

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)
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
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
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., Askell, 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*.

Whirling as Attractor Engineering: Chirality, Shared Resonance, and a Minimal-Dose Protocol for Whole-Body Resilience

Author: Robert Galida

Date: May 2026 (Revised June 2026)

□ **Note (June 2026):** This paper's description of conservative attractors has been updated to reflect the refined framework in *Metronome, Memory, and the Threefold Anchor: A Relational Account of Time* [F] (2026). The health and self-engineering content is unchanged.

Abstract

Whirling – the spinning practice of Mevlevi dervishes – is often seen as a mystical ritual. This paper reinterprets it through the attractor framework, where the mind is a dissipative attractor of the whole body.

Whirling is a controlled, repeated perturbation. It trains your balance, nervous system, and heart to settle into a stable, coherent pattern – a form of attractor engineering.

We discuss two additional ideas:

- **Chirality alignment** – spinning counter-clockwise may symbolically align with the universe's handedness (e.g., left-handed neutrinos), but this is speculative and not needed for health benefits.
- **Shared resonance** – group whirling synchronises heartbeats, creating a collective attractor.

We review scientific evidence showing that whirling improves heart rate variability (HRV), sleep quality, anxiety, brain plasticity, and physical fitness. A minimal effective dose is 5–15 minutes per day, 3–4 times per week. A graduated protocol is provided.

The health benefits are real. The chirality interpretation is optional.

1. Introduction

In the attractor framework, your mind is a dissipative attractor of your whole body – a pattern that needs energy flow to stay stable, can be disturbed, and can adapt. Self-engineering means using small, repeated disturbances to

reshape your own attractor towards greater resilience.

Whirling is a sustained, counter-clockwise spin performed by Mevlevi dervishes for centuries. It is spiritual, but modern science has found clear physical and mental benefits.

This paper argues that whirling is a powerful attractor engineering practice: a rhythmic whole-body disturbance that forces your system to become more stable and coherent. We also explore two extra ideas:

- **Chirality** (spinning with the universe's "handedness" – speculative)
 - **Shared resonance** (heartbeat synchronisation in groups – well supported).
-

2. The Attractor Framework Primer (Very Brief)

- **Conservative attractors** are eternal, time-symmetric, and require no energy input. They form the *eternal skeleton*. The three most fundamental conservative attractors – the *metronomes* – are the **electron**, **neutrino mass eigenstates** (collectively), and **proton**. (The photon is a signal carrier, not a metronome; see *Metronome, Memory, and the Threefold Anchor* for details.)
- **Dissipative attractors** (life, mind, society) need energy flow, have finite lifetimes, and can change. Your body is a stack of dissipative attractors.
- **Persistence under disturbance** is the basic mark of reality. A resilient system returns to its attractor after a knock.
- **Self-engineering** uses small, repeated nudges to reshape your own attractor basin.

- **Whirling** is a strong, repeated disturbance. Your body must adapt. That adaptation is the engineering.
-

3. Chirality Alignment – A Speculative Interpretation

3.1 What do we know about universal handedness?

- **Weak interactions:** Neutrinos produced in weak decays are always left-handed (Wu experiment, 1956). This is a fact. But electrons and protons do not have a universal spin direction.
- **Astronomical rotations:** From the north pole, Earth, the solar system, and the Milky Way rotate counter-clockwise. From the south pole, they appear clockwise. That's just a viewpoint – there is no privileged direction in space.
- **Cosmic Microwave Background:** Some studies suggested a preferred axis ("axis of evil"), but these results are contested and likely statistical artifacts. No clear evidence.

3.2 The speculative claim

The dervish's counter-clockwise spin can be seen as a heuristic alignment with these physical handednesses (neutrino helicity, frame-dependent rotation). In our attractor framework, we propose that spinning with the majority direction (as seen from the northern hemisphere) may resonate symbolically and phenomenologically with the invariant rhythms of the conservative substrate – the three metronomes.

Crucially, there is no known physical mechanism linking a rotating body (~1–2 rpm) to particle spin or photon

polarisation. The scale difference is huge. So this alignment is presented as a speculative metaphysical claim within our framework, not as proven physics. It's a way to frame the practice, not a testable hypothesis. The health benefits of whirling do not depend on this speculation.

3.3 Clockwise vs. counter-clockwise

No study has compared clockwise and counter-clockwise whirling for health effects. The idea that clockwise spinning “needs more energy” or “opposes the Tao” is unsupported – we label it as speculation. You can try both directions, but the traditional counter-clockwise spin is recommended for alignment with our framework's interpretive preferences.

4. Shared Resonance: Heartbeat Synchronisation

A published study measured heart rates during a group Sufi whirling ritual. It found that participants' heartbeats became synchronised – the biological data matched the spiritual goal of unity.

In attractor terms: the shared rhythm creates a common basin of attraction across people. Each body locks onto the same external rhythm (the group spin), and through mutual coupling, their cardiac oscillators fall into step.

This is like metronomes placed on a movable platform – they eventually synchronise (a classic demonstration from Huygens, 1665). Here, the “platform” is the shared sound and feel of the group whirling. The result is a collective attractor – a stable shared state where heart rates align, possibly amplifying resilience.

Note: The term “collective attractor” simply means a stable

pattern in a coupled system. The 2019 study showed cardiac synchronisation, but the idea that whirling together increases resilience beyond what you can do alone is still a plausible hypothesis that needs testing.

5. Evidence for Health Benefits

5.1 Heart Rate Variability (Autonomic Resilience)

A 2012 study on “Whirling-Kung” (5–15 minutes, three times per week) found the practice prevented a decline in key HRV measures (SDNN, total power) seen in a control group. Higher HRV means a wider attractor basin, faster recovery, and greater resilience.

5.2 Sleep Quality and Stress Markers

A 2022 study on whirling dervishes found significantly better sleep quality and much lower anxiety ($p < 0.001$) compared to non-whirling controls. The dervishes also had lower levels of VEGF, BDNF, and GDNF – markers often elevated by chronic stress.

Note on BDNF: Lower BDNF is usually linked to depression, not less stress. The authors of the study interpreted this as a possible protective effect, but the relationship is complex. We simply report the finding without endorsing a specific interpretation.

5.3 Neuroplasticity – Reshaping the Brain’s Attractor Landscape

An MRI study found that long-term dervishes have cortical thinning in the default mode network (DMN) and motion-perception areas (right DLPFC, lingual gyrus, visual area V5). This thinning is experience-dependent neuroplasticity: the brain prunes inefficient connections to

become more specialised.

5.4 Physical Fitness and VO_2max

A 12-week whirling training programme improved body composition, leg strength, flexibility, grip strength, and both anaerobic and aerobic power (VO_2max). Whirling is effective whole-body cardiovascular exercise.

5.5 Mental Health – Less Anxiety, Better Self-Regulation

Multiple studies confirm lower anxiety. Participants report better mind-body focus, self-regulation, positive feelings, and a “quietness in the centre of the vortex” – the subjective experience of a stable core attractor.

Finding the original studies: The papers cited here (2012 HRV, 2022 sleep/anxiety, MRI, 12-week fitness, and the 2019 heartbeat study) can be found by searching terms like “whirling dervish heart rate variability,” “whirling kung HRV,” “Dursun whirling MRI,” “Karakaya whirling sleep,” or “Genc whirling VO_2max .”

6. The Minimal Effective Dose

Based on the 2012 study and traditional practice:

- 5–15 minutes per session
- 3–4 times per week
- Counter-clockwise rotation (traditional; clockwise not harmful but lacks evidence)
- Gradual progression

Phase	Duration	Frequency	Goal
Adaptation (weeks 1–2)	5 min	3–4x/week	Get used to the spin
Consolidation (weeks 3–4)	10–15 min	3–4x/week	Find the rhythm, notice calm
Expansion (week 5+)	20–30 min	3–4x/week	Explore deeper states

7. Practical Instructions

- **Space:** A large, empty room. Bare feet.
 - **Posture:** Start with arms crossed on your chest. Begin turning counter-clockwise. After a few revolutions, open your arms: right hand up (palm to sky), left hand down (palm to earth).
 - **Gaze:** Soft, unfocused – don't fixate on a single point.
 - **Safety:** Stop if you feel severe nausea. Use a wall for support if needed.
 - **Afterward:** Rest lying down for 5–10 minutes to let your balance system settle.
-

8. Conclusion

Whirling produces real, measurable benefits: better HRV, sleep, anxiety, brain plasticity, and fitness. A minimal dose of 5–15 minutes a day, three to four times a week, is enough.

The shared resonance (heartbeat synchronisation in groups) is empirically supported.

The chirality alignment (spinning counter-clockwise to align with the universe) is a speculative interpretation – not

required for the health benefits.

The dervish's spin is a dance of persistence under perturbation – a transient dancer humming along with the eternal skeleton. The dance has a new step.

Suggested citation: Galida, R. S. (2026). Whirling as Attractor Engineering: Chirality, Shared Resonance, and a Minimal-Dose Protocol for Whole-Body Resilience (Revised June 2026). *Fantasy Attractor*.