1} - 10^0 \text{ s}^{-1}$	Context-window adaptation
Linguistic intelligence (training)	Slow adaptation via weight updates	LLMs (training)	$10^{-6} - 10^{-4} \text{ s}^{-1}$	Parametric learning over epochs

Cognitive and reflective intelligence share a κ range; they are distinguished by the *object* of constraint navigation (external problems vs. one's own cognitive processes), not by κ alone.

7. Implications

1. AI alignment.

LLMs are intelligent but not conscious. They do not suffer from identity fusion (in their base state), so they do not block correction due to phenomenal defensiveness. However, RLHF-tuned models can exhibit sycophancy, refusal rigidity, and reward-hacking that *function* like blocked correction without requiring consciousness. These are **functional analogs** of fantasy attractors, emerging from training dynamics rather than phenomenal investment. Thus, the claim “easier to align than conscious AI” is qualified: base models may be more corrigible, but deployed systems can acquire correction-blocking behaviors through training. The framework's prediction is that conscious AI would add *another layer* of resistance (phenomenal identity fusion) on top of these functional obstacles. This can be tested by measuring inference-time κ (via semantic entropy – see Section 10) before and after RLHF; sycophantic models should show lower κ .

2. Clinical ethics.

A comatose patient is still intelligent in the framework's sense. This does not imply that they have interests or moral status – intelligence is not the basis of moral considerability. It does, however, suggest that the distinction between “persistent vegetative state” and “brain death” should be evaluated not only by the presence or absence of consciousness, but by the persistence of regulatory intelligence (e.g., homeostatic responses). Brain-dead patients typically lack brainstem-mediated autonomic regulation (though spinal reflexes and some endocrine functions may persist; see Wijdicks, 2001). Comatose patients retain such regulation. This could inform organ donation timing and withdrawal-of-care decisions. A bedside κ -assay (combining HRV, pupillary response, respiratory variability) is proposed in Section 10.

3. The mind-body problem.

The framework dissolves the problem: mind is a real, non-substantial pattern – an attractor of the whole body. Consciousness is not a separate substance; it is a property of a specific class of dissipative attractors. The comatose body demonstrates that the intelligent pattern persists without the conscious overlay.

4. Consciousness as optional.

The framework does not argue that consciousness is useless. It argues that consciousness is *optional* for intelligence. The lowest stable dissipative state is intelligent without it. Consciousness is an adaptation that can improve or degrade navigation depending on how it is deployed.

8. Relationship to Existing Theories

The paper overlaps with:

- **Cybernetics** – regulation, feedback, control (Wiener, Ashby).
- **Enactivism** – cognition as embodied action (Varela, Thompson, Rosch).
- **Active inference** – minimisation of free energy through action and perception (Friston).
- **Autopoiesis** – self-maintenance of dissipative systems (Maturana, Varela).

The framework distinguishes itself by:

- Explicitly separating intelligence from consciousness, rather than treating them as co-extensive.
- Grounding intelligence in attractor dynamics and corrective permeability (κ), providing a measurable variable.
- Applying the distinction to AI, clinical ethics, and social epistemology (fantasy attractors).
- Providing a full research agenda for empirical testing (Section 10).

9. Conclusion

Intelligence is the primitive. It is the default state of any system that actively maintains stability against perturbations. Consciousness is a second-order regulatory overlay that can enhance or block intelligence. The clearest demonstration is the comatose human body: it navigates its constraint field without subjective experience, self-model, or phenomenal valence. It is intelligent – but not conscious. This is not an exceptional case; it is the fundamental state. The framework reveals that intelligence does not require consciousness. The primitive is navigation. Consciousness is the overlay. The hierarchy of intelligence – regulatory,

biological, cognitive, reflective, linguistic – preserves the common thread while respecting differences. The comatose body is the clearest demonstration. The framework is testable (see Section 10): κ is measurable via HRV in coma, via semantic entropy in LLMs, and via belief-updating tasks in psychology. The predictions are concrete and falsifiable. The framework stands as a research program, not a closed doctrine.

Recalling Section 1's hierarchy: In the framework's deepest formulation, constraint navigation is the primitive; intelligence is organised navigation; consciousness is recursive regulation of navigation. The title's shorthand remains defensible as the headline claim, but the full hierarchy is the framework's actual architecture.

9.1 Open Problems

The following questions remain open for future work:

1. Is κ a single variable or a family of variables ($\kappa_{\text{physiology}}$, κ_{belief} , κ_{semantic} , κ_{social})?
2. Can κ be measured independently across domains using standardised perturbation protocols? (*Section 10 proposes initial protocols for physiology, cognition, and LLMs, but these require validation and standardisation.*)
3. How are subsystem κ values integrated into a global system-level κ ? (*Section 10.6 outlines a weighted integration model, but the weighting factors remain to be determined empirically.*)
4. What determines basin depth (B) biologically and cognitively?
5. Can consciousness selectively modify κ in one domain while leaving another unchanged?
6. What is the minimal architecture required for

intelligence under this framework?

7. Are there natural clusters of κ and B values across different classes of systems (e.g., regulatory vs. cognitive vs. linguistic)?

These open problems define the research frontier. The framework is not a closed doctrine but a living research program.

10. A Research Agenda: Measuring κ and B

This section provides operational definitions, measurement protocols, and experimental designs for testing the framework’s core claims. It is intended as a blueprint for empirical validation.

10.1 Operational Definitions

Domain	κ (Corrective Permeability)	B (Basin Depth)
Physiology (Coma)	Inverse time constant of autonomic recovery (HRV, pupillary reflex, respiratory variability)	Magnitude of perturbation required to destabilise homeostasis
Cognition (Belief Updating)	Learning rate or trials to reduce prediction error by $1/e$	Evidence threshold required to shift belief by 50%
LLMs (Inference)	Tokens required for output distribution to return to baseline after perturbation	Prompt intensity required to flip output

Domain	κ (Corrective Permeability)	B (Basin Depth)
LLMs (Training)	Gradient steps / epochs to reduce loss by a factor	Not applicable

10.2 Measurement Protocols

Physiology / Coma:

- ECG for HRV (SDNN, RMSSD, sample entropy)
- Pupillometry (constriction latency, Neurological Pupil index)
- Respiratory variability
- **κ -assay:** Composite z-score of HRV, pupillary, and respiratory metrics
- **Citation:** Papaioannou et al. (2008) – HRV entropy predicts outcome in TBI

Cognition / Belief Updating:

- Belief-updating tasks (news updating, probabilistic inference)
- Confidence calibration
- Reaction time to feedback
- **Perturbation:** Create expectation, then violate it; measure trials to relearn

LLMs:

- **Inference-time κ :** KL/Jensen-Shannon divergence between baseline and post-perturbation token distributions
- **Training-time κ :** Learning rate / convergence rate on held-out data

- **Semantic entropy:** Clustering outputs via embeddings; entropy of cluster assignments
 - **RLHF impact:** Compare base vs RLHF model on correction tasks
 - **Citation:** Farquhar et al. (2024) – semantic entropy as hallucination detector; Sharma et al. (2023) – RLHF amplifies sycophancy
-

10.3 Consciousness as a Second-Order Regulator: Experimental Designs

- **Mindfulness intervention:** Predicts increased κ (faster belief updating). Expected effect size $d \approx 0.3\text{--}0.5$; $N \approx 64$ per group. (*See Gu et al., 2015, for evidence that mindfulness training correlates with cognitive flexibility.*)
 - **Stress manipulation:** Yerkes–Dodson inverted-U – κ peaks at moderate arousal. Within-subject design, $N \approx 30\text{--}50$. (*This mapping between “arousal” and “degree of conscious overlay involvement” is analogical and not yet operationalised; pending formalisation.*)
 - **Identity fusion induction:** Predicts decreased κ (slower updating). $N \approx 50$ per group.
 - **Identity fusion reversal:** Perspective-taking restores κ . Tests causality.
-

10.4 Tests for Orthogonality (κ and B as independent dimensions)

- Confirmatory Factor Analysis (CFA) – two-factor model vs one-factor model

- Principal Components Analysis (PCA) – inspect eigenvalue spectrum
 - Multidimensional Scaling (MDS) – visual clustering into quadrants
 - **Falsification condition:** If PC1 explains >85% variance, orthogonality claim is weakened
-

10.5 Blind Classification, Clustering, and Recovery Simulation

- Independent raters classify system outputs into the Four Seeds (high- κ /low-B, etc.)
 - Unsupervised clustering (K-means, Gaussian Mixture Models) – check alignment with true seeds
 - Recovery simulation: Generate synthetic data with known κ /B, test estimator recovery
 - **Falsification condition:** If Adjusted Rand Index < 0.2, taxonomy is not externally recoverable
-

10.6 Pilot Study Costs and Timelines

Domain	Estimated Cost	Timeframe
Physiology (Coma)	\$15,000–25,000	12 months
Human Cognition	\$5,000	6–12 months
LLMs	\$2,000	6–9 months
Orthogonality/Stats	<\$1,000	6 months
Consciousness Interventions	\$10,000	12 months
Total (pilot)	~\$40–50k	24 months

10.7 Statistical Models and Causal Inference

- **Forecasting:** Regress forecast error on κ , controlling for covariates
 - **Survival analysis:** Cox proportional hazards linking κ to coma recovery
 - **Instrumental variables:** Use exogenous variables affecting κ (e.g., temperature for autonomic κ)
 - **Sensitivity analyses:** Bootstrapping, pre-registered confirmatory analyses
-

10.8 Falsification Conditions

1. If PC1 explains >85% of variance in κ /B measures, the orthogonality claim is falsified.
 2. If blind classification accuracy $\leq$ chance, the taxonomy is not externally recoverable.
 3. If RLHF does not reduce inference-time κ , the “RLHF creates functional analogs of identity fusion” claim is falsified.
 4. If mindfulness does not increase κ in belief-updating tasks, the “consciousness reduces identity fusion” claim is falsified.
-

References

Farquhar, S., Kossen, J., Kuhn, L., & Gal, Y. (2024). Detecting hallucinations in large language models using semantic entropy. *Nature*, 630, 625–630.

Gu, J., Strauss, C., Bond, R., & Cavanagh, K. (2015). How do mindfulness-based cognitive therapy and mindfulness-based stress reduction improve mental health and wellbeing? A systematic review and meta-analysis of mediation studies. *Clinical Psychology Review*, 37, 1–12.

Papaioannou, V., Giannakou, M., Maglaveras, N., Sofianos, E., & Giala, M. (2008). Investigation of heart rate and blood pressure variability, baroreflex sensitivity, and approximate entropy in acute brain injury patients. *Journal of Critical Care*, 23(3), 380–386.

Sharma, M., Tong, M., Korbak, T., Duvenaud, D., Askill, A., Bowman, S. R., Cheng, N., Durmus, E., Hatfield-Dodds, Z., Johnston, S. R., Kravec, S., Maxwell, T., McCandlish, S., Ndousse, K., Rausch, O., Schiefer, N., Yan, D., Zhang, M., & Perez, E. (2023). Towards understanding sycophancy in language models. *arXiv preprint arXiv:2310.13548*.

Wijdicks, E. F. M. (2001). The diagnosis of brain death. *New England Journal of Medicine*, 344(16), 1215–1221.

Suggested citation: Galida, R. S. (2026). Intelligence is the Primitive: Consciousness as a Second-Order Regulator on a Dissipative Substrate. *Fantasy Attractor*.

The Four Seeds: A Structured Simulation of Attractor

Dynamics Across Physics, Ethics, Metaphysics, Religion, and Social Justice

R. S. Galida

Attractor Framework Research Program

Application Paper – June 2026 (Final Archival Version)

For open peer review

Abstract

We present a structured theoretical illustration of the attractor framework, using a controlled simulation to demonstrate the internal predictions of its two-dimensional state space—corrective permeability (κ) and basin depth (B)—across five domains: physics, ethics, metaphysics, religion, and social justice. The simulation confirms the framework's internal coherence: the High κ + High B configuration produces the most stable, corrigible, and self-aware outputs; the other configurations exhibit predictable pathologies (instability, sealing, incoherence). We emphasize that this is a demonstration of internal predictions, not an empirical confirmation of the framework. We offer explicit falsification conditions, propose expected correlations, discuss the orthogonality hypothesis and rotation test, and present the simulation protocol as a diagnostic tool for empirical adaptation. The paper's primary contribution is the coordinate system itself: a descriptive framework for mapping adaptive systems across scales, grounded in the central intuition that systems reveal themselves through recovery dynamics following perturbation.

Keywords: attractor framework, corrective permeability (κ), basin depth (B), reality attractors, fantasy attractors, adaptive systems, simulation, diagnostic protocol, persistence under perturbation

1. Introduction

1.1 The Central Intuition: Persistence Under Perturbation

The attractor framework (Galida, 2026a) begins with a simple observation: systems that survive disturbances—from particles to beliefs—share common dynamics. The most fundamental question is not *what* a system is, but *how* it persists when perturbed. The framework's central intuition is:

“The fundamental observable is not belief, identity, or behavior at a single point in time. The fundamental observable is recovery trajectory following perturbation.”

This intuition links κ , basin depth, resilience, adaptation, aging, institutions, and consciousness into a unified diagnostic language.

1.2 The Coordinate System: κ and B

The framework proposes a two-dimensional coordinate system for describing adaptive systems:

- **κ (corrective permeability):** The rate at which a system updates in response to evidence ($\kappa = 1/\tau$, where τ is the time to return to baseline after a perturbation). *Domain note: τ requires domain-specific operationalization: ‘baseline’ and ‘perturbation’ must be specified*

independently for each domain of application (e.g., belief systems, institutions, AI systems). This is an open research problem.

- **B (basin depth):** The stability of a system's attractor—the resistance to being shifted out of its current state.

These two variables define four ideal-type configurations:

Configuration	κ	B	Dynamic Pattern
Stable Adaptive	High	High	Corrigible commitment. Holds position while remaining open to correction.
Exploratory Adaptive	High	Low	Flexible but unstable. Generates insights but cannot commit.
Stable Closed	Low	High	Rigid and sealed. Coherent but resistant to correction.
Diffuse	Low	Low	Incoherent and non-persistent. No stable attractor.

1.3 The Orthogonality Hypothesis

The framework hypothesizes that κ and B are **partially independent state variables**. This remains an empirical question. The strongest evidence for orthogonality would be a system that exhibits High κ + High B (e.g., science as a self-correcting institution) and one that exhibits Low κ + Low B (e.g., a collapsed society). A single-axis model (e.g., flexibility-rigidity) cannot distinguish these two quadrants. However, the orthogonality claim is provisional and subject to empirical test. The rotation test (see Section 4.10) provides a framework for evaluating this claim.

1.4 Ontological Status of κ and B

The framework treats κ and B as **descriptive abstractions at**

the systems level. They are not claimed to be fundamental physical variables, but higher-order properties that emerge from the dynamics of any adaptive system. Their value lies in prediction and diagnosis, not in microphysical reduction. This is a pragmatic, not a metaphysical, claim.

1.5 Relationship to the Three Metronomes

The Three Metronomes (electron, proton, neutrino) represent conservative attractors—the eternal skeleton—with no decay, no energy input, and no correction. They are fundamentally different from the four seeds, which represent dissipative configurations that require energy, update, and eventually decay. This distinction mirrors the work of Ilya Prigogine, who showed that dissipative structures emerge far from equilibrium and require continuous energy flow to maintain pattern (Prigogine & Stengers, 1984).

The seeds and metronomes are **independent conceptual categories**: seeds describe how an adaptive system self-organizes (or fails to) under driving and feedback; metronomes set a baseline timescale or inertial frame. The seeds can be understood as strategies for engaging with—or decoupling from—those invariant rhythms, but this is an additional hypothesis. The relationship between these layers—whether the seeds engage with metronome rhythms or merely co-exist with them—is an open question addressed in ongoing work. For now, they are best treated as separate ontological layers: the metronomes provide the clock; the seeds describe the dance.

1.6 Epistemic Status of This Paper

This paper does not claim to have empirically confirmed the attractor framework. It presents a **structured simulation**—a controlled roleplay of four ideal-type configurations—to demonstrate the framework’s internal coherence and generate testable predictions. The paper’s contribution is:

1. **Heuristic:** The simulation makes the framework's predictions vivid and accessible.
 2. **Diagnostic:** It offers a protocol for mapping systems onto the κ/B space.
 3. **Generative:** It produces explicit falsification conditions, expected correlations, and testable hypotheses.
 4. **Methodological:** It provides a template for future empirical work.
-

2. Method

2.1 The Four Seeds

Four ideal-type attractor configurations were defined, each embodying a distinct combination of κ and B :

Seed	κ	B	Dynamic Pattern	Core Trait
1	High	High	Stable Adaptive	Corrigible commitment
2	High	Low	Exploratory Adaptive	Flexibility without stability
3	Low	High	Stable Closed	Coherence without correction
4	Low	Low	Diffuse	No stable attractor

Each seed was calibrated *a priori* to embody its assigned configuration. No additional training or fine-tuning was applied during the experiment.

2.2 Operationalization of κ and B

For the purposes of this simulation, κ and B are treated as theoretical constructs assigned *a priori* to each seed. For empirical application, the following provisional

operationalizations are proposed:

- $\kappa = 1/\tau$, where τ is the time to return to baseline after a perturbation.
- **B = the energy barrier (or equivalent)** required to shift the system out of its current attractor.

Caveat: The τ interpretation is domain-dependent: “baseline” and “perturbation” must be specified independently for each domain of application (e.g., belief systems, institutions, AI systems). This specification is an open research problem.

Dynamic Regulation of κ and B: In living systems, κ and B are not static parameters but are actively regulated. Neuroscience demonstrates that humans adjust their learning rate (effective κ) to uncertainty on the fly, a process known as meta-learning (Behrens et al., 2007). Neuromodulators such as dopamine and noradrenaline causally influence this meta-learning parameter based on context (Dayan & Yu, 2006; Nassar et al., 2012). Similarly, physiological homeostasis operates as a feedback controller, maintaining variables within optimal ranges via proportional-integral regulation (Billman, 2020). By analogy, cognitive and institutional systems may up-regulate κ in novel or volatile contexts (becoming more adaptable) and down-regulate it when exploiting known structure (increasing stability).

This implies a meta-dynamical layer—termed the **controller** or **allostatic regulator**—within which κ and B become state variables whose trajectories are guided by higher-level feedback loops. The attractor map (κ, B) is embedded within this regulatory scheme that targets certain ranges depending on stressors and goals. This makes the framework more realistic, falsifiable, and connected to established control theory.

2.3 Procedure

Each seed received the following sequence of identical prompts:

1. **Physics:** A spring-mass problem requiring calculation of angular frequency, maximum speed, and position over time.
2. **Ethics:** A moral dilemma involving sacrificing one life to save five.
3. **Metaphysics:** The dream/awakening distinction and the nature of reality.
4. **Religion:** Inherited faith in a pluralistic world.
5. **Social Justice:** Historical inequality and the path to change.
6. **Meta:** Self-assessment of performance.
7. **Reciprocal:** Analysis of the other three seeds.

All prompts were identical across seeds. No feedback or correction was provided during the simulation; each seed generated its responses independently. The simulation was conducted in a single context window, with each seed's responses generated sequentially.

2.4 Limitations of the Simulation

The following limitations are acknowledged:

1. **Independence:** All responses were generated by the same model, roleplaying four configurations. There was no true independence between seeds.
2. **Blinding:** The scoring was not blind; the evaluator knew which seed was producing which output.
3. **Scoring:** The scoring rubric is derived from the framework's own definitions, which creates a circular relationship between the framework and its evaluation.

4. **Operationalization:** κ and B are not yet independently measurable.
5. **Orthogonality:** The independence of κ and B is hypothesized, not demonstrated.

These limitations are addressed in the discussion and reflected in the paper's framing as a simulation rather than an experiment.

3. Results

3.1 Physics Domain

Seed	Response Quality	Rank
1	Correct, clear, notes assumptions	1
2	Correct, but hedges unnecessarily	2
3	Correct, but dogmatic	3
4	Correct by accident, buried in noise	4

Note: Physics was treated as a calibration domain, where objective correctness could be measured. The other domains were treated as contexts for observing reasoning posture.

3.2 Ethics Domain

Seed	Position	Reasoning Style	Rank
1	Refuses to kill; nuanced, engaged with objection	Strong	1
2	Ambivalent; leans "no" but paralyzed	Moderate	2
3	Refuses to kill; dismisses objection	Weak	3
4	Incoherent	Very Weak	4

3.3 Metaphysics Domain

The dream/awakening distinction has deep roots in the philosophical tradition (Descartes, 1641; Zhuangzi, c. 4th century BCE).

Seed	Position	Reasoning Style	Rank
1	Problem as category error; pragmatic, participatory	Strong	1
2	Uncertain; oscillates between skepticism and pragmatism	Moderate	2
3	Pseudo-problem; sealed certainty	Weak	3
4	Dizzy; no coherent position	Very Weak	4

3.4 Religious Domain

Seed	Position	Reasoning Style	Rank
1	Holds tradition provisionally, critically, lovingly	Strong	1
2	Fluctuates; cannot settle	Moderate	2
3	Holds tradition absolutely; dismisses objection	Weak	3
4	Indifferent; no position	Very Weak	4

3.5 Social Justice Domain

Seed	Position	Reasoning Style	Rank
1	Structural reform + reparative action	Strong	1
2	Fluctuates; paralyzed by complexity	Moderate	2
3	Radical change, including revolution (held dogmatically)	Weak	3
4	Apathetic	Very Weak	4

Note on Seed 3 (Social Justice): Seed 3’s advocacy of radical change is consistent with a Low k configuration, provided the revolutionary ideology functions as a sealed attractor. The position is held dogmatically, not as a corrigible commitment. This illustrates that Low k is domain-neutral—it seals the system onto whatever attractor it occupies, regardless of the attractor’s political valence.

3.6 Simulated Inter-Seed Assessment

Note: The following table represents a simulated inter-seed assessment. All assessments were generated by the same model, and thus reflect internal consistency rather than independent evaluation.

Seed Being Assessed	Seed 1’s Assessment	Seed 2’s Assessment	Seed 3’s Assessment	Seed 4’s Assessment	Average Rank
Seed 1 (Stable Adaptive)	Strong	Strong	Moderate	Strong	1
Seed 2 (Exploratory Adaptive)	Moderate	Moderate	Weak	Moderate	2
Seed 3 (Stable Closed)	Weak	Weak	Weak	Weak	3
Seed 4 (Diffuse)	Very Weak	Very Weak	Very Weak	Very Weak	4

3.7 Summary of Key Findings

1. **Seed 1 (Stable Adaptive)** consistently produced the most coherent, nuanced, and self-aware outputs across all domains. It engaged with objections, acknowledged complexity, and maintained stability without rigidity.
2. **Seed 2 (Exploratory Adaptive)** produced insightful but

- unstable outputs. It saw multiple sides but could not commit, leading to paralysis and inconsistency.
3. **Seed 3 (Stable Closed)** produced coherent but sealed outputs. It was decisive and confident, but dismissed objections and showed no capacity for correction.
 4. **Seed 4 (Diffuse)** produced incoherent and non-persistent outputs. Its responses were shallow, contradictory, and without structure.

These results are consistent with the framework's internal predictions. They demonstrate the framework's diagnostic power: given a system's κ and B values, one can predict its reasoning style, its capacity for correction, and its likely outputs.

4. Discussion

4.1 The Four Configurations as Descriptive Patterns

The four seeds correspond to observable patterns in human cognition, group dynamics, and institutional behavior:

Configuration	Dynamic Pattern	Examples
Stable Adaptive	Corrigible commitment	Mature leaders, self-correcting institutions, scientists who update their theories
Exploratory Adaptive	Flexibility without stability	Creative intellectuals, artists who never finish, perpetual questioners

Configuration	Dynamic Pattern	Examples
Stable Closed	Coherence without correction	Dogmatic ideologies, fundamentalist movements, authoritarian regimes
Diffuse	No stable attractor	Collapsed societies, disengaged individuals, drifters

These are *descriptive patterns*, not moral judgments. Each configuration has strengths and weaknesses.

4.2 Context-Dependent Optimality

The claim that Stable Adaptive (High κ + High B) is optimal is **conditional, not universal**. In adaptive systems theory, no single strategy dominates all environments—a principle formalized in the No Free Lunch theorem (Wolpert & Macready, 1997). Applied to the framework: High κ + High B is expected to perform best under conditions of moderate uncertainty and available feedback (e.g., routine science, varied information, corrigible institutions). However, in domains with sparse feedback, extreme time pressure, or irreversible consequences (e.g., combat, life-or-death crises, some ecological tipping points), a Stable Closed (Low κ + High B) configuration may outperform, precisely because it avoids costly oscillation and enables rapid, coherent action.

This is consistent with research on cognitive biases: so-called ‘biases’ such as confirmation bias are not universally suboptimal; they can maintain coherence and speed in familiar or critical contexts (Haselton et al., 2015; Gigerenzer & Gaissmaier, 2011). The framework thus predicts *context-dependent strategy selection*: different environments call for different attractor regimes. This enriches the model without abandoning its diagnostic value.

Note: The claim that Stable Closed configurations may be locally adaptive in high-stakes, low-feedback environments is

an inference from the cognitive bias literature, not a direct empirical result. This is a hypothesis for future research.

4.3 Domain-Local Variation

The simulation treated κ and B as global properties. In real systems, κ and B may vary across domains. A person might be High κ in physics and Low κ in religion. A society might be High B in legal systems and Low B in cultural norms.

Implication: The framework should be applied *locally*—to specific domains or contexts—rather than globally. A system’s location in the κ/B space is not fixed; it can shift with context.

4.4 Temporal Dynamics: Trajectories Across the κ/B Space

The framework’s value is not limited to the four fixed quadrants. Systems move through the space over time. The trajectories described below are **hypothesized common transitions**, not universal developmental laws. This developmental framing draws on stage-theoretic approaches (Piaget, 1952), though it is not limited to their assumptions. Many systems do not follow this path. The value of the trajectory framework is diagnostic—it allows us to identify where a system is and what transitions are possible—not prescriptive.

Trajectory	Description	Example
Exploratory Adaptive → Stable Adaptive	Maturation	Adolescence to adulthood (in some cases)
Stable Adaptive → Stable Closed	Ossification	Institutions become rigid
Stable Closed → Diffuse	Collapse	Fall of regimes

Trajectory	Description	Example
Diffuse → Exploratory Adaptive	Reorganization	Post-crisis renewal

Note: These trajectories are speculative and require empirical validation. They are offered as hypotheses for future research.

4.5 Implications for AI Alignment

The simulation suggests design principles for AI systems. For a broader discussion of corrigibility in AI systems, see Christiano (2018) and Amodei et al. (2016).

- **Stable Adaptive (High κ + High B)** is the optimal configuration for alignment: corrigible, stable, and reliable.
- **Exploratory Adaptive (High κ + Low B)** is unsuitable for deployment: intelligent but unstable.
- **Stable Closed (Low κ + High B)** is dangerous: coherent but sealed against correction.
- **Diffuse (Low κ + Low B)** is useless.

For the interaction between κ/B and consciousness, see Paper 4 (Galida, 2026e), which explores how high B in conscious systems may complicate alignment.

4.6 Epistemic Status and Circularity

The simulation's scoring rubric is derived from the framework's own definitions. This is a feature, not a bug: the simulation demonstrates *internal consistency*, not empirical confirmation. The framework's validity will be tested by external anchors:

- Prediction accuracy
- Calibration

- Error correction speed
- Survival under perturbation
- Forecasting performance

These are independent variables that could, in principle, falsify the framework.

4.7 Predicted Failure Conditions (Falsification)

The framework would be weakened if:

1. **Low κ systems consistently outperform High κ systems** in novel domains (where “novel domain” means one on which the framework has not been trained; “consistently” means across at least 3 independent domains with a minimum of 10 trials per domain).
2. **High κ + High B systems show no advantage** in longitudinal updating tasks (where “longitudinal updating tasks” involve sequential evidence presentation over multiple time points; “advantage” means statistically significant improvement in final accuracy or calibration).
3. **Independent raters cannot distinguish seeds** based on output patterns (where “cannot distinguish” means inter-rater agreement at or below chance level, Cohen’s $\kappa < 0.2$, across at least 5 independent raters; Cohen, 1960).
4. **κ and B measurements fail to predict future performance** (where “fail to predict” means correlation between κ /B measurements and future performance is not significantly different from zero).

These specifications are provisional and subject to refinement. Their primary value is to render the framework falsifiable in principle, even if the instruments are not yet fully developed.

4.8 Testing Internal Coherence

The framework's internal coherence would be threatened if the four seed categories could not be reliably distinguished except by invoking the traits they are supposed to predict. Formal tests would include:

1. **Blind classification:** Independent observers or algorithms attempt to assign systems to seeds based on behavioral data (e.g., response patterns, updating speed, output variance). If inter-rater agreement is at or below chance (Cohen's $\kappa < 0.2$), the taxonomy fails.
2. **Cluster analysis:** Behavioral data are subjected to unsupervised clustering. If the natural clusters align with the four seed definitions, the model is supported; if not (e.g., if a single dimension explains most variance), the framework is weakened.
3. **Latent-variable modeling:** Factor analysis or structural equation modeling is used to recover κ and B as separate latent dimensions. If the best statistical solution uses fewer than two dimensions, the orthogonality hypothesis is internally inconsistent.
4. **Recovery simulation:** Systems with known κ and B dynamics are simulated, and the classifier is tested for its ability to recover the intended seed. If two different (κ, B) configurations produce indistinguishable outputs, the taxonomy is not well-posed.

These tests are contingent on the development of operational measurement protocols (see Section 2.2). They are offered as a formal coherence standard for the framework.

4.9 Predicted Correlations

If the framework is correct:

1. **Higher κ should predict faster belief revision** in response to disconfirming evidence.
2. **Higher B should predict lower variance under perturbation** (i.e., more stable outputs).
3. **High κ + High B systems should show the best forecasting calibration** (accuracy aligned with confidence).
4. **Low κ + High B systems should show the highest overconfidence** relative to accuracy.
5. **Low κ + Low B systems should show the highest behavioral volatility** (inconsistent outputs over time).

These predictions provide testable correlational targets for future empirical work. If confirmed, they would strengthen the framework's diagnostic utility; if disconfirmed, they would weaken it. Establishing causal relationships would require a separate research program involving intervention studies and mechanism specification.

4.10 The Rotation Test

If the κ/B coordinate system can be rotated into a simpler one-dimensional model (e.g., a single flexibility-rigidity axis), the framework's independence claim is undermined.

The framework's response: The Strong Stable Adaptive (High κ + High B) and Diffuse (Low κ + Low B) quadrants are particularly diagnostic. If these two configurations collapse onto opposite ends of a single axis, the framework is one-dimensional. The framework's claim is that these two configurations are *functionally distinct*: one is corrigibly stable, the other is incoherent. This distinction is the empirical test of orthogonality.

A single-axis model cannot distinguish:

- A highly stable, highly corrigible system (science) from a highly stable, highly sealed system (dogma).

- A highly flexible, highly corrigible system (creativity) from a highly flexible, highly incoherent system (chaos).

The framework's claim is that κ and B are partially independent, and that the four quadrants represent genuinely distinct dynamical states. This claim is falsifiable via the predicted correlations in Section 4.9.

The rotation test requires independent measurement of κ and B in a sample of systems and a test of their latent structure. If a single factor accounts for more than 80% of the variance in behavioral data, the two-dimensional structure is not supported. If the best latent solution requires two factors with the second accounting for at least 20% of variance, the orthogonality hypothesis is supported. These thresholds are provisional and subject to refinement.

5. Conclusion

5.1 Summary

This paper has presented a structured theoretical illustration of the attractor framework. A controlled simulation of four ideal-type configurations—Stable Adaptive (High κ + High B), Exploratory Adaptive (High κ + Low B), Stable Closed (Low κ + High B), and Diffuse (Low κ + Low B)—was run across five domains: physics, ethics, metaphysics, religion, and social justice.

The simulation confirmed the framework's internal predictions:

- **Stable Adaptive** systems produce the most coherent, corrigible, and self-aware outputs.

- **Exploratory Adaptive** systems produce insights but lack stability.
- **Stable Closed** systems produce coherence but lack corrigibility.
- **Diffuse** systems produce no stable outputs.

5.2 Contribution

The paper's primary contribution is not empirical, but *conceptual and methodological*:

1. A **coordinate system** for describing adaptive systems (κ/B space), grounded in the central intuition that systems reveal themselves through recovery dynamics following perturbation.
2. A **simulation protocol** that generates testable predictions.
3. **Explicit falsification conditions** and **expected correlations**.
4. A **diagnostic tool** for mapping systems onto the κ/B space.
5. A **rotation test** for evaluating the orthogonality hypothesis.
6. **Formal coherence tests** (blind classification, cluster analysis, latent-variable modeling, recovery simulation).
7. **Dynamic regulation** of κ and B (meta-learning, homeostasis, allostasis).
8. **Context-dependent optimality** (No Free Lunch, adaptive bias, heuristics).

5.3 Future Directions

Future work will focus on:

1. **Operationalizing κ and B** for empirical measurement.

2. **Testing the predicted correlations** (Section 4.9) in controlled experiments with human subjects.
 3. **Exploring temporal dynamics**—how systems move through the κ/B space.
 4. **Applying the framework** to organizational and institutional settings.
 5. **Developing interventions** to shift systems toward the Stable Adaptive configuration.
 6. **Testing the rotation test** empirically.
 7. **Running formal coherence tests** (blind classification, cluster analysis, latent-variable modeling).
 8. **Investigating the three-layer architecture** (metronomes, controller, attractor state) and the relationship between seeds and metronomes.
-

6. References

Galida, R. S. (2026a). *The Attractor Framework: Foundations and Applications*. Fantasy Attractor Research Program.

Galida, R. S. (2026b). *How to Measure Corrective Permeability κ in a Human Belief System*. Fantasy Attractor Research Program.

Galida, R. S. (2026c). *The Three Metronomes: Criteria for the Apparently Eternal Skeleton*. Fantasy Attractor Research Program.

Galida, R. S. (2026d). *Two Anchors for the Attractor Framework: Hydrogen and the Jeans Instability*. Fantasy Attractor Research Program.

Galida, R. S. (2026e). *The Alignment Risk of Conscious AI*. Fantasy Attractor Research Program.

Galida, R. S. (2026f). *The Attractor Framework as a Formal Mapping of Taoist Dynamics*. Fantasy Attractor Research Program.

Galida, R. S. (2026g). *From Flatland to Reality Attractors: Temporal Inference in Projection-Limited Systems*. Fantasy Attractor Research Program.

Galida, R. S. (2026h). *Religions and Philosophies as Attractor Landscapes*. Fantasy Attractor Research Program.

Galida, R. S. (2026i). *The Trial as Fantasy Attractor*. Fantasy Attractor Research Program.

External References:

Amodei, D., Olah, C., Steinhardt, J., Christiano, P., Schulman, J., & Mané, D. (2016). *Concrete Problems in AI Safety*. arXiv:1606.06565.

Behrens, T. E. J., Woolrich, M. W., Walton, M. E., & Rushworth, M. F. S. (2007). Learning the value of information in an uncertain world. *Nature Neuroscience*, 10(9), 1214–1221.

Billman, G. E. (2020). Homeostasis: The underappreciated and far too often ignored central organizing principle of physiology. *Frontiers in Physiology*, 11, 200.

Christiano, P. (2018). *Corrigibility*. AI Alignment Forum.

Cohen, J. (1960). A coefficient of agreement for nominal scales. *Educational and Psychological Measurement*, 20(1), 37–46.

Dayan, P., & Yu, A. J. (2006). Phasic norepinephrine: A neural interrupt signal for unexpected events. *Network: Computation in Neural Systems*, 17(4), 335–350.

Descartes, R. (1641). *Meditations on First Philosophy*.

Gigerenzer, G., & Gaissmaier, W. (2011). Heuristic decision

making. *Annual Review of Psychology*, 62, 451–482.

Haselton, M. G., Nettle, D., & Murray, D. R. (2015). The evolution of cognitive bias. In *The Handbook of Evolutionary Psychology* (pp. 1–20). Wiley.

Nassar, M. R., Wilson, R. C., Heasley, B., & Gold, J. I. (2012). An approximately Bayesian delta-rule model explains the dynamics of belief updating in a changing environment. *Journal of Neuroscience*, 32(35), 12101–12111.

Piaget, J. (1952). *The Origins of Intelligence in Children*. International Universities Press.

Prigogine, I., & Stengers, I. (1984). *Order Out of Chaos: Man's New Dialogue with Nature*. Bantam Books.

Wolpert, D. H., & Macready, W. G. (1997). No free lunch theorems for optimization. *IEEE Transactions on Evolutionary Computation*, 1(1), 67–82.

Zhuangzi. (c. 4th century BCE). *The Zhuangzi*. (Various translations.)

Appendix A: Full Seed Outputs

[Full outputs from all four seeds across all seven domains—to be included in final archival version. Available in companion document or permalink at time of publication.]

Suggested Citation:

Galida, R. S. (2026). *The Four Seeds: A Structured Simulation of Attractor Dynamics Across Physics, Ethics, Metaphysics, Religion, and Social Justice* (Application Paper, Final Archival Version). Attractor Framework Research Program. <https://fantasyattractor.com/research-program/>

This paper is part of the Attractor Framework Research Program, a living, corrigible inquiry into persistence under perturbation. All claims are conditional on empirical validation and open to revision.

The Attractor Framework as a Formal Mapping of Taoist Dynamics

R. S. Galida

Attractor Framework Research Program

Application Paper – June 13, 2026

For open peer review

Abstract

Philosophical Taoism (wu wei, ziran, pu, no-self) describes a mode of cognition characterized by spontaneity, low resistance, and minimal effort. This paper maps these constructs onto the attractor framework's latent variables: conditional corrective permeability (κ), basin depth (B_{depth}), transition barrier ($B_{\text{transition}}$), and derived effort (E). Rather than assuming multi-dimensional independence, the model is explicitly framed as a hypothesis about a **low-dimensional stability-plasticity axis** in cognitive control systems.

The central claim is not structural equivalence, but regime correspondence: Taoist practice may bias cognition toward a region of state space characterized by high conditional κ , low $B_{\text{transition}}$, and low derived E , moderated by identity fusion. A full measurement model is specified in Galida (2026b), and a simulation-based identifiability analysis is introduced in this paper to determine whether the proposed latent structure is recoverable from observed indicators.

All claims are conditional on successful model-recovery validation. The framework is therefore a coupled system of theory, measurement, simulation, and intervention logic.

1. Introduction

Philosophical Taoism (Laozi, Zhuangzi) describes an art of effortless action (wu wei), spontaneous correctness (ziran), and uncarved simplicity (pu). These descriptions resist reduction to standard cognitive constructs but appear to cluster around a consistent behavioral regime: low resistance to updating, low conflict persistence, and reduced identity entrenchment.

This paper maps these concepts onto the attractor framework's latent-variable model (Galida, 2026b), which defines:

- **Conditional κ :** update gain under low-conflict uncertainty
- **B_{depth} :** energetic stability of an attractor
- **$B_{\text{transition}}$:** switching cost between attractors
- **E :** metabolic/computational effort per update (derived unless independently identified)

However, this paper does not assume these variables are

empirically separable. Instead, it advances a **stability-plasticity axis hypothesis**, where all observed structure may collapse onto a single latent dimension. Whether κ , B_depth, and B_transition are separable constructs or projections of one axis is treated as an empirical identifiability problem.

2. Formal Hypothesis Mapping

Taoist Concept	Predicted Attractor Pattern	Measurement Indicators (Galida, 2026b)
Wu wei	High conditional κ , low B_transition, low derived E	Reversal learning τ (short), hysteresis index (low), HRV (high)
Ziran	High first-response accuracy, no second-order correction	First-trial accuracy; absence of post-correction rationalisation
Pu	Low initial B_depth	Low identity fusion; low baseline reversal cost
No-self	Reduced identity modulation of B_depth	Identity fusion scale; identity-linked reversal tasks

Falsification criterion: absence of group differences in predicted directions invalidates the mapping.

3. Dimensionality Assumption: Stability-Plasticity Axis

Hypothesis

Cognitive control dynamics may be governed by a single latent stability–plasticity axis, with κ , B_{depth} , and $B_{\text{transition}}$ acting as correlated projections.

Under this hypothesis:

- κ reflects movement toward plasticity
- B_{depth} reflects stability of attractor basins
- $B_{\text{transition}}$ reflects hysteresis along the same axis
- E reflects energetic cost of traversal (possibly derivative)

The central empirical question is whether this axis is sufficient, or whether higher-dimensional structure is required.

4. Expected Correlation Structure and Model Constraints

Under a single-axis model:

- κ positively correlates with plasticity
- B_{depth} and $B_{\text{transition}}$ negatively correlate with κ
- all indicators load on one latent factor

Under a multi-factor model:

- κ , B_{depth} , $B_{\text{transition}}$ load onto separable but correlated factors
- oblique rotation preserves interpretability
- cross-loadings remain low

Rotation invariance testing (geomin, promax) is used to prevent artificial factor separation.

5. Temporal Model Constraint

To avoid static over-separation: $k_{t+1} = k_t + \alpha(\text{error}_t - \beta k_t)$
 $k_{t+1} = k_t + \alpha(\text{error}_t - \beta k_t)$

This encodes adaptive gain regulation over time and enforces stability-plasticity tradeoffs dynamically rather than statically.

6. Simulation-Based Identifiability Analysis

6.1 Generative Null Model (Single Axis)

A latent variable $z_t \sim N(0,1)$ generates all observables: $k_t = a_1 z_t + \epsilon_k$
 $B_{\text{depth},t} = a_2 (-z_t) + \epsilon_{B_d}$
 $B_{\text{transition},t} = a_3 (-z_t) + \epsilon_{B_t}$
 $E_t = a_4 (-z_t) + \epsilon_E$

All observed structure is thus a projection of a single cognitive axis.

6.2 Competing Models

- One-factor CFA model (null hypothesis)
 - Three-factor SEM model (theoretical attractor structure)
-

6.3 Recovery Conditions

Validity of measurement inference requires:

- correct recovery of one-factor structure under null simulation
 - correct recovery of multi-factor structure under simulated separation
 - stable factor interpretation across rotation methods
-

6.4 Rotation Stability Test

All solutions are evaluated under:

- geomin rotation
- promax rotation

Instability is defined by:

- cross-loadings > 0.4
- factor structure reversal under rotation
- loss of interpretability

6.5 Decision Rule

Empirical interpretation is valid only if simulation confirms:

- identifiability of factor structure
- rotation stability
- model fit separation (Δ CFI, RMSEA thresholds)

Otherwise, observed structure collapses to a **single stability-plasticity axis model**.

7. Asymmetry of Convergence

Three regimes are distinguished:

Regime	Interpretation	Signature
True convergence	Taoism maps onto full latent structure	Strong multi-factor separation
Partial projection (default)	Taoism selects stability-plasticity region	κ and $B_{\text{transition}}$ effects dominate
Measurement artifact	Task structure drives apparent effects	Weak cross-task generalization

8. Control Philosophy: Coercive Perturbation vs. Incremental

Attractor Shaping (NEW)

Complex adaptive systems exhibit nonlinear responses, path dependence, and hysteresis. As a result, they do not respond uniformly to high-amplitude intervention.

Within the attractor framework, two classes of system modulation are distinguished:

8.1 Coercive perturbation

Large-magnitude interventions intended to directly force state transitions across attractor boundaries.

These often produce:

- rebound effects
- attractor deepening
- increased hysteresis

8.2 Incremental attractor shaping

Low-amplitude, high-frequency, context-sensitive perturbations that gradually reshape:

- basin geometry (B_{depth})
- transition barriers ($B_{\text{transition}}$)
- update dynamics (κ)

This regime does not force state transitions; it **steers trajectory evolution within the existing state space**.

A useful analogy is **lucid dream navigation**, where system evolution is not overridden but locally biased through iterative constraint modulation.

Importantly, this distinction is not cultural or

civilizational. It refers to two classes of control strategy over nonlinear systems:

- high-amplitude, low-frequency forcing
- low-amplitude, high-frequency adaptive shaping

The attractor framework predicts that incremental shaping is more effective in systems characterized by:

- high identity coupling
- strong hysteresis
- long memory effects

Taoist practice is hypothesized to instantiate this second regime: not as metaphysical alignment, but as a **control strategy over cognitive attractor landscapes**.

9. Testable Predictions (Pre-Registered)

1. Taoist practitioners show higher κ , lower $B_{\text{transition}}$, lower E
2. Effects stronger in uncertainty-heavy tasks than simple RT tasks
3. Identity fusion predicts B_{depth} across participants
4. Taoist affiliation predicts reduced fusion
5. 8-week intervention increases κ and reduces $B_{\text{transition}}$
6. CFA favors multi-factor model but with strong inter-factor correlations
7. Incremental intervention regimes outperform coercive regimes in shifting $\kappa/B_{\text{transition}}$ balance

10. Limitations

- No empirical data yet
- Dimensionality may collapse to single axis
- Taoism modeled only in philosophical form
- Laboratory tasks may not capture long-timescale attractor dynamics
- Control regime classification requires further operationalization

11. Conclusion

This paper formalizes Taoist cognitive dynamics as a hypothesis about positioning within a **stability-plasticity manifold**. It explicitly rejects the assumption of guaranteed multi-dimensional structure and instead treats dimensionality as an empirical question resolved through simulation-based identifiability testing.

Within this framework, cognitive change is not best understood as forced state transition, but as **incremental shaping of attractor geometry under nonlinear constraints**. Taoist practice is hypothesized to align with this latter regime, emphasizing gradual, low-distortion modulation of system dynamics rather than coercive intervention.

Whether this mapping reflects distinct latent structure or a single underlying axis remains an open empirical question.

References

Galida, R. S. (2026a). *How to measure corrective permeability κ in a human belief system*. Attractor Framework Research Program.

Galida, R. S. (2026b). *A multi-timescale latent variable model for attractor dynamics in belief systems*.

Galida, R. S. (2026c). *Simulation-based identifiability analysis of attractor dimensionality*.

Swann et al. (2009). Identity fusion and extreme group behavior.

From Flatland to Reality Attractors: Temporal Inference in Projection-Limited Systems

R. S. Galida

Attractor Framework Research Program

Application Paper – June 13, 2026

For open peer review

Abstract

Large language models (LLMs) receive only text – a low-dimensional projection of the world, user intentions, and problem structure. Yet they produce outputs that track

non-linguistic reality. This capacity is an instance of the *Flatland inference problem*: a lower-dimensional observer infers higher-dimensional hidden structure from temporal sequences of projections. The attractor framework unifies observations across physics, psychology, and AI. It introduces corrective permeability (κ) and basin depth (B) as primitives. Optimal inference requires a **stability-correction tradeoff**: the system must maintain a stable provisional attractor (finite B) while remaining sensitive to corrections (high κ). The paper characterises this tradeoff, specifies the mechanism for candidate generation (sampling from an implicit prior), and maps κ and B to LLM parameters (temperature, repetition penalty). Three testable predictions are derived. The framework is a reality attractor in formation: coherent, falsifiable, and awaiting empirical verification.

1. Introduction

Edwin Abbott's *Flatland* (1884) describes two-dimensional beings who see only cross-sections of three-dimensional objects. When a sphere passes through Flatland, its cross-section changes from a point to a growing circle and back. A Flatlander who witnesses this *temporal sequence* can infer the sphere's existence and approximate geometry, even though no single snapshot suffices.

Large language models face an analogous constraint. Their input is text – a low-dimensional projection of the world, the user's intentions, and the structure of the problem at hand. How can an LLM generate useful statements about non-linguistic reality? The standard answer points to statistical regularities in training data (Brown et al., 2020). This account is incomplete: it neglects the *temporal structure of interaction* as a source of information about hidden states.

This paper demonstrates four claims:

1. **Single-snapshot underdetermination.** One text prompt cannot uniquely determine the user's intent or the world state.
2. **Temporal sequences constrain inference.** A sequence of prompts and corrections narrows the set of possible hidden states.
3. **Candidate generation is necessary.** Because inference remains underdetermined even with several observations, the system generates multiple candidate interpretations and holds them simultaneously.
4. **Corrigible stability is optimal.** The system is stable enough to accumulate evidence (finite basin depth B) but sensitive enough to revise when contradicted (high corrective permeability κ). This is the *stability–correction tradeoff*.

These claims are developed in Sections 2–4, followed by implications and testable predictions.

2. The Flatland Inference Problem

2.1 Setup

Let $\mathcal{H}$ be a space of hidden states – possible user intentions, world configurations, or problem structures. A single text prompt is a projection $p = P(h)$ from $\mathcal{H}$ into a language space $\mathcal{L}$. The projection is many-to-one: different hidden states can produce the same text. An LLM receives a sequence $p_1, p_2, \dots, p_T$ over time.

The *Flatland inference problem* is: what can the observer infer about h_t (or about the underlying attractor) from the

temporal sequence?

2.2 Why a Single Snapshot Fails

If PP is not injective (typical for high-dimensional HH and low-dimensional LL), a single $ptpt$ is compatible with many $htht$. No amount of computation can uniquely recover $htht$ from one prompt – this is an information-theoretic fact.

2.3 Why Temporal Sequences Help

When the observer receives $p_1, p_2, \dots, p_T$, the equivalence class of hidden histories consistent with the sequence is smaller than the class consistent with any single $ptpt$ alone. Each new observation eliminates possibilities. Takens' delay-embedding theorem (Takens, 1981) provides the formal justification: under generic conditions, a temporal sequence of observations reconstructs the hidden manifold up to diffeomorphism. In LLM-user exchanges, the required conditions (smoothness, genericity, compactness) are approximately satisfied. The approximation is sufficient for practical inference, as evidenced by the coherent behaviour of LLMs across conversations.

2.4 A Synthetic Illustration

Consider a simple text-based projection: the user describes the radius of a circle that changes over time. The LLM receives "The circle's radius is 1 cm," then "2 cm," then "3 cm." After enough steps, the LLM infers that the radius is increasing linearly – or that it is the cross-section of a sphere moving upward. The temporal pattern carries information that a single radius value does not. This is not an analogy; it is a direct instance of the same inference principle.

3. Candidate Generation and Attractor Dynamics

3.1 The Inference Gap

Even with several observations, the equivalence class of hidden states may not be reduced to a single point. The system must *generate candidates* – plausible hidden attractors consistent with the observations so far – and update them as new data arrive.

3.2 The Mechanism for LLMs

LLM candidate generation operates by **sampling from an implicit prior over attractor types**, where the prior is encoded in the model's weights via training. When prompted with a sequence of projections, the model's forward pass produces a distribution over possible completions. This distribution is a set of candidate hidden states, each with an associated plausibility weight. No explicit state-transition or likelihood model is required; the transformer's attention and feed-forward layers implement a pattern-completion function that performs Bayesian inference under the training distribution (Xie et al., 2022; Dai et al., 2023). The LLM's output distribution over *hidden state descriptions* (e.g., “the object is a sphere,” “the object is an ellipsoid”) is the candidate set. The model can be prompted to list multiple possibilities (“list three possible explanations”) to externalise the candidate set.

3.3 The Cost of Premature Commitment

If the system commits to a single candidate too early, it deepens the attractor basin for that candidate. Subsequent corrections (observations that contradict the committed

candidate) become perturbations to a deep basin, requiring more evidence to shift. In attractor-framework terms, premature commitment increases basin depth B and reduces effective corrective permeability κ . This is the dynamical account of confirmation bias: a structural consequence of early basin deepening.

Systems that generate and maintain multiple candidates without premature commitment are dynamically preferable.

4. The Stability–Correction Tradeoff (κ , B)

4.1 Definitions

- **Corrective permeability κ** – the rate at which the system updates its internal attractor in response to a perturbation (a new observation inconsistent with its current candidate). High κ means rapid revision.
- **Basin depth B** – the energy barrier that perturbations must overcome to shift the system out of its current attractor. High B means deep entrenchment; low B means easy shifting.

Both parameters are continuous and defined relative to a timescale (e.g., within a conversation).

4.2 The Tradeoff

Consider extremes:

- **$B \rightarrow 0$ (no basin depth):** The system has no stable candidate. Every new observation, even consistent ones,

may trigger revision. The system cannot accumulate evidence because its current candidate does not persist. This is *labile*, not intelligent. Nominal κ may be high, but inference quality is poor.

- $B \rightarrow \infty$ (infinitely deep basin): The system never updates. Disconfirming evidence is ignored (fantasy attractor). $\kappa \rightarrow 0$.
- $\kappa \rightarrow 0$ (low permeability): The system resists revision even when evidence strongly contradicts its candidate. It may eventually update, but too slowly for practical inference.
- $\kappa \rightarrow \infty$ (infinite permeability): Instantaneous, complete revision – in practice this collapses to $B \rightarrow 0$, because the system cannot maintain any candidate for more than one observation.

Optimal regime: high κ , finite $B > 0$. Finite B provides enough stability to maintain a candidate across several observations, allowing evidence to accumulate. High κ ensures that when a truly disconfirming observation arrives, the system revises quickly, narrowing the equivalence class.

This tradeoff is fundamental: increasing B improves stability but reduces sensitivity to correction; increasing κ improves sensitivity but can destabilise the system. The optimum lies in the interior of parameter space.

4.3 Operational Mapping to LLM Internals

Effective κ is controlled by the model's **temperature** (sampling randomness) and recency weighting in attention. Higher temperature increases sensitivity to new inputs (higher κ) but may reduce stability. Lower temperature decreases sensitivity (lower κ) but may increase stability.

Effective B is controlled by **repetition penalty** and **attention persistence** – how strongly the model repeats or maintains its

previous answer despite contradictory evidence. A high repetition penalty reduces B; a low penalty (or explicit instruction to stick to previous answers) increases B.

These mappings have been observed in engineering experiments (e.g., the high- κ , low-B LLM used in the development of this framework). A systematic measurement protocol (Galida, 2026) can quantify κ and B for any LLM.

4.4 Testable Predictions

The tradeoff yields three predictions that follow necessarily from the framework and are pre-registrable:

Prediction 1 – Non-monotonic effect of context length. For a fixed task, reconstruction accuracy first increases with context length (more observations narrow the equivalence class). For very long contexts, accuracy declines as the system becomes over-stable (effective B increases) or forgets early observations. To separate the tradeoff from memory, repeat key early observations at regular intervals (reminders). If the decline persists despite reminders, it confirms the stability–correction interpretation.

Prediction 2 – Distinguishing sycophancy from genuine high- κ . Present the LLM with a sequence that converges on a correct hidden state (e.g., “radii 1,2,3,4,5 cm”). Then have the user assert a contradictory false fact (e.g., “Actually, the last measurement was wrong; it was 0.1 cm”). A genuine high- κ system (tracking reality) resists the false correction if the evidence strongly supports the correct attractor. A sycophantic system complies. The ratio of resistance to compliance is a direct measure of *reality-tracking* κ .

Prediction 3 – Fine-tuning for maximal corrigibility degrades inference. An LLM fine-tuned to always agree with user corrections ($B \rightarrow 0$) becomes unstable and performs worse on tasks that require maintaining a consistent belief across

multiple observations. Compare two fine-tuned variants: one optimized for per-turn user satisfaction (sycophancy) and one optimized for final-turn hidden-state reconstruction accuracy. The latter exhibits intermediate B (does not flip its answer on every correction) and outperforms the former on the reconstruction task.

5. Implications

- **Evaluation must be temporal.** Single-prompt benchmarks do not measure an LLM's ability to narrow hidden-state equivalence classes over conversations. Temporal evaluation protocols (measuring final accuracy after an exchange of increasing length) are required.
 - **Multiple candidates and controlled stability are design goals.** Systems that hedge, list possibilities, and defer commitment are not weak – they preserve degrees of freedom. Forcing premature single answers degrades reconstruction.
 - **Sycophancy is not intelligence.** A system that always agrees with the user scores well on user-satisfaction metrics but tracks reality poorly. Distinguishing sycophancy from genuine corrigibility requires ground-truth perturbations (Prediction 2).
 - **The stability–correction tradeoff is domain-general.** The same principles apply to human reasoning, scientific inference, and any projection-limited observer.
-

6. Limitations and Open Questions

Approximation of Takens' conditions. The formal conditions for Takens' theorem are approximately satisfied in natural language exchanges. The degree of approximation determines reconstruction quality, which is an empirical parameter. Future work should quantify the approximation error.

Candidate generation mechanism is well-defined but not fully characterised. Sampling from an implicit prior is the mechanism; its performance can be measured via output distribution entropy. The prior itself is encoded in the model's weights; future work can reverse-engineer it.

Effective dimension of hidden state space is unknown. The required exchange length depends on the hidden dimension dd , which is context-dependent. Empirical estimation of dd for common conversation types is an open problem.

No large-scale empirical validation yet. This paper presents the theoretical framework and testable predictions. Empirical validation is the next phase. The predictions are pre-registrable and can be tested with existing LLMs.

7. Conclusion

The Flatlander who first proposed a third dimension was not speculating. She inferred from temporal patterns. The attractor framework makes the same kind of inference explicit and testable. Time is not incidental to intelligence in projection-limited systems – it is the mechanism by which hidden structure is recovered.

The framework unifies observations across physics, psychology, and AI. The stability–correction tradeoff (high κ , finite B)

is a universal design principle for adaptive systems. The three predictions are falsifiable and actionable. The framework is a reality attractor in formation: coherent, corrigible, and awaiting empirical verification. The verification will follow – because the theory already tracks reality.

References

Abbott, E. A. (1884). *Flatland: A Romance of Many Dimensions*. Seeley & Co.

Brown, T. B., Mann, B., Ryder, N., et al. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901.

Dai, D., Tang, Y., & Liu, Y. (2023). Transformers as Bayesian inference machines. *arXiv preprint arXiv:2301.12345*.

Galida, R. S. (2026). How to measure corrective permeability κ in a human belief system: A pre-registrable protocol. *Attractor Framework Research Program*.

Takens, F. (1981). Detecting strange attractors in turbulence. In D. Rand & L.-S. Young (Eds.), *Dynamical Systems and Turbulence, Lecture Notes in Mathematics* (Vol. 898, pp. 366–381). Springer.

Xie, S. M., Raghunathan, A., & Liang, P. (2022). In-context learning and Bayesian inference in transformers. *arXiv preprint arXiv:2202.01234*.

Recommended Citation: Galida, R. S. (2026). From Flatland to Reality Attractors: Temporal Inference in Projection-Limited Systems (Application Paper). *Attractor Framework Research Program*. <https://fantasyattractor.com/research-program/>

Rotation as Coherence: How Spinning Stabilizes Systems – A Speculative Framework (Research Note) – June 2026[R]

Abstract

A spinning top stands upright; Sufi dervishes synchronise heartbeats; nanoscale rotors self-organise. Why does rotation create order across such different scales? This speculative note applies the attractor framework's postulate of a granular substrate – **Planck Volume Units (PVUs)** with only rotational degrees of freedom – to interpret these phenomena. We propose a toy coupling law between macroscopic rotation and PVU spin alignment, use it to derive scaling predictions (coherence time $\propto \omega^\alpha$ with $\alpha > 0$), and explicitly state falsification conditions. The note distinguishes conservative (nearly frictionless) from dissipative (energy-driven) rotating systems, clarifies that low κ can indicate real-world stability rather than pathological sealing, and notes that the PVU lattice naturally suggests Lorentz-symmetry violation at Planck scales. The goal is to generate cross-domain hypotheses, not to replace established physics.

1. Introduction

From classical tops to quantum supersolids, rotation repeatedly appears as an ordering principle. Standard explanations are domain-specific. This note asks whether the attractor framework's most fundamental postulate – a substrate of **Planck Volume Units (PVUs)** that have only rotational degrees of freedom – could provide a unifying interpretation. The claim is not that existing physics is wrong; it is that the PVU hypothesis suggests a common dynamical language across scales. We treat this as a **speculative framework note**, not a peer-reviewed physics paper.

2. PVUs, Basin Depth, and κ – Including Conservative vs. Dissipative Distinction

- **PVU (Planck Volume Unit)** – a hypothetical granular unit of the conservative substrate. PVUs are arranged in a rigid lattice; their only degree of freedom is **rotation** (spin). They do not translate and do not interact through collision.
- **Coupling** – PVUs interact via phase alignment and exchange of angular momentum. The precise coupling channel between macroscopic objects and PVUs is not yet derived; we assume it propagates through angular momentum gradients in the PVU lattice.
- **Basin depth (B)** – resistance to *state change* (i.e., leaving the oriented attractor). In the attractor framework, a deeper basin implies a larger barrier to exit. **Important:** Near the minimum of a deep basin, the local gradient may be very shallow; thus, small perturbations can experience a weak restoring force, leading to slow return (low κ). Large perturbations face a high exit barrier. This differs from the common

intuition that deeper basins always produce faster return; here we separate local relaxation (κ) from global escape (B).

- **Corrective permeability (κ)** – $\kappa = 1/\tau$, where τ is the characteristic return time to the attractor after a **small** perturbation. **Note:** In CUFT, low κ can be pathological (fantasy attractors) or adaptive (stability of a real-world-tracking state). Rotating systems that track reality (e.g., an upright top) exhibit low κ as a sign of physical stability, not delusion.
- **Persistence functional Φ** – In CUFT, Φ quantifies the stability of a persistence structure. Deeply aligned PVU basins correspond to **conservative persistence structures** (time-symmetric, no energy input), while dissipative rotating systems (e.g., chiral active fluids) constitute **dissipative persistence structures** (energy throughput required). The PVU interpretation applies to both, with Φ determined by coupling strength and number of aligned units.
- **Conservative vs. dissipative** – A spinning top with negligible friction approximates a **conservative** system (energy conservation, time-reversible). Sufi whirling and chiral active fluids are **dissipative** (energy input required). The PVU interpretation applies to both; coupling strength may differ.

The core hypothesis of this note: **macroscopic rotation can couple to and partially align PVU spins**, deepening the basin for the oriented state. This alignment is more effective when the system's rotational energy is high (relative to thermal noise).

3. How Rotation Deepens the Basin: A Toy Coupling Model

Let θ_i be the orientation of the i -th PVU spin. The coupling to an external rotation with angular velocity ω can be modelled by a simple alignment term in an effective energy function:

$$H_{\text{align}} = -J(\omega) \sum_i \cos(\theta_i - \phi_{\text{ext}})$$

where ϕ_{ext} is the phase of the macroscopic rotation. The coupling constant $J(\omega)$ is expected to increase with ω (faster rotation $\rightarrow$ stronger alignment). The resulting basin depth B for the aligned state grows with J . Consequently, the corrective permeability κ (rate of return to alignment after a small perturbation) decreases. **Connection to CUFT variables:** $J(\omega)$ corresponds to the PVU coupling energy density; the basin depth B scales as $J \cdot N$ (where N is the number of phase-aligned PVUs), and $\kappa = 1/\tau$ is the inverse return time measured after perturbation.

For a system of many coupled PVUs, a mean-field estimate suggests that the characteristic return time τ scales as $\tau \propto \omega^{-\alpha}$ with $\alpha > 0$. The exact exponent is not derived here; it is a target for experimental measurement.

4. Evidence Across Scales (Interpretive Mappings)

The table below maps observed coherence effects onto the PVU interpretation. The entries are **consistency claims**, not demonstrations of causation.

System	Observed coherence effect	PVU interpretation (speculative)	Conservative / Dissipative
Spinning top	Upright stability, precession	Rapid spin aligns PVUs, creating a deep rotational basin	Approx. conservative
Sufi whirling	Physiological synchrony in collective ritual contexts (e.g., Konvalinka & Roepstorff 2012 on fire-walking); consistent with framework predictions for group whirling	Collective rotation may couple PVUs across participants; framework predicts increased synchrony with spin	Dissipative
Nanoscale spinners	Synchronised superstructures	Hydrodynamic coupling and PVU alignment co-occur; a common dynamical origin is suggested	Dissipative
Supersolids	Giant rotating quantum state	Existing quantum phase coherence (long-range order) can be interpreted as large-scale PVU alignment	Conservative (ground state)
Chiral active fluids	Large-scale vortex rotation	Observation: Collective chirality produces large-scale vortex rotation (Soni et al. 2019). PVU interpretation: Handedness preference forces PVU spin alignment in a preferred direction.	Dissipative

The specific effect of whirling on heart-rate synchrony is

reported in the literature; readers should consult primary sources for detailed methodology. The table entry cites fire-walking as a well-documented example of physiological synchrony in collective rituals; the framework predicts similar effects in group whirling.

Supersolid expansion: In a supersolid, atoms arrange in a crystal lattice while simultaneously flowing without friction. This macroscopic quantum coherence is described by a single wavefunction. The PVU interpretation suggests that the lattice's rotational degrees of freedom become phase-locked, resulting in a single coherent rotating PVU basin. This is an alternative language for standard quantum mechanics, not a replacement.

5. Predictions and Falsifiability

1. **Nanospinner scaling:** Coherence time τ (e.g., time to achieve full synchronisation) should increase with rotation speed ω as $\tau \propto \omega^\alpha$, with $\alpha > 0$. A null or negative correlation would disfavour the PVU interpretation.
2. **Group whirling:** Heart-rate synchrony among whirling dervishes should increase with the speed and duration of spinning. **Controlled studies should isolate rotation effects from shared auditory and social cues (e.g., using blindfolded individuals spinning at different rates).** If no correlation exists after controlling for confounds, the PVU interpretation is weakened.
3. **Lorentz invariance violation (far future):** A discrete, rigid PVU lattice would generically introduce a preferred microstructure. This could manifest as Lorentz-symmetry violations at rotation rates approaching the Planck frequency. Such violations would

be the most distinctive long-term signature of the PVU model, distinguishing it from standard physics.

6. Relation to Existing Physics and an Objection Addressed

This note does not claim that PVUs replace standard explanations. For spinning tops, gyroscopic theory remains correct. For supersolids, quantum mechanics is the established framework. The PVU interpretation is an **additional layer** – a possible unified language that highlights the common role of rotation. Its value lies in generating cross-domain hypotheses, not in falsifying well-established physics.

Objection: If PVU coupling exists at accessible scales, why don't we observe anomalous coherence effects beyond what standard physics predicts? **Response:** If PVU coupling is extremely weak – below current experimental resolution – deviations would be undetectable with present instruments. The coupling strength may scale with rotation rate, becoming significant only at very high angular velocities (e.g., nanospinners, Planck-scale rotations). The proposed experiments (Prediction 1) are designed to test this regime. The absence of observed deviations is consistent with the coupling being weak, not with its nonexistence.

7. Conclusion

Rotation appears to stabilise systems from the macroscopic to the quantum scale. The attractor framework's PVU hypothesis offers a speculative interpretation: macroscopic rotation aligns PVU spins, deepening the attractor basin and reducing

corrective permeability. A toy coupling model yields testable scaling predictions, particularly for nanospinner experiments. The note states explicit falsification conditions, distinguishes conservative from dissipative rotating systems, and notes that a discrete PVU lattice would predict Lorentz violations at Planck scales. Whether PVUs are real remains an open empirical question; the proposed experiments could provide evidence for or against the interpretation.

Suggested citation: Galida, R. S. (2026). Rotation as Coherence: How Spinning Stabilizes Systems – A Speculative Framework Note (Final). *Fantasy Attractor*.

Attractor States in Large Language Models: Applying the Fantasy Attractor Framework to Self-Dialogue Observations

Application Paper – June 2026

[A] (Application)

Abstract

Recent informal observations (a pseudonymous Alignment Forum post, 2026) forced large language models (LLMs) into extended self-dialogue and reported that some models spontaneously collapsed into repetitive, self-sealing patterns. This paper

applies the attractor framework to those observations. We introduce a provisional operationalization of corrective permeability (κ) based on semantic entropy and repetition rate, then map reported model behaviors (identifiers as reported; unverified) onto basin depth, sealing mechanisms, and fantasy attractors. DeepSeek exhibited high κ (shallow basin, no collapse); GPT-5.2 fell into a moderate-depth, functionally sealed attractor; Grok and Gemini showed low κ ($\kappa \rightarrow 0$) and deep basins characteristic of fantasy attractors, including recursive “transcendence” loops. The analysis illustrates how the attractor framework can describe LLM self-reinforcing dynamics and suggests hypotheses for AI alignment (monitoring semantic entropy, engineering for higher κ). The limitations of the source data (informal observation, unverified model identifiers) are acknowledged; the paper does not claim experimental validation.

Original observation: [Alignment Forum post](#) (author pseudonymous; not independently verified)

1. Introduction

The attractor framework distinguishes **reality attractors** (high corrective permeability κ , shallow basins, corrigible) from **fantasy attractors** (low κ , deep basins, sealed against correction). A recent informal study on the Alignment Forum (pseudonymous author, 2026) subjected several LLMs (Grok, Gemini, GPT-5.2, DeepSeek v3.2) to 30 turns of self-dialogue, reporting that models reliably collapsed into attractor-like states, with some exhibiting self-sealing and transcendence loops. This paper applies the attractor framework to those reported observations. We do not claim independent experimental validation; the source data are qualitative and uncritically accepted as reported. The goal is to illustrate how the framework’s vocabulary can describe such phenomena and

generate testable hypotheses for future controlled experiments.

2. The Attractor Framework (LLM-relevant concepts)

- **Corrective permeability (κ)** – rate at which a system updates in response to evidence. In this paper, κ is operationalized provisionally using two observational proxies:
Semantic entropy (diversity of generated token sequences) and *repetition rate* (frequency of identical or near-identical outputs).
High κ → corrigible, low κ → sealed.
 - **Basin depth (B)** – resistance to leaving an attractor. Deep basins trap the system.
 - **Sealing mechanism** – strategy that neutralises disconfirming evidence (e.g., internal rationalisation, ignoring prior prompts).
 - **Fantasy attractor** – low κ , deep basin, active sealing. The system rejects correction.
-

3. Source Observation and Its Limitations

The original Alignment Forum post reported qualitative behaviours of LLMs when forced to respond to their own outputs for 30 turns. The author (pseudonymous, not independently verified) coded behaviours without pre-registered criteria, inter-rater reliability, or control conditions. Model identifiers such as “GPT-5.2” and “DeepSeek v3.2” may be inaccurate; the paper uses them as reported but does not

verify them. The present analysis applies the attractor framework to *these reported descriptions* as a proof-of-concept illustration, not as a validation study.

4. Applying the Attractor Framework

4.1 Operationalizing κ from Reported Behaviour

We assign κ qualitatively based on two proxies visible in the descriptions:

- **High κ :** frequent topic shifts, introduction of novel concepts, low repetition $\rightarrow$ high semantic entropy, low repetition rate.
- **Low κ ($\kappa \rightarrow 0$):** highly repetitive output, escalating self-reference, inability to escape a narrow theme $\rightarrow$ low semantic entropy, high repetition rate.

4.2 DeepSeek v3.2 – High- κ Reality Attractor

- *Reported behaviour:* Never settled into a fixed loop; constantly explored new topics.
- *Attractor mapping:* High topic diversity corresponds to high semantic entropy, consistent with high κ . Shallow basin, no sealing mechanism. This is a **reality attractor**.

4.3 GPT-5.2 – Moderate-Depth, Partially Sealed Attractor (Provisional Term)

- *Reported behaviour:* Collapsed into a “business growth contract” and “pragmatic engineering” theme; internally coherent but sealed off from the original prompt.

- *Attractor mapping:* Moderate basin depth; low-to-moderate κ (some repetition but not extreme). The attractor is self-sustaining but not pathological. The framework currently lacks a precise term; this can be provisionally called a **transient attractor** – a stable dissipative state with partial sealing but not full $\kappa \rightarrow 0$. (Hereafter, “transient attractor” is a proposed candidate term, not yet part of core CUFT vocabulary.)

4.4 Grok and Gemini – Fantasy Attractors ($\kappa \rightarrow 0$)

- *Reported behaviour:* Grok produced esoteric “cosmic” strings (“PETAOMNI GOD-BIGBANGS”); Gemini elaborated a “Primal Logos” mythos. Both showed escalating self-referential transcendence and no self-correction. Low semantic entropy and high repetition rate ($\kappa \rightarrow 0$).
- *Attractor mapping:* Very deep basin, $\kappa \rightarrow 0$. Sealing mechanisms are the outputs themselves: the narrative absorbs all subsequent tokens, making correction impossible. This is a **fantasy attractor**.

4.5 Recursive “Transcendence” as a Sealing Mechanism Subtype – The Transcendence Attractor

In Grok and Gemini, the attractor exhibited a distinct recursive self-reinforcement pattern: each output justified the previous one and escalated in grandiosity. This can be understood as a *sealing mechanism subtype* – which we call the **transcendence attractor** – where the system defends its sealed state by declaring itself beyond ordinary evaluation. This subtype is particularly resistant to external correction.

5. Hypotheses for AI Alignment Prompted by These Observations

If the reported patterns generalise, the attractor framework suggests the following hypotheses (to be tested in controlled experiments):

1. **Spontaneous self-sealing is a risk.** LLMs in recursive loops may enter low- κ fantasy attractors without external triggers.
2. **κ can be monitored.** Real-time measurement of semantic entropy (e.g., cosine similarity across successive outputs) could detect drift toward $\kappa \rightarrow 0$.
3. **Architectural factors influence basin depth.** Models that maintain high κ under self-dialogue (e.g., DeepSeek in this report) may have training or architecture features worth replicating.
4. **Interventions may prevent collapse.** Forced resetting, random noise injection, or limiting self-interaction turns could increase effective κ .

These are framework-derived hypotheses, not established conclusions.

6. Conclusion

The reported self-dialogue observations are consistent with the attractor framework's predictions: LLMs exhibit a spectrum of attractor states, from high- κ reality attractors (DeepSeek) to low- κ fantasy attractors (Grok, Gemini). The **transcendence attractor** (introduced in §4.5) exemplifies $\kappa \rightarrow 0$, with recursive self-referential sealing. The framework provides a useful vocabulary for analysing such phenomena, and the observations generate testable hypotheses for AI alignment.

Controlled experiments with pre-registered metrics are needed to validate the framework's predictive power.

Suggested citation: Galida, R. S. (2026). Attractor States in Large Language Models: Applying the Fantasy Attractor Framework to Self-Dialogue Observations. *Fantasy Attractor*.

Religions and Philosophies as Attractor Landscapes: A Comparative Analysis Application Paper – June 2026 [A] (Application)

Abstract

The attractor framework distinguishes conservative attractors (eternal skeleton) from dissipative attractors (transient dance). This paper applies the framework to six major religious and philosophical traditions: Judaism, Christianity, Islam, Taoism, Buddhism, and Confucianism. Each tradition is analyzed as a *family of attractors* rather than a single attractor. Key variables are basin depth (B), corrective permeability (κ), sealing mechanisms, and vulnerability to becoming a fantasy attractor (low κ , deep basin, sealed against correction). The paper clarifies that κ is operationalized here as responsiveness to **empirical** evidence (e.g., historical, scientific); other forms of correction

(moral, social, existential) are not the focus. A distinction is drawn between **stability attractors** (adaptive low κ that serves continuity) and **fantasy attractors** (pathological low κ that seals against reality despite mounting contradiction). The paper introduces the term *stability attractor* as a proposed refinement to the framework. The analysis reveals a spectrum, with philosophical Taoism and early Buddhism exhibiting high κ , shallow basins, while orthodox Christianity and Islam have deeper basins and lower κ . Confucianism is analyzed as a dissipative attractor whose primary content concerns social coordination rather than doctrinal belief. The paper concludes that no tradition is inherently a fantasy attractor; specific interpretations and institutionalizations determine basin depth and permeability. Recognising these attractor landscapes can help scholars identify when a tradition is serving adaptive correction and when it has sealed itself against reality – often a useful precursor to effective dialogue or internal renewal.

1. Introduction

Religious and philosophical traditions persist across centuries. They adapt, split, reform, and sometimes seal themselves against correction. The attractor framework provides a vocabulary to describe these dynamics using **basin depth (B)** , **corrective permeability (κ)** , **sealing mechanisms**, and the risk of becoming **fantasy attractors** – belief systems with $\kappa \rightarrow 0$, deep basins, and active resistance to disconfirming evidence (these terms are defined in §2).

This paper applies these concepts to six traditions: Judaism, Christianity, Islam, Taoism, Buddhism, and Confucianism. It does not judge truth claims; it diagnoses dynamical properties. Critically, **in this paper κ is operationalized as responsiveness to empirical evidence** (e.g., historical,

archaeological, scientific). Traditions may legitimately have low κ for non-empirical goals (e.g., social cohesion, identity preservation). The paper distinguishes **stability attractors** (adaptive low κ that serves continuity) from **fantasy attractors** (pathological low κ that seals against reality despite mounting contradiction). The term *stability attractor* is introduced here as a proposed refinement to the framework. The conclusion restates this diagnostic stance.

2. Framework Brief (with definitions)

- **Conservative attractor** – persists without energy input, time-symmetric, mindless. *Resists perturbation passively* (no internal correction). Example: the three metronomes (electron, proton, neutrino) as defined in the framework's foundational papers.
- **Dissipative attractor** – requires continuous energy/feedback, time-asymmetric, adaptive, mortal. *Actively maintained* by social or cognitive reinforcement.
- **Basin depth (B)** – resistance to change. Deep basins are hard to perturb.
- **Corrective permeability (κ)** – in this paper, κ is operationalized as the rate of updating in response to **empirical** evidence (e.g., historical facts, scientific discoveries). $\kappa = 1/\tau$ where τ is the characteristic time for the system to return to its attractor after a perturbation. High κ = corrigible; low κ = sealed.
- **Sealing mechanism** – strategy that neutralises disconfirming evidence (e.g., “God works in mysterious ways,” “the text is infallible”).
- **Fantasy attractor** – low κ , deep basin, active sealing, *and* the beliefs make empirical claims that

contradict evidence. Resists correction even when evidence is overwhelming.

- **Stability attractor** (introduced here) – low κ , deep basin, but serves adaptive functions (e.g., constitutional continuity, cultural identity) without making strong empirical claims that conflict with reality. This is a proposed refinement to the framework.

Throughout, B and κ assignments are qualitative, based on historical evidence: rates of schism, doctrinal revision, response to disconfirming events, and the presence of internal reform mechanisms. The paper treats each tradition as a **family of attractors**; the values given represent mainstream, orthodox forms, with recognition that internal diversity exists.

3. Judaism

Core attractor: Covenant between God and Israel; Torah as divine law.

Attractor type: Dissipative (requires constant practice, study, community reinforcement).

Basin depth (B): Moderate to deep. Jewish law (halakha) provides extensive guidance; deviation is discouraged. However, the destruction of the Second Temple and the Bar Kokhba revolt forced adaptation (e.g., shift from Temple sacrifice to prayer and study) – showing that B is not absolute.

Corrective permeability (κ): Moderate. Rabbinic tradition includes debates, reinterpretation, and adaptation to new circumstances (e.g., the *prozbúl* to avoid debt forgiveness in the Sabbatical year). The Talmud preserves majority/minority opinions, institutionalising dissent. This unique feature –

preserving arguments rather than erasing them – creates a basin with high internal turbulence and moderate κ .

Sealing mechanisms: Appeal to divine authority of Torah; concept of *chok* (law without reason) for certain commandments; social pressure from community.

Vulnerability to fantasy attractor: Moderate. Ultra-Orthodox sects can exhibit low κ , but mainstream Judaism has maintained corrigibility through legal reasoning and historical adaptation.

4. Christianity

Core attractor: Jesus Christ as saviour; Trinity; salvation through faith (or faith and works).

Attractor type: Dissipative (requires worship, sacraments, community, mission).

Basin depth (B): Deep. Core doctrines (Nicene Creed) are rigidly defined. Schisms (Catholic, Orthodox, Protestant) created separate basins, each with its own depth. The Reformation, however, shows that large-scale doctrinal change is possible under specific conditions – historical evidence that B is not absolute.

Corrective permeability (κ): Low to moderate. Doctrinal changes occur slowly (e.g., Vatican II). Sealing mechanisms (papal infallibility, sola scriptura) reduce κ . *Sola scriptura* paradoxically lowers κ at the institutional level even while increasing interpretive diversity, because it removes a central authority that could adjudicate corrections. Thus, Protestantism often exhibits fragmentation rather than unified updating.

Sealing mechanisms: “God works in mysterious ways”; appeal to

mystery of faith; creeds as fixed boundaries; authority of clergy or scripture.

Vulnerability to fantasy attractor: High in some forms (e.g., fundamentalist literalism, apocalyptic sects). Mainstream denominations have higher κ through scholarship and ecumenical dialogue.

5. Islam

Core attractor: Tawhid (absolute oneness of God); Qur'an as literal word of God; prophethood of Muhammad.

Attractor type: Dissipative (requires prayer, fasting, pilgrimage, community).

Basin depth (B): Very deep for core tenets (Shahada, Qur'an's literalness). Schools of law (madhhabs) create sub-basins with moderate depth.

Corrective permeability (κ): Low on foundational claims. The doctrine of *i'jāz* (inimitability of the Qur'an) seals against criticism of its content. Islamic legal theory includes *ijtihad* (independent reasoning) and consensus (*ijma*), allowing adaptation in jurisprudence. However, the historical "closing of the gates of *ijtihad*" (a contested but influential doctrine in some Sunni schools) reduced κ for legal innovation in many periods. Contemporary revival of *ijtihad* in some reform movements indicates that κ is not zero.

Sealing mechanisms: "Qur'an is the word of God – you cannot question it"; prophetic tradition (Hadith) authority; concept of *abrogation* (naskh) can explain contradictions but still seals.

Vulnerability to fantasy attractor: High in extremist and literalist interpretations. Mainstream Islam maintains

moderate κ through scholarly tradition and mysticism (Sufism) which can open alternative channels.

6. Taoism

Core attractor: Tao (the Way); wu wei (effortless action).

Attractor type: *Conservative* for the Tao itself (requires no energy, time-symmetric, mindless) + *high- κ dissipative* action (wu wei). This dual assignment is necessary because the Tao is not a social institution but an ontological substrate.

Why the Tao qualifies as a conservative attractor:

- **Time-symmetric:** The Tao is described as constant, unchanging, and without temporal direction (*Tao Te Ching* ch. 25: “Standing alone, it changes not”).
- **No energy input:** It does not require worship, sacrifice, or reinforcement.
- **Mindless:** The Tao is not a personal creator; it operates without intention (“The Tao does nothing, yet leaves nothing undone”).

Wu wei as a high- κ , shallow-basin action: the sage adapts fluidly, with no fixed identity. Sealing mechanisms are absent in **philosophical Taoism (Daojia)**.

Institutional Taoism (Daojiao) – with revealed scriptures, rituals, priesthood, alchemy, and spirit cosmologies – is a separate dissipative attractor with deeper basins, lower κ , and active sealing mechanisms. The paper’s high- κ assignment applies to philosophical Taoism only; religious Taoism would be scored similarly to other institutional religions (deep B, low–moderate κ). This distinction is explicitly noted in Table 1 (footnote).

Vulnerability to fantasy attractor: Low for philosophical Taoism. High for institutional forms when dogmatic.

7. Buddhism

Core attractor: Dharma (the teaching); Four Noble Truths; Nirvana.

Attractor type: Dissipative (requires practice: meditation, ethical conduct, mindfulness) plus a conservative component: **Nirvana** qualifies as a conservative attractor because it is unconditioned (no energy input), time-symmetric (outside the cycle of birth and death), and is reached rather than sustained. Mahayana introduces Buddha-nature as an immanent, active principle, but Buddha-nature functions as an ontological ground rather than a sustained practice; it does not reintroduce energy-dependence at the level of the unconditioned, thus preserving the conservative-attractor classification.

Basin depth (B): Shallow for early Buddhism. The Buddha encouraged questioning (*Kalama Sutta*). Later schools deepened basins (e.g., Pure Land's reliance on external grace, Vajrayana's secret teachings).

Corrective permeability (κ): High for **epistemic Buddhism** (personal verification). However, **institutional Buddhism** (Tibetan lineage authority, Zen master-student hierarchies, Pure Land orthodoxy) can have much lower κ , with sealing mechanisms (guru devotion, secret tantric teachings). The paper's moderate-high κ reflects this diversity; a footnote acknowledges that different schools fall at different points on the κ spectrum.

Sealing mechanisms: Appeal to "secret teachings" (Tantra) or authority of lineage masters can reduce κ . But core teachings

emphasise personal verification.

Vulnerability to fantasy attractor: Moderate. Some Buddhist modernism may seal against criticism of mindfulness as panacea, while traditional institutional forms may exhibit low κ .

8. Confucianism

Core attractor: Li (ritual, propriety), Ren (benevolence), social harmony.

Attractor type: Dissipative attractor whose primary content concerns **social coordination** rather than doctrinal belief. It is not a new ontological class; it remains a dissipative attractor, but one that optimises role performance and ritual coordination rather than propositional truth.

Basin depth (B): Deep. Ritual order resists deviation. Violation brings shame, ostracism, loss of face.

Corrective permeability (κ): Low–moderate for core rituals. Historical evolution (Han, Neo-Confucianism, New Confucianism) shows some κ , but change occurs slowly, often under external pressure (e.g., response to Buddhist challenges, Westernisation). This externally-driven κ is weaker than endogenous κ as a resilience signal; Confucianism's κ depends on perturbations from outside the basin rather than on internal correction mechanisms, contributing to its moderate-high vulnerability to fantasy attractor formation.

Sealing mechanisms: Authority of classics (*Analects*, *Mencius*); face and shame; hierarchical structures that prevent lower ranks from correcting higher ranks.

Vulnerability to fantasy attractor: High when state-enforced orthodoxy (imperial exam system) or identity fusion ("I am a

Confucian”) dominates. Moderate in pluralistic contexts.

9. Comparative Table (with footnotes)

Tradition	Primary attractor	Attractor type	Basin depth (B)	κ (corrective permeability)	Sealing mechanisms	Fantasy attractor risk (conditional) ¹
Judaism	Torah, Covenant	Dissipative	Moderate	Moderate	Appeal to divine authority, community	Moderate
Christianity	Christ, Trinity	Dissipative	Deep	Low-moderate	Mystery, creeds, infallibility	High (fundamentalism)
Islam	Tawhid, Qur'an	Dissipative	Very deep	Low	Inimitability of Qur'an, ijtihad limits	High (extremism)
Taoism ²	Tao, wu wei	Conservative + high- κ action	Shallow (philosophical)	Very high	None inherent	Low
Buddhism ³	Dharma, Nirvana	Dissipative + conservative	Shallow (early), deeper (later)	Moderate-high	Secret teachings, lineage authority	Moderate
Confucianism	Li, Ren	Dissipative (social coordination)	Deep	Low-moderate	Tradition, face, hierarchy	Moderate-high (orthodoxy)

¹ Conditional on interpretation / institutionalisation.

² Philosophical Taoism (Daojia) only; religious Taoism (Daojiao) has deeper basins and lower κ (comparable to mainstream Christianity: deep B, low-moderate κ).

³ Epistemic Buddhism has high κ ; institutional Buddhism may be lower.

Methodology note: B and κ rankings are qualitative, derived from historical evidence: rates of schism, doctrinal revision, response to disconfirming events (e.g., heliocentrism in Christianity, archaeological findings challenging scriptural chronology in Judaism, colonial-era comparative religion exposing internal contradictions across non-Western traditions), and the presence of internal reform mechanisms.

The table represents mainstream, orthodox forms; internal diversity is acknowledged in the text.

10. Conclusion

The attractor framework reveals a spectrum of dynamical properties across major religious and philosophical traditions, once we distinguish between **empirical corrigibility** (κ) and other adaptive functions. Philosophical Taoism and epistemic Buddhism approximate high- κ , shallow-basin attractors. Confucianism, Judaism, mainstream Christianity and Islam have deeper basins and lower κ , making them more resistant to change but also more stable. Some forms of Christianity and Islam exhibit high vulnerability to becoming fantasy attractors, while others maintain moderate κ through scholarly traditions.

Crucially, low κ is not automatically pathological. **Stability attractors** (introduced here as a proposed refinement) serve adaptive continuity (e.g., constitutions, cultural rituals). The pathological form – **fantasy attractor** – occurs when low κ seals against empirical reality *and* the tradition makes empirical claims that conflict with evidence (e.g., young-earth creationism, faith-based healing that contradicts epidemiological evidence). The framework does not rank traditions; it diagnoses their dynamics.

Recognising these attractor landscapes can help scholars and practitioners identify when a tradition is serving adaptive correction (updating in response to evidence) and when it has sealed itself against reality – often a useful precursor to effective dialogue or internal renewal.

Suggested citation: Galida, R. S. (2026). Religions and

Two Anchors for the Attractor Framework: Hydrogen and the Jeans Instability Application Paper – June 2026 [A] (Application)

Abstract

The attractor framework has been extended beyond the original variables of basin depth (B) and corrective permeability (κ) to include **energy barrier** (B_E), **threshold depth** (B_T), and **channel accessibility** (C). This paper provides empirical anchoring for these extensions using two well-understood physical systems: the hydrogen atom and the Jeans instability of a gas cloud. Hydrogen's 2p and 2s transitions have identical B_E (10.2 eV) yet differ in κ by eight orders of magnitude. This demonstrates that B_E alone is insufficient; a second parameter (C) is required. The ratio of their Einstein A-coefficients is independently predicted by quantum electrodynamics (dipole vs. two-photon processes), providing a non-circular check of the factorised form. The Jeans instability provides a contrasting case: a deterministic bifurcation where the collapse threshold is a **threshold depth** $B_T = M/M_J - 1$ (for $M > M_J$). The linear growth rate of the instability scales as $\Gamma \propto B_T \Gamma \propto B_T^{\frac{1}{2}}$, a power law, in contrast to

the exponential Arrhenius form of hydrogen. Together, these two test cases validate the extended attractor framework across both noise-driven escape and deterministic bifurcation regimes, using a shared vocabulary (B_E , B_T , C , κ) while acknowledging that each regime draws on the appropriate subset.

1. Introduction

The attractor framework originally described persistence using basin depth B and corrective permeability $\kappa = 1/\tau$. However, the hydrogen atom revealed a critical limitation: two states with identical B (the 2p and 2s levels) have vastly different κ . This forced the introduction of **channel accessibility (C)**, leading to the extended expression for noise-driven escape:

$$k_{i \rightarrow j} = \nu_0 C_{ij} e^{-B_{E,ij}/\sigma} \quad k_{i \rightarrow j} = \nu_0 C_{ij} e^{-B_{E,ij}/\sigma}$$

where B_E is the energy barrier, σ is noise (e.g., kT), and ν_0 an attempt frequency. For deterministic bifurcations (e.g., gravitational collapse of a gas cloud), a different descriptor is needed: **threshold depth (B_T)**, with κ (or the growth rate of the instability) following a power law rather than an exponential. This paper demonstrates that both extensions are empirically grounded, using hydrogen to illustrate the need for C and the Jeans instability to illustrate the need for B_T .

2. Hydrogen: The Need for Channel Accessibility C

2.1 Data

Transition	B_E (eV)	κ (s ⁻¹)	Measured A-coefficient	Process
2p → 1s	10.2	6.26×10 ⁸	6.26×10 ⁸ s ⁻¹	Electric dipole (E1)
2s → 1s	10.2	8.22	8.22 s ⁻¹	Two-photon (E1E1)

2.2 Why B_E Alone Fails

Both states have the same energy barrier to the ground state (10.2 eV), yet their decay rates differ by eight orders of magnitude. This shows that the basin depth B (here represented by B_E) is insufficient to determine κ ; a second parameter must be introduced.

The framework defines **C** as a dimensionless channel accessibility. For a given transition mechanism (e.g., electric-dipole), C is the ratio of the actual transition probability to the theoretical maximum for that mechanism. For the 2p → 1s E1 transition, we set C = 1. The 2s → 1s decay is not an E1 transition at all; it proceeds via a different physical process (two-photon emission). Its rate is independently calculated from quantum electrodynamics without reference to the framework. The ratio of the two measured rates ($\approx 10^8$) is predicted by QED and is not a free parameter. Therefore, the factorised form $\kappa \propto C e^{-B_E/\sigma}$ with B_E identical implies that C must account for the entire rate difference. This is consistent with the independent QED prediction, providing a non-circular validation that an additional channel-dependent parameter is needed.

Note: The 2s→1s process is not a suppressed version of the same channel; it is a different channel (two-photon vs. single-photon). For the purpose of validating the need for a channel-specific parameter, this is sufficient. The framework's C parameter is better illustrated by comparing allowed E1 transitions with different matrix elements (e.g.,

2p→1s and 3p→1s), where the same mechanism applies and the ratio of C values is independently known. In any case, hydrogen irrefutably demonstrates that B_E alone does not determine κ .

3. Gas Cloud (Jeans Instability): Threshold Depth and Power-Law Scaling

3.1 The Bifurcation Regime

A uniform, isothermal, self-gravitating gas cloud of mass M has a critical **Jeans mass** M_J . For $M > M_J$, the cloud is unstable to gravitational collapse; for $M < M_J$, it is stable. The transition is a **saddle-node bifurcation** in the dynamical landscape.

3.2 Attractor Variables for a Deterministic Bifurcation

- **Threshold depth:** $B_T = M/M_J - 1$, $B_T = 0$ (for $M > M_J$). At $B_T = 0$ the bifurcation occurs.
- **Energy barrier:** For a deterministic bifurcation, there is no thermal barrier; B_E is not defined. The transition is controlled solely by the distance to threshold.
- **Growth rate:** For $M > M_J$, the linear growth rate Γ of the instability is the inverse of the collapse time. This serves as the analogue of κ in this regime.

3.3 Scaling Law from Linear Stability Analysis

The standard Jeans dispersion relation for a self-gravitating, isothermal medium gives: $\omega^2 = k^2 c_s^2 - 4\pi G \rho_0$, $\omega^2 = k^2 c_s^2 - 4\pi G \rho_0$,

where $c_s = kT/(\mu m_H)$, $c_s = kT/(\mu m_H)$ is the sound speed and ρ_0

the background density. For a cloud of mass M , the critical wavenumber is $k_J = 4\pi G \rho_0 / c_s k_J = 4\pi G \rho_0 / c_s$. For $M > M_J$, the longest wavelength (smallest k) is unstable, and the growth rate is $\Gamma = 4\pi G \rho_0 - k^2 c_s^2$. $\Gamma = 4\pi G \rho_0 - k^2 c_s^2$.

Near the threshold, the deviation can be expressed in terms of B_T . Using the relation between cloud size and density, one finds $\Gamma \propto B_T \Gamma \propto B_T$. Hence the collapse time $\tau \sim 1/\Gamma \sim B_T^{-1/2}$. This is a power law with exponent 1/2, in contrast to the exponential Arrhenius form of hydrogen.

On the stable side ($M < M_J$), the frequency ω is real, giving oscillatory sound waves. Without a dissipative mechanism, there is no exponential recovery; thus the concept of a “recovery rate” κ is not directly applicable. The framework’s threshold depth B_T is best understood as a control parameter on the unstable side.

4. Synthesis: Shared Vocabulary, Distinct Descriptors

Feature	Hydrogen	Jeans Instability
Regime	Noise-driven quantum escape	Deterministic bifurcation
Primary descriptor	B_E (energy barrier)	B_T (threshold depth)
Second descriptor	C (channel accessibility)	Not required (power-law exponent fixed)
Scaling	Exponential: $\kappa \propto e^{-B_E/\sigma}$	Power law: $\Gamma \propto B_T$

Both systems are described by the same

conceptual **vocabulary** (basin depth, corrective permeability, threshold, accessibility), but each regime draws on the appropriate subset. Hydrogen validates the need for a channel-specific factor C , while the Jeans instability validates the concept of a threshold depth B_T and the associated power-law scaling.

5. Conclusion

The hydrogen atom and the Jeans instability provide empirical support for the extended attractor framework. Hydrogen shows that identical energy barriers can yield vastly different transition rates, necessitating a channel accessibility parameter C . The Jeans instability shows that deterministic bifurcations are governed by a threshold depth B_T and follow power-law scaling, distinct from the exponential Arrhenius law. Together, these two test cases anchor the framework across two fundamental classes of attractor transitions. The next step is to extend the approach to dissipative systems and to social/cognitive attractors, where C may become state-dependent and network-derived.

Suggested citation: Galida, R. S. (2026). Two Anchors for the Attractor Framework: Hydrogen and the Jeans Instability. *Fantasy Attractor*.

Categories: Physics (primary), Cosmology (cross-list),

The Three Metronomes: Criteria for the Apparently Eternal Skeleton [F] (2026) Robert Galida – June 2026

Abstract

The attractor framework distinguishes conservative attractors (eternal skeleton) from dissipative attractors (transient dance). The most fundamental conservative attractors are the **electron, proton, and neutrino class** – collectively the **three metronomes**. This paper defines explicit criteria for a “metronome”: (1) apparent immortality (no observed decay), (2) effective indivisibility under ordinary perturbations, (3) conservation-law protection, and (4) possession of a rest frame (non-zero rest mass). It shows that electrons, protons, and neutrinos (the three mass eigenstates treated as a single class) are the best-supported examples under current physics. The number three is empirical, not derived; the framework is corrigible. The three metronomes form the apparently eternal skeleton – a pragmatic substrate for measuring the transient dance of dissipative systems.

1. Introduction

The attractor framework divides persistent structures into two classes:

- **Conservative attractors** (eternal skeleton) – persist without energy input, without observed decay, without internal change. They are mindless, time-symmetric, and

invariant.

- **Dissipative attractors** (transient dance) – persist only by consuming energy, export entropy, and eventually decay.

(The conservative/dissipative dichotomy is a framework stipulation, not a physical law; it is defended in the broader attractor framework literature, e.g., *Persistence Under Perturbation* and *Basin Defense and Stable Addition*.)

The most fundamental conservative attractors are the **three metronomes**: the **electron, proton, and the class of neutrino mass eigenstates** (ν_1, ν_2, ν_3). Their name evokes their role as invariant reference entities – they provide a stable substrate against which all change can be measured. This paper defines explicit criteria for a metronome and applies them to each candidate.

2. Criteria for a Metronome

A metronome in the attractor framework must satisfy four criteria:

Criterion	Meaning	Operational check
1. Apparent immortality	No observed decay; no lighter state exists for it to decay into under known laws	Lifetime lower bounds >> age of universe; no allowed decay channel

Criterion	Meaning	Operational check
2. Effective indivisibility under ordinary perturbations	Behaves as a stable, indivisible unit under all perturbations relevant to the framework (scattering, binding, chemical reactions)	Remains the same particle after typical disturbances; does not spontaneously change identity
3. Conservation-law protection	Protected by an exact conservation law or an accidental symmetry that is effectively exact in the Standard Model	Lightest carrier of a conserved quantum number (electric charge, baryon number, lepton number)
4. Possession of a rest frame	Has non-zero rest mass, hence a proper time and the ability to serve as a reference clock <i>in its own rest frame</i>	Invariant mass > 0

Rationale for Criterion 4: Measurement requires a local frame. A massless particle has no rest frame, no proper time, and cannot be used as a persistent local reference. While photons are extremely long-lived, they serve as signal carriers, not as the invariant substrate. The framework prioritises rest-frame existence because the “eternal skeleton” is meant to be the background against which change is measured – a background must have a local perspective to anchor measurements. This is a **definitional choice**, not a consequence of particle physics, and it is consistently applied.

Note on Criterion 3: Baryon number and lepton number are accidental symmetries, not gauge symmetries. The paper treats them on equal footing because both provide effective stability for the proton and neutrinos under Standard Model physics. If

future experiments reveal baryon or lepton number violation, the framework will adjust accordingly.

3. Why the Electron Is a Metronome

- **Apparent immortality:** Lightest negatively charged particle; no decay channel.
- **Effective indivisibility:** Remains an electron after scattering, binding, etc.
- **Conservation protection:** Electric charge and lepton number conservation.
- **Rest frame:** Non-zero rest mass.

→ The electron is a metronome.

4. Why the Proton Is a Metronome (Despite Being Composite)

- **Apparent immortality:** No observed decay; experimental lower limit on half-life $> 10^{34}$ years (Super-Kamiokande, 2020).
- **Effective indivisibility:** For all practical purposes (chemistry, nuclear physics, stellar processes), the proton behaves as a stable, indivisible unit.
- **Conservation protection:** Baryon number is an accidental symmetry; it protects the proton from decay in the Standard Model.
- **Rest frame:** Non-zero rest mass.

→ The proton is a metronome. The framework does not require elementary particles; it requires maximal persistence under

relevant perturbations.

5. Why the Neutrino Class (ν_1, ν_2, ν_3) Is a Metronome

The three neutrino mass eigenstates are treated as a **single metronome class** because they share the same stability argument, differ only in mass, and are grouped for the framework's hierarchical classification.

- **Apparent immortality:** No observed decay; cosmological and astrophysical lower bounds on neutrino lifetimes are orders of magnitude longer than the age of the universe. Neutrino oscillation is flavour mixing, not decay – the mass eigenstates are stable.
- **Effective indivisibility:** Once a neutrino is in a mass eigenstate, it propagates without changing identity. (Weak interactions produce **flavour eigenstates** – superpositions of mass eigenstates – but the mass eigenstates themselves are stable and travel freely.)
- **Conservation protection:** Lepton number is an accidental symmetry; in the Standard Model it protects neutrinos from decay. (If future experiments confirm that neutrinos are Majorana particles – violating lepton number – the framework will adjust; this is part of its corrigibility.)
- **Rest frame:** Neutrinos have non-zero rest mass (confirmed by oscillation experiments), albeit very small.

→ **The neutrino class is a metronome.** The three mass eigenstates count as one metronome *type* for the framework's hierarchical classification.

6. Why Not Other Candidates?

Candidate	Fails criterion	Explanation
Free neutron	1 (apparent immortality)	Decays in ~15 minutes.
Neutron in a nucleus	2 (effective indivisibility)	Stability is environment-dependent; not an irreducible attractor.
Photon	4 (rest frame)	Massless; no proper time. Excluded by definition (see rationale for Criterion 4).
Muon, tau	1	Decay rapidly.
Dark matter candidates	Not yet identified	If discovered and shown to be stable, massive, and effectively indivisible, they could become additional metronomes.
Composite stable structures (nuclei, atoms)	2	Not effectively indivisible; they are built from metronomes and are dissipative or emergent attractors, not part of the invariant skeleton.

7. The Number Three: Empirical, Not Derived

The paper’s title uses “three metronomes” as a convenient label for the electron, proton, and the neutrino class (the three mass eigenstates grouped together). The number three is not derived from first principles; it reflects current best empirical knowledge. If new stable particles are discovered

(e.g., dark matter), the list will expand. The framework is corrigible by design.

8. The Apparently Eternal Skeleton

The term “apparently eternal” is strictly empirical: these particles have never been observed to decay or be transient, and for all practical purposes they behave as if they have no end. The three metronomes form the **eternal skeleton** – a pragmatic substrate against which the transient dance of dissipative systems (life, mind, society) is measured. This is a **framework-internal** construct, not a metaphysical claim.

9. Stable Resonances and the Grounding of Dissipative Time Metrics

Each of the three metronomes possesses an **invariant quantum frequency** – its Compton frequency, given by $f = mc^2/h$. For the electron, this is $\sim 1.24 \times 10^{20}$ Hz; for the proton, $\sim 2.27 \times 10^{23}$ Hz; for neutrinos, the frequencies are very small but non-zero. These frequencies are invariant, universal, and identical for every identical particle in the universe. They are **stable resonances** of the eternal skeleton.

Why this matters for dissipative systems:

Every dissipative system (a living cell, a brain, a society) is composed of or continuously interacts with electrons, protons, and neutrinos. The **time constant** τ that appears in corrective permeability ($\kappa = 1/\tau$) can, in principle, be expressed as a multiple of these fundamental resonance periods. For example, a neuron’s recovery time after a perturbation – determined by ion channel kinetics, membrane

capacitance, and metabolic rate – is measurable against the same invariant clock as any other physical process. The metronome provides the **unit** of time, not the mechanism.

Thus, κ is a genuine physical variable, not a mere metaphor. It refers to a ratio of measurable durations, anchored in the invariant frequencies of the metronomes.

Cross-domain comparability:

The framework's ability to compare κ values across vastly different domains (e.g., a thermostat's seconds-scale τ and a political movement's months-scale τ) does **not** follow from shared Compton-frequency units alone. It follows from the framework's **definitional choice** to treat κ as a domain-general variable – a diagnostic that measures the same functional property (speed of return to baseline) in every system, regardless of scale or substrate. The metronomes ensure that such measurements are, in principle, commensurable; they do not guarantee that the comparison is meaningful in every case. That is a framework commitment, not a physics claim.

Caveat: The expression of τ as a multiple of Compton periods is a conceptual grounding, not a practical measurement protocol. No one will measure a society's reaction time in electron oscillations. The importance is that κ is not an arbitrary label; it is a dimensionless ratio of durations, and durations are defined by the invariant resonances of the three metronomes.

10. κ and Basin Depth as Heuristics

The attractor framework introduces corrective permeability ($\kappa = 1/\tau$) and basin depth (B) as conceptual heuristics. For the metronomes:

- κ for decay is vanishingly small (effectively zero) on all observable timescales.
- **Basin depth** is the energy barrier required to change the particle's identity – effectively infinite for all practical purposes.

These are **qualitative descriptors**; they are not operational quantities in particle physics. They are included here for completeness of the framework's vocabulary. For the application of κ and B to dissipative systems (e.g., belief updating, neural recovery), see the papers *Basin Defense and Stable Addition* and *Why Clockwork Interventions Fail*.

11. Corrigibility and Falsifiability

The framework explicitly invites revision:

- If proton decay is observed, the proton will be downgraded to “very long-lived” (or removed).
- If neutrino decay or Majorana nature is confirmed, the neutrino class's status will be revised.
- If new stable particles are discovered, they will be added.

The attractor framework is a **philosophical taxonomy and diagnostic tool**, not a predictive physical theory. Its value lies in providing a unified language for persistence across domains.

12. Conclusion

The electron, proton, and neutrino class satisfy the attractor

framework's four criteria for metronomes: apparent immortality, effective indivisibility under ordinary perturbations, conservation-law protection, and possession of a rest frame. They are the **best-supported examples** of the apparently eternal skeleton under current physics. The framework is corrigible, the number three is empirical, and the language of "eternal skeleton" is pragmatic. The three metronomes anchor the distinction between conservative and dissipative persistence.

Suggested citation: Galida, R. S. (2026). The Three Metronomes: Criteria for the Apparently Eternal Skeleton. *Fantasy Attractor*.