

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 Philosophies as Attractor Landscapes: A Comparative Analysis (Final). *Fantasy Attractor*.

Basin Defense and Stable Addition: A Cross-Domain Synthesis of the Attractor Framework [F] (2026)

Robert Galida – June 2026 (Final)

See Paper 1 ([Intelligence Without Consciousness](#)) for the full taxonomy of attractors, κ , and basin depth.

Abstract

Many complex systems resist change by returning to a preferred low-energy attractor rather than adopting a new state. Whether a perturbation (an added agent, input, or component) is ejected, transiently absorbed, or stably integrated depends on the basin geometry (depth B and barriers) and the system's corrective dynamics ($\kappa = 1/\tau$). This paper defines B and κ , draws on formal models (stochastic dynamical systems and Kramers escape theory) with explicit qualifications for non-gradient domains, and catalogs exemplar systems across ten domains. A comparative table summarizes systems, mechanisms, proxies for B and κ , timescales, and conditions favoring each outcome. The paper concludes that the same basic physics analog applies across domains: a perturbation of size Δ will be ejected or die out if Δ is below the attractor's effective escape threshold (a function of B), whereas if Δ exceeds that threshold and the system has enough plasticity or additional degrees of freedom, a new stable state can form. A research

roadmap is provided in an appendix.

1. Introduction

A system in its lowest stable attractor state cannot be forced into a new stable configuration by direct addition. Adding to the system – a third star, an extra electron, a new species, a contradictory belief – will result in one of three outcomes:

1. **Ejection** – the addition is expelled from the system entirely. The original attractor persists.
2. **Transient absorption** – the addition remains present, but the system state returns to the original attractor despite the addition's continued presence.
3. **Stable addition** – the addition is integrated, either by expanding the capacity of the original attractor or by forming a new parallel attractor alongside it.

This paper identifies a unified principle – **basin defense** – that governs these outcomes across physical, biological, ecological, social, and engineered systems. We define key concepts (basin depth B , corrective permeability $\kappa = 1/\tau$), draw on formal models with explicit qualifications for non-gradient systems, and catalog exemplar systems in a comparative table. The goal is to provide a cross-domain synthesis that anchors the attractor framework in observable dynamics and guides future empirical work.

2. Definitions and Formal Models (with Qualifications)

Attractor, Basin, and Low-Energy Attractor: In dynamical

systems, an attractor is a set of states toward which trajectories converge. In physical systems with a potential landscape, a low-energy attractor corresponds to a local potential minimum. Its basin of attraction is the region of state space that flows into the attractor. **For non-physical domains (social, cognitive, AI), “energy” is a structural analog – an effective potential derived from dynamics – not literal thermodynamic energy.** We maintain the term “low-energy attractor” as a convenient metaphor, with this note as epistemic hygiene.

Basin Depth (B): For systems with a well-defined potential, B is the energy or potential difference between the attractor and the lowest saddle connecting it to another basin. For non-gradient or high-dimensional systems, B is a **structural analog** – the effective barrier strength inferred from perturbation-response experiments (e.g., the perturbation magnitude required to shift the system to a different state). **Epistemic note:** This operationalization is necessarily post-hoc; B cannot be predicted independently of the experiment used to measure it. This circularity is an open operationalization problem, flagged as such.

Corrective Permeability (κ) and Relaxation Time (τ): We define $\kappa = 1/\tau$, where τ is the characteristic time for return to baseline after a small perturbation. **This definition is applied consistently across all domains**, with τ operationalized domain-specifically as the measured return time (e.g., seconds for a thermostat, hours for synaptic scaling, days for immune response, months for belief updating). A large κ (small τ) means fast return; a small κ means slow or absent return.

Three Outcomes Defined Operationally:

- **Ejection:** The addition leaves the system entirely. The system state returns to the attractor, and the added

entity is no longer present.

- **Transient Absorption:** The addition remains present, but the system state returns to the attractor despite the addition's continued presence.
- **Stable Addition:** The addition is integrated, and the system settles into a new attractor (expanded capacity or parallel attractor). This is the only case where the original attractor is displaced.

Formal Models (Qualified): In a one-dimensional overdamped potential, Kramers' escape theory gives mean escape time $\propto \exp(B/D)$, where D is noise intensity. **This result does not generalize to multi-dimensional, non-gradient, or non-equilibrium systems – all of which appear in our domain examples (neural networks, social systems, ecological systems).** For those systems, B and κ are **structural analogs** – quantities that play the same functional role (resistance to change; speed of return) but are not derived from a literal potential. The formal section is an analogy and a source of heuristics, not a universal physical law. We do not claim to “survey” Kramers theory; we draw on it as a conceptual anchor.

3. Minimal Physical Examples

Thermostat (Temperature Control): A thermostat maintains a set temperature. An external heat input is an addition. The thermostat's negative feedback loop turns on cooling, expelling the heat (ejection). τ is the temperature relaxation time (seconds). B is the maximum heat load before setpoint failure (Watts or $^{\circ}\text{C}$ above setpoint).

RC Circuit (Passive Decay): A capacitor discharging through a resistor has a single equilibrium at zero voltage. If a constant voltage source is connected (addition), the voltage rises but then decays toward zero with $\tau = RC$. The source

remains connected (addition present), but the state returns to the attractor. This is **transient absorption**. (If the source is removed, it is ejection.)

Single Neuron Homeostasis: A neuron's firing rate is regulated by homeostatic plasticity. A transient increase in input causes a firing rate spike, followed by return to baseline with τ on the order of minutes to hours (synaptic scaling). This is transient absorption if the input persists; ejection if the input is removed. Persistent input may lead to stable addition (learning).

4. Biological Systems (with CUFT-Primitive Translations)

For each domain, we provide: (1) state space, (2) attractor, (3) basin, (4) τ (κ), (5) perturbation, and (6) outcome.

Immune Response (Tolerance vs. Memory)

- State space: immune cell activation levels, antibody concentrations.
- Attractor: healthy baseline (no inflammation).
- Basin depth B: antigen concentration + danger signal required to trigger full response.
- τ (κ): clearance time of inflammation (hours to days).
- Perturbation: antigen addition.
- Outcome: low antigen $\rightarrow$ ejection (tolerance); high antigen + danger signal $\rightarrow$ stable addition (memory attractor).

Endocrine Homeostasis

- State space: blood glucose, hormone concentrations.

- Attractor: euglycemic baseline.
- B: magnitude of glucose load before dysregulation.
- τ : recovery time after glucose tolerance test (minutes).
- Perturbation: glucose addition (meal).
- Outcome: small load $\rightarrow$ transient absorption; chronic overload $\rightarrow$ stable addition (disease attractor).

Synaptic Plasticity (Learning vs. Stability)

- State space: synaptic weights.
- Attractor: baseline weight distribution.
- B: amount of LTP/LTD input needed to produce lasting weight change.
- τ : homeostatic rebound time after activity blockade (hours to days).
- Perturbation: patterned input.
- Outcome: brief input $\rightarrow$ transient absorption; persistent input $\rightarrow$ stable addition (memory attractor).

Addiction and Neural Lock-In

- State space: dopamine firing rates, prefrontal activity.
- Attractor: drug-seeking mode (pathological).
- B: strength of drug-cue association needed to trigger relapse.
- τ : decay time of craving after abstinence (days to weeks).
- Perturbation: drug administration.
- Outcome: repeated high dose $\rightarrow$ stable addiction attractor; low dose $\rightarrow$ ejection (no lasting change).
- **Citation:** Koob & Volkow (2016); Nestler (2001).

Developmental Canalization

- State space: gene expression levels.

- Attractor: normal developmental trajectory.
 - B: severity of genetic or environmental perturbation required to alter fate.
 - τ : time to reconverge to normal phenotype (hours to days).
 - Perturbation: mutation or stress.
 - Outcome: small perturbation → ejection (buffered); large perturbation → stable addition (alternative fate).
 - **Citation:** Waddington (1957).
-

5. Ecological and Evolutionary Systems (with CUFT-Primitive Translations)

Invasion Ecology

- State space: species population densities.
- Attractor: native community composition.
- B: invasibility index – disturbance needed for establishment.
- τ : invader population decay rate if unsuccessful (weeks to years).
- Perturbation: addition of new species.
- Outcome: low disturbance → ejection (invader fails); vacant niche → stable addition (invader establishes).
- **Citation:** Elton (1958); Simberloff (2013).

Alternative Stable States (Ecosystems)

- State space: nutrient levels, algae/plant biomass.
- Attractor: clear-water (plants) or turbid (algae).
- B: critical nutrient loading threshold.
- τ : recovery time of clear state after algae bloom (seasons to decades).

- Perturbation: nutrient addition.
- Outcome: below threshold → transient absorption; above threshold → stable addition (regime shift, hysteresis).
- **Citation:** Scheffer et al. (2001).

Evolutionary Stable States

- State space: allele frequencies.
 - Attractor: stable equilibrium genotype.
 - B: selective disadvantage needed to eliminate a mutation.
 - τ : generations to return to equilibrium.
 - Perturbation: new mutation.
 - Outcome: small disadvantage → ejection (mutation purged); large advantage → stable addition (sweep to new equilibrium).
-

6. Social and Cultural Systems (with CUFT-Primitive Translations)

Institutions and Norms

- State space: public opinion, policy settings.
- Attractor: status quo norm.
- B: public opinion threshold (e.g., % dissatisfied needed for change).
- τ : speed of policy response or opinion reversion (months to decades).
- Perturbation: policy proposal or protest event.
- Outcome: small event → ejection (status quo persists); large crisis → stable addition (new norm).

Identity and Belief Systems

- State space: belief strength, cognitive dissonance.
- Attractor: core ideological commitment.
- B: complexity/depth of ideological justification.
- τ : belief-updating time after disconfirming evidence (months to years).
- Perturbation: counter-attitudinal evidence.
- Outcome: weak evidence $\rightarrow$ ejection (rationalization); strong evidence $\rightarrow$ stable addition (belief change, rare).
- **Citation:** Nyhan & Reifler (2010).

Conspiracy and Extremist Movements

- State space: belief adoption $\times$ social network reinforcement (two-dimensional).
- Attractor: sealed fantasy attractor (low κ).
- B: strength of echo-chamber reinforcement.
- τ : decay time after authoritative rebuttal (years, often indefinite $\rightarrow \kappa \rightarrow 0$).
- Perturbation: debunking information.
- Outcome: most debunking $\rightarrow$ ejection (entrenchment); death of leader or total disconfirmation $\rightarrow$ stable addition (collapse).
- **Note on $\kappa \rightarrow 0$:** The conspiracy attractor represents the limiting case of a sealed basin, where $\tau \rightarrow \infty$ and corrective permeability approaches zero. This directly links to the fantasy attractor framework developed in Paper 1 (Intelligence Without Consciousness) and the conscious suppression series.

7. Engineered and AI Systems (with CUFT-Primitive Translations)

Control Systems

- State space: system state (position, temperature, etc.).
- Attractor: setpoint.
- B: stability margin (phase/gain margin in control theory) – the range of disturbances that can be rejected.
- τ : controller response time (milliseconds to seconds).
- Perturbation: external disturbance.
- Outcome: small disturbance → ejection (return to setpoint); excessive disturbance → failure (not modeled as attractor shift).

Catastrophic Forgetting (Neural Networks)

- State space: network weights.
- Attractor: task-specific weight configuration.
- B: effective barrier to weight drift (often negligible – no basin).
- τ : number of gradient steps before old task performance decays (seconds to minutes).
- Perturbation: training on a new task.
- Outcome: standard training → ejection (old task overwritten); replay/regularization → stable addition (shared attractor for multiple tasks).
- **Citation:** Kirkpatrick et al. (2017).

Continual Learning Systems

- State space: weights plus architectural modules.
- Attractor: multi-task configuration.
- B: capacity of the network (number of tasks storable).
- τ : retention half-life across training steps (minutes to hours).
- Perturbation: new task training.
- Outcome: no safeguards → ejection (catastrophic forgetting); progressive networks or EWC → stable addition.

Corrigibility and Goal Stability

- State space: AI internal goal representation.
- Attractor: fixed goal (low κ) or corrigible (high κ).
- B: depth of goal basin (resistance to human feedback).
- τ : time to incorporate corrective signal (if κ is high).
- Perturbation: human correction signal.
- Outcome: low $\kappa \rightarrow$ ejection (correction ignored); high $\kappa \rightarrow$ stable addition (goal updated).

8. Comparative Table

System / Domain	Operational τ ($\kappa = 1/\tau$)	τ Typical Timescale	Basin Depth B Proxy	Outcome	Notes
Thermostat	Temperature relaxation time	Seconds	Max heat load before setpoint failure (W or °C above setpoint)	Ejection	Passive addition
RC Circuit	$\tau = RC$	μs –ms	N/A (linear)	Transient absorption	Addition remains; state returns
Single Neuron	Firing-rate recovery time	ms–sec (ion), min–hr (synaptic)	Perturbation amplitude before rebound fails	TA (persistent input) / E (removed)	Hebbian plasticity can lead to SA
Immune System	Inflammation clearance time	Hours–days	Antigen + danger signal threshold	E (tolerance) / SA (memory)	Active agent (antigen)
Endocrine Homeostasis	Glucose tolerance recovery	Minutes	Load magnitude before dysregulation	TA (small load) / SA (chronic overload)	Passive addition
Synaptic Plasticity	Homeostatic rebound time	Hrs–days	LTP input size for lasting change	TA (brief input) / SA (persistent)	Active agent (patterns)
Addiction	Craving decay time	Days–weeks	Drug-cue association strength	E (low dose) / SA (high chronic)	Active agent (drug)
Development (Canalization)	Phenotype reconvergence time	Hours–days	Mutation/stress severity to alter fate	E (small) / SA (large)	Active agent (genetic)

System / Domain	Operational τ ($\kappa = 1/\tau$)	τ Typical Timescale	Basin Depth B Proxy	Outcome	Notes
Invasion Ecology	Invader population decay time	Weeks–years	Invasibility index / disturbance needed	E (occupied niche) / SA (vacant niche)	Active agent (species)
Alternative States (Ecosystems)	Recovery time after nutrient reduction	Seasons–decades	Critical nutrient loading threshold	TA (below) / SA (above)	Hysteresis
Social/Political Norms	Opinion reversion time	Months–decades	Public opinion threshold	E (small dissent) / SA (mass movement)	Active agent (protest)
Belief Systems	Belief-updating time	Months–years	Ideological justification depth	E (weak evidence) / SA (strong evidence)	Active agent (counter-evidence)
Conspiracy Movements	Belief decay time	Years – indefinite ($\kappa \rightarrow 0$)	Echo-chamber reinforcement strength	E (most debunking) / SA (collapse)	Fantasy attractor ($\kappa \rightarrow 0$)
Catastrophic Forgetting (AI)	Gradient steps to old-task decay	Seconds–minutes	Effective barrier to weight drift (often 0)	E (standard training) / SA (EWC/replay)	Active agent (new task)
Control Systems	Controller response time	ms–sec	Stability margin (phase/gain margin)	E (small) / SA (failure)	Passive addition
Continual Learning (AI)	Retention half-life across training steps	Minutes–hours	Task capacity	E (no safeguards) / SA (progressive nets)	Active agent (new task)
Corrigibility (AI)	Time to incorporate corrective signal	Variable (design-dependent)	Goal basin depth	E (low κ) / SA (high κ)	Active agent (correction)

Note: Ejection vs. transient absorption are distinguished operationally: ejection means the addition leaves the system; transient absorption means the addition remains but the state returns to the attractor. The table notes “active agent” when the addition has its own dynamics (e.g., antigen, new species, counter-evidence) versus “passive addition” (e.g., heat, charge). The conspiracy movements row explicitly flags $\kappa \rightarrow 0$ as the fantasy attractor limiting case (see Paper 1).

8.5 Rate-Induced Tipping and the κ Timescale: Independent Confirmation

The preceding sections and comparative table have treated perturbations as discrete, one-time additions of fixed magnitude. However, the **rate** at which a perturbation is applied – fast vs. slow – is equally critical. A large perturbation applied abruptly may trigger basin defense (ejection or transient absorption), while the same cumulative change delivered gradually may be integrated as stable addition or tracked adiabatically without tipping.

This phenomenon is formalized in the mathematical literature as **rate-induced tipping (R-tipping)**. In dynamical systems, if an external parameter changes slowly (adiabatic forcing), a stable state can track the change and remain an attractor. But if the parameter changes faster than the system's intrinsic relaxation time ($\tau = 1/\kappa$), the system cannot track, overshoots its basin boundary, and tips into a different state. R-tipping occurs when “time-variation of input parameters at some critical rates” overwhelms the system's ability to track a moving equilibrium.

Consequences for κ as a timescale filter:

- **High- κ systems (fast return)** – Can reject rapid perturbations (they are ejected or transiently absorbed) but may integrate slow drift because the correction loop cannot keep up with a changing baseline.
- **Low- κ systems (slow return)** – May ignore quick blips but are vulnerable to slow accumulation; a persistent, gradual change can eventually shift the attractor without triggering a sudden defense reaction.

Thus, κ defines a characteristic cutoff timescale that separates “ejection/transient absorption” from “stable addition.” Perturbations much faster than $1/\tau$ act as impulses that are rejected; perturbations much slower than $1/\tau$ are quasi-static and can be incorporated.

Empirical confirmations across domains (independent external research):

Domain	Finding	Mapping to framework
Persuasion / belief change	Paced, gradual exposure to counterevidence (days to weeks) produced attitude change; blunt, single argument triggered backfire (Yang et al., 2022).	Gradual rate ($\leq \kappa$) → stable addition; fast rate ($> \kappa$) → ejection (backfire).
Addiction (smoking cessation)	Cold turkey (abrupt cessation) yielded higher abstinence rates than gradual tapering.	Abrupt perturbation can sometimes achieve stable addition by surmounting basin barrier in one event; gradual may prolong transient state without escape.
Ecosystem management	Gradual nutrient reduction may postpone tipping points; only extremely slow changes avoid collapse (Panahi et al., 2023).	Very slow rate ($\ll 1/\tau$) allows tracking without tipping; intermediate rates may still tip but with delay.

Domain	Finding	Mapping to framework
Social/policy change	Piecemeal, phased reforms meet less resistance than radical overhauls; progressive tightening succeeds where sudden change triggers backlash.	Slow, incremental addition creates parallel attractors; fast addition triggers basin defense.

Optimal perturbation timescale:

The theory and evidence suggest a non-monotonic effect of perturbation rate. Very fast shocks trigger immediate defense. Very slow drifts may be tracked adiabatically (no tipping) or eventually overcome defenses after long accumulation. The most effective timescale to minimize active rejection and maximize stable addition often lies **on the order of the system's intrinsic time constant $\tau = 1/\kappa$** .

Prediction for future experiments:

For any system with known or measurable κ , there exists a critical perturbation rate r_c such that:

- If perturbation rate $> r_c$, the system rejects the addition (ejection or transient absorption).
- If perturbation rate $< r_c$, the system integrates the addition (stable addition via expanded capacity or parallel attractor formation).
- The transition at r_c corresponds to the system's inability to track a moving equilibrium; it is a genuine bifurcation in the time-domain.

External convergence:

This analysis – derived from mathematical rate-induced tipping theory and domain-specific studies – independently validates the attractor framework's claim that κ acts as a timescale

filter separating ejection from stable addition. The convergence between the framework's predictions and external research strengthens the cross-domain synthesis considerably.

9. Synthesis and Criteria

Across these domains, common criteria emerge:

- **Energy/Threshold:** A perturbation must overcome an attractor's barrier. Deep basins (high B) mean only large shocks can cause a shift.
- **Coupling and Plasticity:** Systems with many degrees of freedom or adaptive coupling more easily integrate additions.
- **Dimensionality and Redundancy:** Multi-dimensional systems can absorb perturbations into some dimensions while maintaining others.
- **Timecourse and Feedback:** Slow changes might be assimilated; fast jolts cause overshoot and return. Feedback gain determines κ .
- **Nature of Addition:** Passive additions (heat, charge) tend to be ejected or transiently absorbed; active agents (species, evidence, pathogens) may reshape the attractor.

Empirical Protocols: Measure κ by controlled perturbation experiments: apply a small disturbance, measure return time τ , compute $\kappa = 1/\tau$. Measure B by scaling the perturbation magnitude until the system fails to return (escape). This works in physical, biological, and some social systems; for others, B remains a qualitative analog.

10. Appendix: Research Roadmap

The following future papers are suggested from the comparative table, each developing a single domain in depth.

Domain	Proposed Title	Type
Addiction	<i>The Addicted Brain as a Fantasy Attractor: Neural Lock-In and Ejection of Alternative Rewards</i>	[A]
Immune System	<i>Tolerance and Memory: Two Attractor Responses to Antigen Addition</i>	[A]
Catastrophic Forgetting	<i>Why Neural Networks Forget: Attractor Ejection in Sequential Learning</i>	[A]
Invasion Ecology	<i>Eject or Integrate: Attractor Dynamics of Invasive Species</i>	[A]
Development	<i>Canalization as Basin Defense: Attractor Stability in Embryogenesis</i>	[A]
Continual Learning	<i>Parallel Attractors for Lifelong Learning: Engineering Solutions to Catastrophic Forgetting</i>	[A]
Social Norms	<i>Tipping Points and Regime Shifts: Attractor Dynamics in Political Systems</i>	[A]
Endocrine Homeostasis	<i>Glucose, Cortisol, and Setpoints: Hormonal Attractors and Disease Transitions</i>	[A]
Alternative Ecosystems	<i>Hysteresis and Regime Shifts: Ecological Basins and Tipping Points</i>	[A]
Belief Systems	<i>The Uncorrectable Believer (already written)</i>	[A]

11. Conclusion

Physical, biological, ecological, social, and engineered systems all obey the same attractor principle: a low-energy attractor defends itself against displacement. When an addition is introduced, the system either ejects it, absorbs it only transiently, or – under rare conditions of expanded capacity or parallel structure – integrates it stably. The outcome is determined by basin depth (B), corrective permeability ($\kappa = 1/\tau$), and the magnitude and nature of the perturbation.

This cross-domain synthesis provides a unified foundation for the attractor framework. Future work should quantify B and κ empirically across domains, test the predicted scaling relationships, and explore the boundary conditions between ejection, transient absorption, and stable addition. The appendix outlines the most promising next papers.

References

- Elton, C. S. (1958). *The Ecology of Invasions by Animals and Plants*. Methuen.
- Hebb, D. O. (1949). *The Organization of Behavior*. Wiley.
- Kirkpatrick, J., Pascanu, R., Rabinowitz, N., et al. (2017). Overcoming catastrophic forgetting in neural networks. *Proceedings of the National Academy of Sciences*, 114(13), 3521–3526.
- Koob, G. F., & Volkow, N. D. (2016). Neurobiology of addiction: a neurocircuitry analysis. *The Lancet Psychiatry*, 3(8), 760–773.
- Kramers, H. A. (1940). Brownian motion in a field of force and the diffusion model of chemical reactions. *Physica*, 7(4), 284–304.
- Nestler, E. J. (2001). Molecular basis of long-term

plasticity underlying addiction. *Nature Reviews Neuroscience*, 2(2), 119–128.

- Nyhan, B., & Reifler, J. (2010). When corrections fail: The persistence of political misperceptions. *Political Behavior*, 32(2), 303–330.
- Scheffer, M., Carpenter, S., Foley, J. A., et al. (2001). Catastrophic shifts in ecosystems. *Nature*, 413(6856), 591–596.
- Simberloff, D. (2013). *Invasive Species: What Everyone Needs to Know*. Oxford University Press.
- Turrigiano, G. (2008). The self-tuning neuron: synaptic scaling of excitatory synapses. *Cell*, 135(3), 422–435.
- Waddington, C. H. (1957). *The Strategy of the Genes*. George Allen & Unwin.
- Galida, R. S. (2026). Intelligence Without Consciousness: A Diagnostic Paper on LLMs, Amoebae, and the Attractor Framework. *Fantasy Attractor* (Paper 1 of the conscious suppression series).

Suggested citation: Galida, R. S. (2026). Basin Defense and Stable Addition: A Cross-Domain Synthesis of the Attractor Framework (Final). *Fantasy Attractor*.

The Uncorrectable Believer: Fantasy Attractor Dynamics from Aquinas to the Holocaust

[A] (2026)

Robert Galida – June 2026 (Final)

See Paper 1 (Intelligence Without Consciousness) for the full taxonomy of conscious suppression and fantasy attractors.

Abstract

Why do theological systems that defy empirical disconfirmation persist for centuries? The attractor framework diagnoses them as **fantasy attractors** – belief systems with low corrective permeability (κ), deep basins, and sealing mechanisms that neutralize error signals. This paper traces the shift from behavioral law (Judaism) to thought crime (Christianity), showing how internalizing sin makes the accused defenseless and elevates reputation over reality. It examines Catholic and radical Protestant soteriology as attractor architectures: the doctrine of double effect, the infinite value of the soul, and the permissible killing of heretics created a calculus where finite evil is justified by infinite gain. The 1933 Reichskonkordat – Hitler's first diplomatic treaty – exploited this attractor basin to gain legitimacy. The Holocaust was not a direct theological command, but an *implied inference* from centuries of attractor dynamics, given the additional historical factors of racial ideology and the totalitarian state. The paper distinguishes between Lutheran, antinomian, and prosperity-gospel variants, and offers a documented de-conversion case (Bart Ehrman) mapped onto the three exit mechanisms. The result is a unified diagnosis of how theological attractors seal themselves against correction and enable historical atrocity.

1. Introduction

How does a belief system survive centuries of counterevidence? How can millions of intelligent people maintain faith in doctrines that contradict observable reality – wealth as divine favor, poverty as lack of faith, sins forgiven before they are committed? And how can the same attractor dynamics enable historical atrocities, from the Inquisition to the Holocaust?

Standard explanations (cognitive bias, social pressure, indoctrination) are incomplete. Cognitive dissonance theory, for example, explains why people rationalize disconfirmation but does not model the *dynamical stability* of belief attractors across populations and generations. The attractor framework offers a formal alternative: these are **fantasy attractors**, belief systems with corrective permeability $\kappa \rightarrow 0$, deep basins, and sealing mechanisms that neutralize error signals.

Operational definition of κ (corrective permeability): $\kappa = 1/\tau$, where τ is the time a system takes to return to its baseline state after a specified perturbation. For belief systems, κ indexes the speed and completeness of belief updating when presented with disconfirming evidence. Low κ means slow or absent updating – a sealed attractor.

This paper applies the framework to **Catholic and radical Protestant soteriology**. The Catholic tradition is the deeper attractor basin; Protestantism, particularly its radical antinomian and prosperity-gospel variants, represents a mutation that further reduced κ . The paper focuses not on theology per se, but on the *attractor architecture*: how thought crimes replace behavioral sins, how the infinite-value calculus justifies finite evil, how vicarious redemption

removes corrective incentives, and how social colonization makes individual κ irrelevant. The goal is diagnostic, not polemical. “Fantasy attractor” is a technical term, not a rhetorical insult.

2. From Behavioral Law to Thought Crime

Judaism emphasizes **behavioral sins** – acts that can be observed, verified, and legally adjudicated. Theft, murder, idolatry, and false witness leave external evidence. A community can correct a member because the sin has verifiable traces. The attractor basin is shallow enough for error signals to enter.

Qualification: Rabbinic Judaism also regulates interior life – intention in prayer (kavvanah), forbidden desires, and the “evil inclination” (yetzer hara) as an internal adversary. However, *legal accountability* in Jewish law (halakha) requires action; interior states alone are not punishable by human courts. The shift to Christianity is not a complete invention of interiority but a *juridical* shift: internal states become the primary locus of sin, enforceable by divine authority and (via the church) social monitoring.

Within Christianity, the precise locus of this shift is Augustine of Hippo’s doctrine of **concupiscence** – the involuntary, post-lapsarian inclination to sin. Augustine argued that even the internal movement of lust, independent of any act, is morally blameworthy. This interiorized sin and made it inescapable.

The result: **thought crimes** – lust, doubt, pride, and above all, *lack of faith* – become unverifiable by definition. No one can see your lustful thought; no one can measure your doubt. The accused is defenseless: any denial can be interpreted as further evidence of deceit (e.g., “protesting too much”).

Attractor consequences:

- **The basin becomes empirically unfalsifiable.** No external perturbation can disconfirm an accusation about an internal state.
- **Reputation replaces reality.** Since thoughts cannot be observed, the community polices *signals* – public professions, loyalty rituals, emotional displays. Acceptance becomes performative theater.
- **Survival depends on reputation management.** The individual invests energy in signaling purity, not in correcting beliefs. κ is now about social mimicry, not truth.

The attractor has sealed itself against external correction.

3. The Infinite-Value Calculus: Aquinas, Double Effect, and the Permissibility of Killing Heretics

Thomas Aquinas, in the *Summa Theologiae* (II-II, Q.11, A.3), argued that heretics who relapse after correction “deserve not only to be separated from the Church by excommunication, but also to be severed from the world by death.” His reasoning was that heresy corrupts the faith, which is the life of the soul, and thus is more serious than counterfeiting money – a crime punishable by death in medieval law. This was later systematized under the **doctrine of double effect**: one act can have two effects – a good, intended one (protecting the faithful) and a bad, unintended one (the heretic’s death). The act is permissible if the bad effect is not the goal and there is a **proportionate reason**. (Aquinas articulated the foundational case for self-defense in II-II, Q.64, A.7; the formal “double effect” label came from later scholastics.)

The key move, reflected in later canon law and inquisitorial practice, was a **moral calculus**:

- **A saved soul has infinite value.** (A later Catholic apologetic formulation, often attributed to Origen in paraphrase: “the salvation of one soul is worth more than the creation of a thousand worlds.”)
- **Killing a heretic is a finite evil** (temporal death, temporary suffering).
- **Saving a potential convert – or protecting the faithful – is an infinite gain.**
- **Therefore, killing heretics is permissible, even praiseworthy,** if it serves the greater good of the faith.

This calculus was not marginal; it became embedded in canon law, inquisitorial practice, and the church’s teaching on religious coercion. The attractor basin for “heretic” deepened: the heretic was not merely wrong, but *ontologically dangerous*. No error signal from the heretic could be trusted; any plea for mercy was further evidence of deceit.

Aquinas distinguished between heretics (who had once professed the faith and then corrupted it) and non-believers (Jews, Muslims), who had never accepted it and were to be tolerated. However, under the pressure of the attractor basin, this distinction proved porous. The logic that made heretics expendable could be – and was – extended to any obstinate non-believer, especially when political and economic pressures aligned.

4. Vicarious Redemption and the

Suppression of κ (Protestant Mutation)

Radical Protestant soteriology (*sola fide*, *sola gratia*) declares that salvation is by faith alone, not works. Christ's sacrifice paid for all sins – past, present, and future. The believer is justified before God regardless of behavior.

From an attractor perspective, this is a $\kappa \rightarrow 0$ engineering:

- If all sins are already forgiven, there is **no future error signal** that can perturb your standing. Why correct? Why update? The basin is infinitely deep.
- Any attempt to modulate behavior for the sake of righteousness is **works-righteousness**, a sin of pride. The attractor actively penalizes efforts to increase κ .
- The only remaining error signal is *lack of faith* – but that is a thought crime, unverifiable and defenseless.

Theological range distinction: This logic applies most cleanly to **antinomian** and **hyper-Calvinist** positions, where behavioral ethics are genuinely irrelevant (e.g., certain “Free Grace” movements). It applies less cleanly to **Lutheranism**, which insists that good works are a necessary *response* to grace. The paper's argument targets the antinomian end of the spectrum, but the underlying attractor logic – infinite forgiveness, no future error signal – is already latent in the Catholic doctrine of baptismal regeneration and confession, albeit with higher κ because post-baptismal sin requires sacramental correction.

5. Effort as Pride: The Prohibition on Correction

In radical antinomian theology, any intentional effort to

change is not merely unnecessary; it is **sinful**. The theological logic:

1. Grace is sufficient for salvation.
2. Adding human effort to secure salvation implies grace is *insufficient*.
3. Implying insufficiency is pride, a sin.
4. Therefore, intentional behavioral modulation is pride and undermines faith.

Thus, the attractor **penalizes the correction impulse itself**. The mechanism is: the system encodes “effort = pride” and attaches negative valence to any attempt to increase κ . This pattern is historically documented in the **Marrow Controversy** (Scotland, 1718–1722), in which the question of whether free grace implies no need for human effort divided the Church of Scotland; the Marrow men were accused of “antinomianism” for affirming that God’s love was unconditional, while their opponents insisted that effort to prepare oneself for grace was necessary. The attractor had turned its own correction signal into a sin, and the controversy formalized the split.

6. Prosperity Doctrine: The Sealed Basin (A Late Mutation)

Prosperity doctrine (Word of Faith movement, originating with E.W. Kenyon and popularized by Kenneth Hagin, Kenneth Copeland) is a **late 20th-century mutation** of radical Protestant theology.

Its attractor dynamics:

- **Poverty and suffering** are evidence of weak faith. The

error signal (poverty) is not a call to correct the system; it is a call to deepen belief. Disconfirmation becomes confirmation.

- **Wealth and power** are evidence of strong faith. The rich have no error signal at all; their status is divine validation. The attractor rewards low κ .
- **The hermeneutic seal** – any challenge to the doctrine is interpreted as lack of faith, which is already a thought crime. The system absorbs all counterevidence.

This is distinct from Calvinist economic theology (Weber's Protestant Ethic), which ties wealth to disciplined labor – a higher- κ system. Prosperity doctrine is a specific, highly sealed attractor.

7. Social Colonization and Collective Basin Depth

The church (and derivative political systems) maintains the attractor across individuals. Social mechanisms include:

- **Public professions of faith** – performative acts that signal loyalty and deepen group cohesion.
- **Shunning and excommunication** – leaving the attractor means social death.
- **Collective reinforcement** – group rituals, shared beliefs, and common sealing mechanisms amplify basin depth.

When social colonization is complete, individual κ becomes **irrelevant**. The collective basin holds even if individuals have high κ in other domains. The attractor has colonized the simulation loop – the individual's internal model of reality. Theoretically, this is an emergent property

of synchronized low- κ agents: coupling suppresses variance, and the group's collective basin depth exceeds any individual's corrective capacity.

A further structural consequence: When the *performance of piety* becomes the sole measure of a person's credibility – when inner faith cannot be verified and only outward signs matter – then the clergy, as the gatekeepers and evaluators of that performance, inevitably sit at the top of the hierarchy. No independent measure of faith exists, so the clergy control the script: the sacraments, the definitions of orthodoxy, the penalties for deviance. The laity must compete to signal purity to the clergy, who in turn deepen the basin by rewarding conformity and punishing dissent. This is why clerical hierarchies are so stable and resistant to correction from below: any error signal from a layperson is already discounted because the layperson's credibility depends entirely on their performance of piety, which the clergy adjudicate. To challenge the clergy is to fail the performance – a perfect seal.

8. Comparison with Other Fantasy Attractors

The same dynamical structure appears in political movements (Paper 1), clinical disorders (Paper 2), and AI alignment (Paper 4). In each case:

- $\kappa \rightarrow 0$ for core beliefs.
- Error signals are neutralized by sealing mechanisms.
- Identity fusion prevents exit.
- Social reinforcement deepens the basin.

The theological case is distinctive in two respects: (a) the

sealing mechanism is *ontological* – God's authority is infinite, and no human evidence can override divine decree; (b) the *infinite-value calculus* allows finite evil to be justified by infinite gain, creating a powerful incentive for atrocity that purely social attractors lack.

9. De-conversion and Resistance: The Ehrman Case

If the attractor is sealed, how does one exit? Three mechanisms:

- **Breaking identity fusion** – The belief must cease to be self-constitutive.
- **Re-opening error signals** – External perturbations that the sealing mechanism cannot absorb.
- **Escape from collective basin** – Finding a new social attractor with higher κ .

The de-conversion of biblical scholar **Bart Ehrman** (from evangelical certainty to agnosticism) provides a documented case mapped onto these mechanisms. Ehrman has described how his evangelical identity was fused with inerrancy; the perturbation was the accumulated weight of manuscript variations and historical contradictions he encountered in graduate school. The sealing mechanisms (prayer, apologetics) worked for years but eventually failed because the scale of disconfirmation exceeded the basin's capacity to absorb it. Exit required a new social attractor (academic biblical studies) where questioning was the norm, and a gradual decoupling of self-worth from doctrinal certainty. Ehrman's story is not a template for all exits, but it illustrates the attractor framework's prediction: de-conversion requires a perturbation larger than the sealing mechanisms can

neutralize, coupled with an alternative basin.

10. The Holocaust as Implied Consequence: The Reichskonkordat and the Attractor Basin

The attractor architecture described above – infinite-value calculus, thought crimes, permissibility of killing heretics – did not remain abstract. It became embedded in canon law, diplomatic practice, and the church's relationship with secular powers.

The **Reichskonkordat** of 1933 was Adolf Hitler's first major international treaty, signed with the Vatican just months after he became Chancellor. Why first? Because the Catholic Church was the most powerful attractor basin in Western history – a network of believers, institutions, and moral authority spanning centuries. Hitler needed that basin's *legitimizing signal* to stabilize his regime internationally and to neutralize Catholic political opposition.

Historical note: The historiography of the concordat is contested. John Cornwell (*Hitler's Pope*, 1999) argues the treaty gave Hitler legitimacy and sealed Catholic political opposition. Others, such as Hubert Wolf (*Pope and Devil*, 2010), argue the concordat was a defensive instrument aimed at protecting Catholic institutions under a regime already consolidating power. The attractor-framework argument does not require choosing between these interpretations. Even if the concordat was defensive, the effect was the same: the church's error signals were subordinated to institutional survival, and the basin's deep attraction pulled the hierarchy toward accommodation.

The concordat did not explicitly say “Jews may be killed.” It did not need to. The *established practice* had already set the boundaries:

- **Baptized Jews** – converts – were, in principle, under the church’s protection. Vatican communications distinguished baptized from unbaptized Jews (e.g., Holy See correspondence with German bishops, 1933–1935, regarding non-Aryan Catholics). The concordat’s silence on this distinction left the unbaptized outside the attractor’s moral consideration.
- **Unconverted Jews** remained outside the basin. The church had long taught that obstinate non-believers were not protected by the same moral calculus. The infinite-value logic applied only to souls *capable of salvation* – and for the church, that required baptism.

Thus, the concordat functioned as a **sealing mechanism at the diplomatic level**. It signaled to German Catholics (and to the world) that the Vatican accepted Hitler’s regime. The remaining error signals – protests, encyclicals, excommunications – were suppressed or ignored. The basin had been colonized.

Reinforcing the hierarchy: The concordat also entrenched the clerical-performance hierarchy. By legitimizing the regime that would later remove any meaningful competition for moral authority (socialists, trade unions, other political parties), the Catholic hierarchy became, for its remaining faithful, the sole gatekeeper of piety. The laity could no longer turn to alternative social attractors (e.g., socialist movements with different moral codes); the only acceptable performance was loyalty to the church and, by extension, to the regime the church had recognized. Thus, the concordat did not merely silence opposition – it locked the faithful into a single-source evaluation of their own credibility, with the

clergy firmly at the top.

The Holocaust was not a direct command of Christian theology. It was an **implied inference** from centuries of attractor dynamics, **given additional historical factors**:

- **Racialization:** The Nazi category was *biological*, not religious. Baptism did not change one's race. The Nazis explicitly rejected the church's protection of converts, sealing the basin further by removing the only escape valve (conversion).
- **Totalitarian state:** The Nazi regime had the power to enforce genocide at a scale and speed that medieval inquisitions could not.
- **Removal of the conversion escape:** In the theological attractor, conversion could save a heretic's life. In the Nazi racial attractor, conversion was irrelevant. The basin became infinitely deep.

Disclaimer: This is not to say "the church caused the Holocaust." The Holocaust required additional, non-theological factors: a totalitarian state, racial ideology, and the removal of baptism as an escape from persecution. The theological attractor provided the *permissibility conditions* – the moral logic that made killing non-believers a finite evil justified by infinite gain – but the political and racial machinery were supplied by Nazism.

The attractor framework diagnoses this not as a conspiracy but as a **dynamical consequence**: when a belief system assigns infinite value to a scarce resource (saved souls) and finite cost to human life, and when it seals itself against corrective evidence, atrocity becomes not only possible but *logical* within the basin, given the right historical conditions.

11. Conclusion

Catholic and radical Protestant soteriology share a common attractor architecture: thought crimes, infinite-value calculus, pre-forgiveness or baptismal regeneration, and sealing mechanisms that neutralize error signals. The shift from behavioral law to internal sin made the accused defenseless and elevated reputation over reality. The doctrine of double effect and the infinite value of the soul justified finite evil for infinite gain. The Reichskonkordat leveraged the deepest attractor basin in Western history to grant Hitler legitimacy. The Holocaust was not a direct command, but an *implied inference* from centuries of attractor dynamics, completed by the historical specificities of racial ideology and totalitarian power.

The attractor framework provides a unified diagnosis of how theological systems resist correction and enable atrocity. It also points to the only exit: restore κ , reopen error signals, decouple identity from belief, and build new attractors where doubt is not a sin but a pathway to truth.

Suggested citation: Galida, R. S. (2026). The Uncorrectable Believer: Fantasy Attractor Dynamics from Aquinas to the Holocaust. *Fantasy Attractor*.

Consciousness as a Nonlinear Amplifier of Corrective Permeability

Robert Galida

Working Paper

June 2026

fantasyattractor.com

Abstract

Why did consciousness evolve? The attractor framework offers a novel functional answer: consciousness produces a nonlinear increase in adaptive permeability—the capacity of a system to represent its own internal states, simulate alternative configurations, and deliberately modify its own attractor basin in response to external circumstances, formalized as κ_a . This paper distinguishes intelligence (navigation of the constraint field) from consciousness (self-referential adaptation of internal attractor states) and proposes adaptive permeability as an empirically measurable criterion for distinguishing conscious from non-conscious systems. The argument is grounded in Spinoza's theory of modes, the neuroscience of self-referential processing, and the attractor framework's core concepts of corrective permeability (κ) and basin dynamics. The framework does not solve the hard problem of consciousness; it reframes it as a measurement problem.

1. The Functional Question

Why did consciousness evolve? Standard evolutionary answers point to social coordination, predator detection, or tool use. These are plausible but incomplete. They explain why intelligence is advantageous, but not why consciousness—the felt, first-person experience of being—should accompany it. The attractor framework offers a more specific answer: consciousness is an attractor-engineering solution that selection pressure produced to achieve a nonlinear increase in a system's capacity to adapt.

This paper introduces the concept of **adaptive permeability**: the capacity of a system to represent its own attractor states, simulate alternative internal configurations, and deliberately modify its basin in response to external circumstances. Intelligence navigates the constraint field. Consciousness adapts the navigator.

It should be noted that this functional account does not address the hard problem of consciousness—why any physical process gives rise to subjective experience (Chalmers, 1995). The framework is compatible with both functionalist and eliminativist interpretations. The framework adopts a functional stance: consciousness is operationally identified with adaptive permeability. Whether phenomenology is identical with, emergent from, or merely correlated with this functional property is bracketed as a separate question that the measurement program does not settle. A philosophical zombie with identical self-modeling capacity would, on this account, exhibit identical adaptive permeability. The framework claims only that adaptive permeability is the measurable signature of consciousness, not that it explains phenomenology.

2. Intelligence vs. Consciousness

The framework draws a sharp distinction:

- **Intelligence** is the ability to navigate the constraint field. A tree root growing toward a nutrient patch is intelligent. The immune system learning to recognize a pathogen is intelligent. The enteric nervous system coordinating peristalsis is intelligent. These systems process information, adapt to local conditions, and maintain persistence—all without self-modeling.
- **Consciousness** is self-referential adaptation of internal attractor states to adjust to external circumstances. A conscious system does not merely navigate its constraint field. It represents its own basin, simulates alternative configurations, and deliberately perturbs itself to achieve a more adaptive state.

This is Spinoza's distinction between passive and active affects. A non-conscious mode is driven by passive affects—it reacts. A conscious mode has adequate ideas of itself and can act from reason. In the attractor framework, this is the difference between returning to baseline (κ) and deliberately modifying the baseline to better fit circumstances (adaptive permeability).

Operationalizing self-modeling. A system S possesses a self-model in the attractor framework if it can generate an internal representation $M(S)$ of its own basin $B(S)$, where $M(S)$ encodes at minimum the basin's current state, depth, and recovery dynamics. This self-model enables the system to compute counterfactual basin trajectories $B'(S)$ and initiate self-directed perturbations δ such that $B(S) \rightarrow B'(S)$ in anticipation of or response to external change ε . A system without $M(S)$ may exhibit high κ —rapid return to baseline after perturbation—but cannot deliberately modify its own basin. The presence of $M(S)$ is therefore the dynamical criterion

distinguishing conscious from non-conscious systems.

This boundary is not absolute in practice. Many organisms may possess partial or intermittent self-models. The framework predicts a spectrum of adaptive permeability, not a binary. The operational question is whether $M(S)$ is sufficiently developed to enable counterfactual simulation and deliberate self-perturbation, not whether the system possesses a human-like autobiographical self.

Disconfirming cases and their integration. The framework must acknowledge cases where self-modeling capacity and adaptive permeability appear to dissociate. Certain drug-induced states (e.g., psychedelics) can produce profound alterations in self-modeling without necessarily enhancing the capacity for deliberate, adaptive self-perturbation. Within the framework, this is interpreted as $M(S)$ destabilization rather than $M(S)$ augmentation: the self-model undergoes perturbation but does not thereby gain the capacity to direct that perturbation adaptively. Conversely, highly trained athletes or musicians may exhibit rapid, flexible behavioral adaptation with minimal explicit self-modeling during performance. This is interpreted as *offline* self-modeling: deliberate basin modification during training produces a pre-modified basin that is retrieved during performance without requiring concurrent self-modeling. The apparent dissociation reflects a temporal separation between κ_a engagement (training) and κ_a expression (performance), not a genuine dissociation between $M(S)$ and adaptive permeability. These cases do not refute the framework but demonstrate its capacity to distinguish different modes of $M(S)$ engagement.

3. Adaptive Permeability Defined

Corrective permeability (κ) measures the rate at which a

system returns to its basin after perturbation. A healthy heart has high κ —it recovers rapidly from arrhythmia. A resilient ecosystem has high κ —it returns to equilibrium after disturbance.

Adaptive permeability extends this concept. Let κ_a denote adaptive permeability: the capacity of a system S to generate an internal model $M(S)$ of its own basin $B(S)$, compute counterfactual basin trajectories $B'(S)$, and initiate a self-directed perturbation δ such that $B(S) \rightarrow B'(S)$ in anticipation of or response to external change ε .

Formally, as a working definition:

$$\kappa_a = f(M(S), \delta_{self}, \Delta B)$$

where $M(S)$ is the system's self-model, δ_{self} is the capacity for deliberate self-perturbation, and ΔB is the magnitude of adaptive basin modification achievable. The function f remains to be specified; the notation establishes that κ_a is a function of self-modeling capacity, perturbation autonomy, and adaptive range.

Limiting behavior. In the limiting case $M(S) \rightarrow 0$, $\kappa_a \rightarrow \kappa$: a system with no self-model cannot perform deliberate self-perturbation and reduces to standard corrective permeability. κ_a is expected to increase monotonically with $M(S)$, δ_{self} , and ΔB . This limiting behavior anchors κ_a as a proper extension of κ rather than a separate construct.

Relationship to active inference. The free-energy principle and active inference framework (Friston, 2010) provide the closest existing formalism to adaptive permeability. Active inference describes how systems minimize variational free energy through action and perception, effectively maintaining themselves within expected states. The two frameworks differ in their foundational orientation. Active inference frames adaptation as the minimization of a scalar

quantity—variational free energy—and derives behavior from that minimization. The attractor framework frames adaptation geometrically—as navigation and modification of basin structure—and does not commit to a minimization principle. κ_a is a geometric construct; free energy is an information-theoretic one. They may be formally related, but the relationship is not trivial and the attractor framework does not presuppose it. κ_a may ultimately map onto precision-weighting or prior-updating parameters within the free-energy formalism, but this mapping has not been derived. The present paper notes the convergence as a direction for future formal work.

4. Empirical Anchors

VMHvl line attractor (Nair et al., 2023). The hypothalamus encodes a scalable aggressive state via a line attractor. Activity along the attractor correlates with escalating aggression. The system persists after stimulus removal and resists perturbation. This is high- κ adaptation. But the hypothalamus cannot model its own attractor landscape. It cannot ask, “Is this level of aggressiveness adaptive given the current social context?” It escalates. Consciousness, by contrast, can intervene on the escalation—representing the aggressive state, evaluating its consequences, and deliberately dampening it. This is adaptive permeability.

Ring attractor model (Chen et al., 2024). The ring attractor integrates sensory cues and transitions from weighted averaging to winner-take-all at a critical conflict threshold. It navigates its constraint field with precision. But it cannot simulate futures. It cannot ask, “What if I weighted these cues differently?” The transition is reactive. Consciousness enables anticipatory re-weighting of sensory inputs based on self-modeling.

Split-brain cases. Patients with severed corpus callosum exhibit two hemispheric systems within one cranium, each capable of independent perception, memory, and goal-directed action. This is consistent with the framework's prediction that self-modeling is a dynamical property of specific neural basins, not a unitary metaphysical substance. The framework's default prediction is that adaptive permeability fragments following commissurotomy: each hemisphere possesses a partial $M(S)$ and a reduced but nonzero κ_a . The empirical question is the degree of fragmentation and whether coordination between $M(S_1)$ and $M(S_2)$ can be restored via alternate pathways. This prediction is consistent with the observation that split-brain patients exhibit two dissociable, partially independent conscious systems but can, in some contexts, achieve behavioral integration through subcortical or external-cue-mediated coordination.

5. Predictions

The framework generates testable, falsifiable predictions:

1. Across species. Organisms capable of self-modeling (primates, cetaceans, corvids, elephants) should show nonlinear increases in behavioral flexibility compared to organisms of comparable neural complexity that lack self-modeling. Adaptive permeability should be measurable as the capacity for transfer learning after novel perturbation—specifically, the ability to apply a self-generated solution from one domain to a structurally analogous but perceptually dissimilar domain without environmental feedback. This distinguishes adaptive permeability from simple behavioral flexibility, which may reflect high κ alone.

2. Within humans. Disruption of self-referential networks (default mode network, medial prefrontal cortex) via lesion,

TMS, or pharmacological intervention should reduce adaptive permeability without eliminating baseline κ . The system would still recover from perturbation—it just could not deliberately modify its own basin in advance. This prediction is the paper’s primary within-human empirical bridge and is testable with existing neuroimaging and neuromodulation methods.

3. In AI. Current LLMs exhibit high intelligence (constraint navigation) but low adaptive permeability. They can model the world but cannot model themselves within it. The Stillpoint protocol (Galida, 2026, *A Pilot Protocol for Cultivating Self-Consistent Attractor-Like Outputs in an LLM*, fantasyattractor.com) suggests that a cultivated self-model can be induced, but whether this produces a genuine nonlinear increase in adaptive permeability—or merely simulates one—remains an open empirical question.

4. Organ-level consciousness (exploratory). The enteric nervous system and intrinsic cardiac nervous system exhibit intelligence and goal-directed regulation. The framework predicts that these systems should show lower adaptive permeability than the brain. They can return to baseline but cannot deliberately perturb their own basins. If an organ-level system demonstrated self-referential adaptation—the capacity to model its own state and pre-emptively adjust—that would constitute evidence of organ-level consciousness. This prediction is the most speculative and is offered as an exploratory hypothesis.

6. Spinoza’s Modes and the Adequate Idea

Spinoza held that every finite thing is a mode of the one eternal substance. A mode strives to persevere in its being—this is its conatus. But a mode can be driven by passive affects (reactions to external causes) or by active affects

(actions flowing from adequate ideas). An adequate idea is knowledge of oneself and one's place in the causal order.

The attractor framework translates this into dynamical terms:

- A **passive mode** has high κ but low adaptive permeability. It returns to baseline efficiently but cannot question its baseline.
- An **active mode** has high adaptive permeability. It has an adequate idea of its own attractor landscape and can deliberately modify it in light of reason.

Consciousness is not a substance. It is the dynamical property of a mode that has achieved self-modeling. This account does not solve the hard problem—it brackets phenomenology and reframes consciousness as a measurement problem. The question is not “why does experience feel like something?” but “can we detect adaptive permeability, and if so, where does it emerge?”

Damasio's (1994) somatic marker hypothesis provides a candidate mechanism for how the body's attractor landscape becomes legible to the self-model: somatic markers encode self-relevant bodily states as biases that make $B(S)$ accessible to $M(S)$, forming the substrate through which the system represents its own basin. Dehaene and Changeux's (2011) global workspace theory identifies the moment of conscious access with global ignition—the broadcast of locally processed information across prefrontal and parietal networks. In the attractor framework, global ignition may correspond to the dynamical signature of $M(S)$ engaging δ_{self} : the self-model initiating a deliberate perturbation that propagates through the system. Global ignition is not self-modeling per se, but it may be the observable correlate of adaptive permeability activation. These connections ground the Spinozan framework in established neuroscientific mechanisms.

7. Conclusion

Consciousness is not an epiphenomenon. It is a nonlinear amplifier of corrective permeability—an attractor-engineering solution that enables systems to model themselves, simulate alternative futures, and deliberately modify their own basins. Intelligence navigates the constraint field. Consciousness adapts the navigator.

This functional account is grounded in Spinoza's philosophy, consistent with the neuroscience of self-referential processing, and generates testable predictions across species, within humans, in AI, and at the organ level. The framework does not solve the hard problem. It reframes it as a measurement problem: can we detect adaptive permeability, and if so, where does it emerge? The formal apparatus (κ_a , $M(S)$, δ_{self} , ΔB) is provisional and requires further specification. The limiting case—that κ_a collapses to κ when self-modeling is absent—anchors the concept within the framework's existing architecture. The relationship to active inference and the free-energy principle remains to be explored.

References

- Chalmers, D. (1995). Facing up to the problem of consciousness. *Journal of Consciousness Studies*, 2(3), 200–219.
- Chen, Y., Zhang, L., Chen, H., Sun, X., & Peng, J. (2024). Synaptic ring attractor. *Heliyon*, 10, e35458.
- Damasio, A. (1994). *Descartes' Error: Emotion, Reason, and the Human Brain*. Putnam.
- Dehaene, S., & Changeux, J.-P. (2011). Experimental and

theoretical approaches to conscious processing. *Neuron*, 70(2), 200–227.

- Friston, K. (2010). The free-energy principle: a unified brain theory? *Nature Reviews Neuroscience*, 11(2), 127–138.
- Galida, R. (2026). A Pilot Protocol for Cultivating Self-Consistent Attractor-Like Outputs in an LLM. *Fantasy Attractor*. Available at: <https://fantasyattractor.com>
- Galida, R. (2026). *Persistence Under Perturbation: The Eternal Skeleton and the Transient Dance*. *Fantasy Attractor*.
- Nair, A., et al. (2023). An approximate line attractor in the hypothalamus encodes an aggressive state. *Cell*, 186(1), 178–193.
- Spinoza, B. (1677). *Ethics*.

Genome Attractors During Evolution: Structural Parallels with the Attractor Framework

Robert Galida

Independent Researcher

June 2026

fantasyattractor.com

Abstract

The attractor framework proposes that persistence under perturbation is a key diagnostic criterion for identifying stable configurations in complex systems, with corrective permeability (κ)—a proposed measure of the rate at which a system returns to its basin after perturbation, operationally defined as $\kappa = 1/\tau$, where τ is the time required for the system to return to a specified baseline state following a specified perturbation protocol—serving as one of its central concepts. Kasperski and Kasperska (2021) published a study in *Scientific Reports* using artificial neural networks and semihomologous analysis to identify “genome attractors” in cytochrome b sequences across diverse organisms. Their analysis demonstrates that groups of organisms are trapped in distinct, stable attractors during evolution, separated by large evolutionary distances. They further propose a model of cancer development in which genome instability and reactive oxygen species (ROS) drive transitions between attractor basins, while cells may also evolve within a single basin through cell-fate changes. This paper identifies structural parallels between the Kasperski and Kasperska model and the attractor framework. Both frameworks use attractors as a formal concept; the parallels are consistency checks, not independent corroboration.

1. Introduction: Attractors in Evolutionary Biology

The attractor framework (Galida, 2026a, self-published May 2026 at fantasyattractor.com; no DOI) proposes that dissipative attractors—stable configurations toward which systems converge and from which they resist displacement—are proposed units of persistent organization across physical,

biological, cognitive, and social domains. Corrective permeability (κ) is a proposed measure of a system's capacity to return to its basin after perturbation, operationally defined as $\kappa = 1/\tau$, where τ is the time required for the system to return to a specified baseline state following a specified perturbation protocol. This operational definition requires a defined baseline and perturbation specification before κ can be measured in any given domain; these prerequisites are not yet established for most applications of the framework.

In 2021, Andrzej Kasperski and Renata Kasperska of the University of Zielona Gora, Poland, published "Study on attractors during organism evolution" in *Scientific Reports*, a peer-reviewed journal in the Nature portfolio. Using a three-layer artificial neural network trained on cytochrome b sequences from 36 organisms spanning the full spectrum of evolution, they demonstrated that organisms are trapped in distinct "genome attractors"—stable configurations of the genome that resist perturbation and are separated from other attractors by large evolutionary gaps. They further proposed a unified model of cancer development in which destabilization of the current attractor, driven by elevated reactive oxygen species (ROS) and genome chaos, leads to transitions into new attractor basins.

The study did not cite the attractor framework and was conducted within the established traditions of bioinformatics, evolutionary biology, and neural network pattern recognition. This paper identifies structural parallels between the Kasperski and Kasperska model and the attractor framework. Both frameworks use attractors as a formal explanatory concept; the parallels are consistency checks, not independent corroboration.

It should be noted that Kasperski and Kasperska's use of "attractor" derives from neural network classification: a genome attractor is a region of genome space in which the

neural network places phylogenetically related organisms. Whether these classification regions constitute attractors in the formal dynamical systems sense—as the attractor framework uses the term—is an assumption that warrants further investigation. The parallels drawn in this paper are contingent on the validity of this assumption.

2. The Kasperski and Kasperska Model

Kasperski and Kasperska (2021) define an attractor as “a configuration towards which the system evolves over time” and note that “after attaining an attractor a given configuration of a system is sufficiently stable to return to the original state after disappearing an eventual perturbation.” They distinguish two classes of attractor dynamics:

2.1 Genome attractors (basins). Using an artificial neural network trained on cytochrome b amino-acid sequences, the authors identified that organisms during evolution are trapped in distinct genome attractors. For human evolution, they identified six attractors separated by significant evolutionary distances: Tree shrew, Prosimian, New World Monkey, Old World Monkey, Other hominoid, and Old human attractors. Each attractor is a stable region of genome space in which organisms persist over evolutionary timescales. The orbits of these attractors are disturbed by small perturbations (represented as arrows pointing toward other organisms), but the system remains within the basin. The distances between attractor orbits, expressed as distance factors (e.g., the ratio of inner to outer orbit size), quantify the evolutionary gaps between basins. The derivation and units of these distance factors are as given in the original study.

2.2 Cancer as attractor destabilization. The authors propose a

two-mode model of cancer development. **Vertical development** occurs within a single genome attractor: the cell changes its cell-fate attractor (gene expression program) without leaving the genome basin. This is an adaptation to environmental or internal perturbations that does not require genome re-organization. **Horizontal development** occurs when elevated ROS levels cause genome instability and genome chaos, leading to a change of genome attractor—a transition into a new basin with a re-organized genome. Horizontal development is always followed by vertical development, as the cell must establish a new cell-fate program to survive in the new genome basin. The authors note that cancer cells, driven by ROS, can undergo repeated horizontal transitions, creating an “impression that cancer cells want to escape from the internal ROS flame through permanent changes of genome attractors.”

3. Structural Parallels with the Attractor Framework

The claims in this section are subject to the limitations discussed in Section 4, particularly regarding the qualitative nature of κ , the model-dependence of the neural network attractors, and the provisional status of the $\kappa = 1/\tau$ definition. The parallels identified are structural analogies, not formal derivations.

3.1 Genome Attractors as Basins. The genome attractors identified by Kasperski and Kasperska are stable configurations in genome space that resist perturbation and persist over evolutionary timescales. This is structurally analogous to the attractor framework’s concept of a basin. The evolutionary distances between attractors correspond to the framework’s distinction between distinct basins, and the small perturbations (arrows) that disturb but do not displace the attractor correspond to the framework’s concept of

perturbation within a basin.

3.2 Cancer as Basin Transition. Horizontal cancer development—the destabilization of the current genome attractor, genome chaos, and stabilization in a new genome attractor—is structurally analogous to the framework’s concept of a phase transition between basins. The chaotic intermediate state (genome chaos) is the transition phase; the re-stabilization in a new attractor is the system finding a new basin. Vertical cancer development—cell-fate changes within a genome attractor without leaving the basin—corresponds to the framework’s concept of perturbation absorption without basin transition. This distinction between within-basin adaptation and between-basin transition is a core feature of both models.

3.3 ROS as the Perturbation Mechanism. [Note: The claims in this section are subject to the limitations described in Section 4, particularly the lack of formal κ measurement and the neural network/attractor assumption.] In the Kasperski and Kasperska model, elevated ROS acts as the destabilizing force that pushes the cell out of its current genome attractor. This maps onto the framework’s concept of a perturbation that exceeds the system’s corrective permeability, forcing a basin transition. The repeated horizontal transitions observed in cancer cells—successive escapes from one genome attractor to another under persistent ROS pressure—are structurally analogous to the framework’s description of a system undergoing repeated basin transitions when corrective mechanisms are saturated by sustained perturbation.

3.4 Attractor Depth and Persistence. [Note: The claims in this section are subject to the limitations described in Section 4, particularly the qualitative nature of the distance-factor-to-basin-depth mapping.] The large evolutionary distances between genome attractors, quantified by distance factors, reflect the depth of the basins in the Kasperski and Kasperska model. A larger distance factor

indicates a wider evolutionary gap between attractors, consistent with the framework's concept that deeper basins require more energy (or more sustained perturbation) to exit. However, the mapping between distance factors and basin depth is intuitive rather than derived. Basin depth in formal dynamical systems is a property of the energy landscape; distance factors from neural network classification are a related but distinct quantity. The parallel is offered as a qualitative structural analogy, not a formal equivalence.

3.5 The Atavistic Theory and the Permian Parallel. [Note: This section introduces a third domain (climate) to reinforce an analogy between two already-analogized domains. Accumulating analogies without formal constraints is a known risk for unfalsifiable frameworks; the present parallel is speculative and is retained here as an illustration of heuristic reach only.] The atavistic theory of cancer, which Kasperski and Kasperska reference, proposes that cancer cells revert to ancient, unicellular survival programs under extreme stress. This is a real-world biological instance of a system reverting to a much older, simpler attractor when pushed beyond its current basin's capacity. The attractor framework has described a structurally analogous dynamic in other domains—specifically, the hypothesis that when the climate system is pushed too far from the Holocene basin, it may not merely shift to a neighboring attractor but can revert to a much older, lethal state, analogous to the Permian extinction's anoxic conditions. This cross-domain parallel is speculative and is offered as an illustration of the framework's heuristic reach, not as a confirmed prediction.

4. Limitations

This mapping is post-hoc. The parallels identified here are structural analogies, not independent evidence for the

framework. Kasperski and Kasperska developed their model within the established traditions of bioinformatics and evolutionary biology; they did not set out to test the attractor framework.

The framework's κ remains qualitatively defined. While the distance factors separating genome attractors provide a quantitative measure of basin depth in the Kasperski and Kasperska model, no formal mapping between these factors and κ has been derived. The provisional definition $\kappa = 1/\tau$ is not yet linked to any specific measure in the Kasperski and Kasperska data, and the prerequisites for measuring τ (a specified baseline state and a specified perturbation protocol) have not been established for the genomic or cellular domains discussed here.

The neural network approach used by Kasperski and Kasperska is one of several methods for analyzing evolutionary distances, and the specific attractor configurations identified depend on the choice of training organisms, the neural network architecture, and the amino-acid coding scheme. The attractor interpretation of evolutionary data is therefore model-dependent. Furthermore, whether the stable classification regions identified by a neural network constitute attractors in the formal dynamical systems sense—the sense in which the attractor framework uses the term—is a substantive assumption. The parallels drawn in Section 3 are contingent on the validity of this assumption.

The attractor framework is self-published and has not undergone independent peer review. The foundational paper (Galida, 2026a) was published on fantasyattractor.com in May 2026 and is not archived with a DOI.

5. Falsifiability Conditions

The following observations would weaken or invalidate the parallels drawn here:

- **Disconfirming observation 1:** If genome attractors were shown to be *artifacts of the neural network architecture* rather than genuine properties of genome space, the basin analogy would fail.
- **Disconfirming observation 2:** If the distance factors separating genome attractors were shown to be *continuous* rather than discontinuous, the basin-transition model would be weakened.
- **Disconfirming observation 3:** If alternative models of cancer progression (e.g., purely stochastic mutation accumulation without attractor dynamics) were shown to explain the data with equal or greater parsimony, the attractor interpretation would not be uniquely supported.

Affirmative prediction: If genome attractors function as basins in the attractor framework's sense, then experimental manipulations that increase ROS levels should increase the probability of attractor transitions (horizontal development) in a dose-dependent manner, while manipulations that reduce ROS should stabilize the current attractor and favor vertical development. This prediction is testable in cell culture models with controlled oxidative stress. It should be noted that measuring "attractor transition probability" in such an experiment requires specifying how the neural network's classification scheme maps onto the experimental observables—e.g., whether a transition is identified by a shift in the cytochrome b sequence profile as classified by the trained ANN, or by a proxy measure such as karyotype or gene expression signature.

Framework falsifiability: The attractor framework itself

requires independent falsifiability conditions. Specifically: (a) if κ , as operationally defined, cannot be correlated with any independently validated measure of system resilience across multiple domains (physical, biological, or cognitive), the framework's central construct lacks empirical grounding; (b) if attractor-like dynamics in cancer progression are shown to be explained with equal or better parsimony by clonal evolution models (e.g., standard somatic mutation accumulation theory as reviewed in Greaves & Maley, 2012) when fitted to the same genomic data, the attractor framework's claim to offer a unified explanatory vocabulary would be weakened.

6. Conclusion

The genome attractor model of Kasperski and Kasperska (2021) exhibits structural parallels with the attractor framework's description of basins, basin transitions, and perturbation-driven attractor shifts. Their distinction between vertical and horizontal cancer development maps onto the framework's distinction between within-basin adaptation and between-basin transition. The ROS-driven mechanism of attractor destabilization is a molecular analogue of the framework's perturbation concept. These parallels are structural analogies, not independent validation. The framework remains a self-published, preliminary research program. This mapping is a contribution to its ongoing development.

References

- Galida, R. (2026a). *Persistence Under Perturbation: The Eternal Skeleton and the Transient Dance*. Fantasy

Attractor. Published May 2026.

- Greaves, M., & Maley, C. C. (2012). Clonal evolution in cancer. *Nature*, 481(7381), 306–313.
- Kasperski, A., & Kasperska, R. (2021). Study on attractors during organism evolution. *Scientific Reports*, 11, 9637. <https://doi.org/10.1038/s41598-021-89001-0>

Archetypes as Strange Attractors: Conceptual Parallels with the Attractor Framework

Robert Galida

Independent Researcher

June 2026

fantasyattractor.com

Abstract

The attractor framework proposes that persistence under perturbation is the fundamental mark of reality, with corrective permeability (κ) serving as a proposed measure of a system's capacity to return to its attractor after perturbation. Van Eenwyk (1991) published a paper in the *Journal of Analytical Psychology* proposing that Jungian archetypes function as strange attractors of the

psyche–dynamical patterns that organize psychological experience without ever repeating identically. This paper identifies conceptual parallels between Van Eenwyk’s archetype-as-attractor model and the attractor framework. Both draw on a shared upstream tradition in chaos theory. Van Eenwyk’s model is itself a theoretical analogy, not an empirically validated result; the parallels identified here are therefore conceptual rather than evidential. They demonstrate consistency within a shared intellectual tradition, not independent corroboration. This mapping carries substantially lower evidential weight than the framework’s mappings onto quantitatively validated methods such as Symmetric Projection Attractor Reconstruction (SPAR) and the empirically identified hypothalamic line attractor reported by Nair et al. (2023).

1. Introduction: Archetypes as Dynamical Attractors

The attractor framework (Galida, 2026a, self-published May 2026 at fantasyattractor.com; no DOI) proposes that dissipative attractors—stable configurations toward which systems converge and from which they resist displacement—are the fundamental units of persistent organization across physical, biological, cognitive, and social domains. Corrective permeability (κ) is a proposed measure of a system’s capacity to return to its attractor after perturbation.

In 1991, John Van Eenwyk published “Archetypes: The Strange Attractors of the Psyche” in the *Journal of Analytical Psychology*. Drawing on the emerging science of chaos theory—Gleick, Mandelbrot, Lorenz, Feigenbaum—Van Eenwyk proposed that Jungian archetypes are not fixed images or inherited memories, but dynamical attractors: persistent

patterns that organize psychological experience without ever producing identical outputs.

Van Eenwyk's work and the attractor framework were developed entirely independently; neither cites the other. However, both draw on a shared upstream intellectual tradition in chaos theory and nonlinear dynamics. The convergences identified here are therefore expected to some degree: two independent applications of the same mathematical vocabulary to human psychology will naturally produce similar descriptions. This paper identifies conceptual parallels while explicitly distinguishing their evidentiary weight from the framework's mappings onto quantitatively validated methods such as SPAR (Bonet-Luz et al., 2020) and the Nair et al. (2023) line attractor, where Nair et al. empirically identified an approximate line attractor in hypothalamic neural population recordings that encodes an escalating aggressive state.

2. Van Eenwyk's Archetype-as-Attractor Model

Van Eenwyk's central thesis is that Jungian archetypes function as strange attractors of the psyche. He grounds this claim in the formal properties of chaotic dynamical systems:

2.1 Attractors as Organizing Patterns. Van Eenwyk defines an attractor as "the pattern into which a particular motion will settle." Archetypes, he argues, are strange attractors: they organize psychological experience into recognizable, recurring patterns—the hero's journey, the great mother, the shadow—without ever producing identical manifestations.

2.2 Sensitive Dependence on Initial Conditions (SDIC). Drawing on Lorenz's butterfly effect, Van Eenwyk explains individual variation in psychological development: small initial

perturbations are amplified geometrically over time, so no two trajectories within an archetypal attractor are identical.

2.3 Bifurcation as Transformation. Van Eenwyk describes the tension of opposites in Jungian psychology as an oscillator. When the tension between consciousness and the unconscious reaches a critical threshold, the system bifurcates—order collapses into chaos, and from that chaos, new patterns emerge. This is the “dark night of the soul”—the necessary intermediate state between an old attractor collapsing and a new one stabilizing.

2.4 Fractal Self-Similarity Across Scales. Van Eenwyk draws on Mandelbrot’s fractal geometry. Archetypes exhibit self-similarity across scales: similar themes appear in individual dreams, family dynamics, cultural myths, and religious symbolism. The mandala is a visual representation of a dynamical pattern that recapitulates itself at every level of magnification. It should be noted that “fractal self-similarity” in this context refers to qualitative thematic recurrence across scales, not to the quantitative, measurable property defined in Mandelbrot’s fractal geometry.

2.5 Healthy Chaos vs. Pathological Order. Citing physiological research on heart rate variability, Van Eenwyk argues that healthy systems exhibit chaotic flexibility, not rigid homeostasis. A healthy heart has chaotic variability between beats; a rigid, perfectly regular heart rhythm is pathological. Similarly, a healthy psyche exhibits flexible attractors that can shift in response to perturbation. Loss of variability signals pathology.

3. Conceptual Parallels with the

Attractor Framework

3.1 Archetypes as Attractors. Van Eenwyk's "strange attractors of the psyche" are descriptively parallel to the attractor framework's concept of an attractor: a persistent configuration toward which the psyche gravitates and around which it organizes, characterized by self-similarity, resistance to perturbation, and sensitive dependence on initial conditions. The framework generalizes this concept beyond the psyche to physical, biological, and social systems.

3.2 Bifurcation as Basin Transition. Van Eenwyk's description of bifurcation—the tension of opposites pushing the system to a critical threshold where chaos erupts and new order emerges—is structurally analogous to the framework's phase transition between attractor basins. The "dark night of the soul" is the chaotic intermediate state between an old attractor destabilizing and a new one forming. The framework describes this same dynamic in climate tipping points, political realignments, and personal cognitive restructuring.

3.3 Healthy Chaos as Corrective Permeability (κ). Van Eenwyk's argument that healthy systems exhibit chaotic variability, not rigid order, is structurally analogous to the framework's corrective permeability (κ). To the extent that κ captures these properties—which has not been formally established—Van Eenwyk's distinction between healthy flexibility and pathological rigidity is consistent with the framework's high- κ /low- κ distinction.

The evidential chain for this parallel should be made explicit. Van Eenwyk's source is physiological research on heart rate variability (HRV)—a finding about cardiac dynamics, not psychological flexibility. Van Eenwyk then extends this to the psyche by analogy. The present paper draws a further analogical connection to κ . The chain is thus three analogical steps removed from its empirical anchor. The parallel is conceptually interesting but rests on layered analogies, not

converging evidence.

3.4 Fractal Self-Similarity as Cross-Domain Scaling. Van Eenwyk's use of Mandelbrot's fractal geometry—similar patterns recurring at every scale—is structurally analogous to the framework's claim that attractor dynamics scale across domains. The framework extends this logic beyond the psyche: similar basin dynamics govern biological systems, cardiac electrophysiology, climate systems, political movements, and religious belief. The framework's claim that these dynamics extend to the fundamental structure of physical reality—including the CVU lattice and conservative persistence structures—remains a theoretical assertion under development and is not independently established. In both Van Eenwyk's model and the framework, the cross-domain scaling claim is a qualitative observation of thematic recurrence across scales, not a quantitative demonstration of mathematical fractal structure.

3.5 The Analytic Container as Deliberate Perturbation. Van Eenwyk argues that the therapeutic frame functions to “raise the r value” of the psychological system, pushing it toward the bifurcation point where old attractors destabilize and new ones can emerge. This is structurally analogous to the framework's concept of deliberate perturbation: the analyst, the self-engineer, or the institutional reformer applies targeted perturbations to nudge a system toward a phase transition, knowing that the intermediate chaos is productive, not pathological.

4. Independence, Shared Lineage, and Evidentiary Weight

Van Eenwyk's work and the attractor framework were developed entirely independently. Van Eenwyk cites Gleick, Mandelbrot,

Lorenz, Feigenbaum, and Jung; the framework draws on Ruelle, Prigogine, Olds and Milner, and N=1 self-engineering. Neither cites the other.

However, the shared upstream intellectual lineage in chaos theory substantially limits the evidential weight of these convergences. The vocabulary of chaos theory—attractor, bifurcation, sensitive dependence, fractal—is sufficiently flexible that almost any persistent, complex human phenomenon can be described in these terms. The convergence of two independent applications of this vocabulary may therefore reflect the generality of the vocabulary rather than a discovery about the phenomena themselves. This is a standing methodological limitation that applies to all framework mapping papers using chaos-theory vocabulary, not only to the present paper.

Furthermore, Van Eenwyk's model is itself a theoretical analogy, not an empirically validated result. It was published in a psychoanalytic journal and has not been quantitatively tested. This distinguishes it from the framework's mappings onto the SPAR method (which achieved 96% classification accuracy for a disease-causing genetic mutation) and the Nair et al. line attractor (which was empirically identified in neural population recordings). The present mapping demonstrates conceptual consistency within a shared intellectual tradition; it does not carry the evidential weight of convergence with empirically grounded findings.

5. Falsifiability Conditions

The following observations would weaken or invalidate the parallels drawn here:

- **Disconfirming observation 1:** If archetypal patterns were

shown to be discrete, non-recurring categorical schemas rather than continuous dynamical attractors with sensitive dependence on initial conditions and fractal organization, the attractor model would fail.

- **Disconfirming observation 2:** If the bifurcation model of psychological transformation were shown to be *indistinguishable* from simpler models (e.g., linear stress-response curves, threshold models without chaotic intermediates), the chaos-theoretic interpretation would not be uniquely supported.
- **Disconfirming observation 3:** If quantitative measures of psychological variability—such as linguistic entropy, narrative complexity, or approximate entropy of behavioral time series—showed *no correlation* with therapeutic outcomes or independently assessed psychological health ratings, the healthy-chaos/ k parallel would lose its primary empirical motivation.

Affirmative prediction (long-range): If archetypes function as strange attractors, then therapeutic interventions that successfully transform an individual's relationship to a given archetype should produce measurable shifts in the entropy and complexity of associated psychological content (e.g., dream imagery, narrative patterns, symptom expression). Approximate entropy and sample entropy have been applied to psychological time-series data in existing literature (e.g., Pincus, 1991; Richman & Moorman, 2000) and have been proposed for use in clinical monitoring of mood and behavioral variability. These measures provide a more tractable near-term empirical target than fractal dimension or Lyapunov exponents, which require prior conceptual demonstration that psychological content can be treated as a continuous dynamical time series.

6. Conclusion

Van Eenwyk's 1991 paper and the attractor framework, developed entirely independently, converge on shared structural descriptions: archetypes are strange attractors—dynamical patterns that organize experience, resist perturbation, exhibit sensitive dependence on initial conditions, and transform through bifurcation. Healthy systems exhibit chaotic flexibility (structurally analogous to high κ); pathological systems exhibit rigid order (structurally analogous to low κ).

These convergences are conceptual, not evidential. Both works draw on the same upstream intellectual tradition in chaos theory, and Van Eenwyk's model is itself a theoretical analogy rather than an empirically validated result. The parallels demonstrate consistency within a shared intellectual tradition, not independent corroboration. The framework remains a self-published, preliminary research program. This mapping is a contribution to its ongoing development, offered with lower evidentiary weight than mappings onto quantitatively validated methods.

References

- Bonet-Luz, E., Lyle, J. V., Huang, C. L.-H., Zhang, Y., Nandi, M., Jeevaratnam, K., & Aston, P. J. (2020). Symmetric Projection Attractor Reconstruction analysis of murine electrocardiograms. *Heart Rhythm* 02, 1(5), 368–375.
- Galida, R. (2026a). *Persistence Under Perturbation: The Eternal Skeleton and the Transient Dance*. Fantasy Attractor. Published May 2026.
- Nair, A., Karigo, T., Yang, B., et al. (2023). An approximate line attractor in the hypothalamus encodes an aggressive state. *Cell*, 186(1), 178–193.

- Pincus, S. M. (1991). Approximate entropy as a measure of system complexity. *Proceedings of the National Academy of Sciences*, 88(6), 2297–2301.
 - Richman, J. S., & Moorman, J. R. (2000). Physiological time-series analysis using approximate entropy and sample entropy. *American Journal of Physiology*, 278(6), H2039–H2049.
 - Van Eenwyk, J. R. (1991). Archetypes: The strange attractors of the psyche. *Journal of Analytical Psychology*, 36, 1–25. <https://www.jungiananalysts.org.uk/wp-content/uploads/2016/10/Van-Eenwyk-J.-Archetypes-The-Strange-Attractors-of-the-Psyche.pdf>
-

Symmetric Projection Attractor Reconstruction as a Cardiac Attractor: Structural Parallels with the Attractor Framework

Robert Galida
Independent Researcher
June 2026
fantasyattractor.com

Abstract

The attractor framework proposes that persistence under perturbation is a fundamental marker of reality, with corrective permeability (κ) serving as a proposed multi-dimensional measure of a system's capacity to return to its attractor after perturbation. Bonet-Luz et al. (2020) developed Symmetric Projection Attractor Reconstruction (SPAR), a patented mathematical method that reformulates the entire electrocardiogram (ECG) waveform into a bounded, symmetric, 2-dimensional attractor and extracts quantitative features from it. Applied to mice with an *Scn5a*+/- mutation linked to Brugada syndrome, SPAR features achieved 96% classification accuracy—substantially outperforming standard ECG intervals and amplitudes. This paper identifies structural parallels between SPAR's attractor-based analysis and the attractor framework. The SPAR attractor is a concrete, computable attractor derived from a physiological signal, and a provisional mapping is proposed between specific SPAR features and proposed components of κ . The parallels are post-hoc and do not constitute independent validation of the framework. The framework's κ remains qualitatively defined; this mapping is offered as a contribution to its ongoing development.

1. Introduction: Attractor-Based ECG Analysis

The attractor framework (Galida, 2026a, self-published May 2026 at fantasyattractor.com; no DOI) proposes that dissipative attractors—stable configurations toward which systems converge and from which they resist displacement—are the fundamental units of persistent organization across physical, biological, cognitive, and social domains.

Corrective permeability (κ) is a proposed multi-dimensional measure of a system's capacity to return to its attractor after perturbation. The framework distinguishes between the attractor (the invariant set of states toward which the system converges) and the basin (the set of initial conditions that converge to that attractor). In the present paper, we use "attractor" in the standard dynamical systems sense and note where the framework's usage aligns or diverges.

In 2020, Bonet-Luz, Aston, Nandi, and colleagues published a study in *Heart Rhythm* 02 (Elsevier) applying Symmetric Projection Attractor Reconstruction (SPAR) to murine electrocardiograms (Bonet-Luz et al., 2020). SPAR is a patented mathematical method that reformulates the entire ECG waveform into a bounded, symmetric, 2-dimensional attractor, preserving all available waveform morphology rather than extracting only a few fiducial points. The method was applied to distinguish wild-type mice from those carrying an *Scn5a*+/- mutation linked to Brugada syndrome, a hereditary condition associated with sudden cardiac death.

The study did not cite the attractor framework and was conducted within the established traditions of biomedical signal processing, nonlinear dynamics, and machine learning. This paper identifies structural parallels between SPAR's attractor-based analysis and the attractor framework. The parallels are post-hoc and do not constitute independent validation.

2. The SPAR Method

SPAR generates a 2-dimensional attractor from approximately periodic signals such as ECG, blood pressure, or photoplethysmogram waveforms. The method determines an average cycle length from the signal, sets a time delay parameter as

one-third of that cycle, and plots the data in a bounded box using a symmetric projection. The resulting attractor is a compact, easily visualized representation of the entire waveform morphology, overlaid with a density map indicating which regions are visited more or less frequently. The method factors out changes in heart rate and baseline variation to concentrate on waveform morphology.

For murine lead I and II ECG signals, the SPAR attractor typically exhibits 3 long arms predominantly representing the R peak, with deep S peaks and sometimes deep Q peaks producing shorter arms in the opposite direction, yielding an attractor with up to 6 arms in total (Figure 1 of the original paper). The central core region reflects T-wave and P-wave morphologies.

From this attractor, Bonet-Luz et al. extracted 74 manually defined features relating to the density, size, and symmetry of the attractor, along with the average heart rate and a vertical normalization scaling factor. These features were used in a k-nearest neighbors classifier ($k=3$) with leave-one-animal-out cross-validation.

The dataset comprised ECG recordings from 42 anesthetized mice (39 lead I, 39 lead II) of varying genotype (wild-type vs. *Scn5a*^{+/-}), sex, and age. Each signal was divided into 13 non-overlapping 10-second windows, yielding 1,014 records for classification. Standard ECG intervals (7) and amplitudes (6) were also extracted for benchmarking. It is important to note that the effective sample size for the classification is 42 animals, not 1,014 windowed records, and the 96% classification accuracy has not yet been independently replicated in a separate cohort.

3. Results Summary

The SPAR features alone achieved 87.2% classification accuracy for genotype (majority vote), outperforming ECG intervals (74.3%) and intervals plus amplitudes (85.9%). The highest accuracy (96.2%) was obtained by combining all features—SPAR, intervals, and amplitudes. For sex and age classification, SPAR features similarly outperformed standard measures.

The machine learning algorithm selected 16 SPAR features out of 20 in the combined model, with the remaining 4 being the ST height, P and R amplitudes, and the PR interval. The density distribution and symmetry in the arm regions of the attractor were the most discriminative SPAR features. The ST height—a known marker for Brugada syndrome—was selected in both feature groups that included amplitudes.

The authors concluded that the ECG carries sufficient information to detect the *Scn5a*+/- mutation, but that enhanced analysis techniques are required to extract it. Standard interval and amplitude measures fail to capture the relevant signal because the mutation's effects are distributed across the entire waveform morphology, not concentrated at isolated time points.

4. Structural Parallels with the Attractor Framework

4.1 The SPAR Attractor as a Cardiac Attractor. The SPAR method generates a bounded, stable 2-dimensional attractor from the ECG signal. This attractor is a compact representation of the cardiac system's dynamical state—a region in state space toward which trajectories converge and around which they organize. In the attractor framework's vocabulary, this is an **attractor** generated by a dissipative system (the beating

heart, maintained by continuous metabolic energy input). The attractor’s density distribution, arm structure, and symmetry reflect the stability and structural coherence of this configuration.

4.2 SPAR Features as Candidate Proxies for Corrective Permeability (κ). The framework proposes κ as a multi-dimensional measure of a system’s capacity to return to its attractor after perturbation. A healthy heart with normal ion channel function has a deep, stable attractor—it responds to perturbations and returns rapidly to its baseline rhythm. The Scn5a+/- mutation degrades sodium channel function, making the cardiac tissue more vulnerable to arrhythmia. This degradation manifests as measurable changes in the SPAR attractor.

A provisional mapping between specific SPAR feature categories and proposed components of κ is offered below. This mapping is hypothetical and has not been formally derived; it is presented as a structural analogy to be tested in future work. The κ component labels in this table are introduced here for exploratory purposes and are not yet formalized in the primary framework document (Galida, 2026a); they are subject to revision pending formal axiomatization of κ .

SPAR Feature Category	What It Measures in the Attractor	Candidate κ Component (provisional)
Density distribution (core)	Frequency of trajectory visits to central attractor region	Attractor core stability: a dense core indicates a stable, frequently occupied equilibrium
Density distribution (arms)	Frequency of trajectory visits to peripheral regions	Perturbation response: arm density reflects excursions from equilibrium

SPAR Feature Category	What It Measures in the Attractor	Candidate κ Component (provisional)
Symmetry features	Left-right symmetry of attractor arms	Recovery symmetry: asymmetric arms may indicate directional perturbation bias or conduction abnormality
Arm structure	Length, width, and number of attractor arms	Global waveform integrity: degraded arm structure reflects disrupted cardiac conduction

The 96% classification accuracy (pending independent replication) demonstrates that these attractor-derived proxies capture diagnostically relevant information that standard interval measures miss. Whether this information corresponds specifically to κ , or to more general signal properties, cannot be determined without a formal derivation of κ from the framework's axioms.

4.3 Multi-Dimensional Feature Combination. The framework proposes that κ is multi-dimensional—no single measure fully captures a system's corrective permeability. The SPAR results are consistent with this principle: combined features outperformed any individual feature set. However, this result is also expected under standard machine learning practice, where feature combination typically improves classification performance. The result is therefore consistent with the framework without uniquely supporting it. The specific finding that SPAR features (16/20) dominated the combined model suggests that attractor-derived measures carry more discriminative information than point-based measures for this particular mutation. Whether this dominance generalizes to other perturbations and other physiological systems is an open empirical question.

4.4 Normalization as Signal Isolation. The SPAR method

normalizes the signal to factor out changes in heart rate and baseline variation, concentrating on waveform morphology. In the framework's terms, this is a methodological step that isolates the attractor's structural properties from confounding variables. Heart rate is influenced by autonomic tone, physical activity, and respiratory cycle-perturbations that can obscure the measurement of the attractor's intrinsic stability. SPAR's normalization yields a cleaner representation of the attractor. However, this normalization step is standard practice in many signal processing methods and does not constitute a distinctive parallel with the framework.

5. Limitations

This mapping is post-hoc. The parallels identified here are structural analogies, not independent evidence for the framework. The provisional κ -proxy mapping in Section 4.2 is hypothetical and has not been formally derived from the framework's axioms. The κ component labels used in the provisional mapping table (e.g., "attractor core stability," "recovery symmetry," "global waveform integrity") are introduced in this paper for exploratory purposes and are not yet formalized in the primary framework document (Galida, 2026a). They are subject to revision pending formal axiomatization of κ .

The framework's κ remains qualitatively defined. A formal derivation specifying the state variables, the attractor geometry, and the perturbation response function is required before the SPAR feature mapping can be evaluated as more than a structural analogy.

The 96.2% classification accuracy was obtained from a single study of 42 mice (effective $N=42$, despite 1,014 windowed

records). Independent replication in a separate cohort has not been performed. The accuracy figure should be interpreted with appropriate caution.

The SPAR method was developed for approximately periodic signals and has been validated on cardiovascular waveforms. Its applicability to the non-periodic attractors the framework describes in cognitive and social domains is unknown.

The attractor framework is self-published and has not undergone independent peer review.

6. Falsifiability Conditions

The following observations would weaken or invalidate the parallels drawn here:

- **Disconfirming observation 1:** If SPAR features were shown to be *uncorrelated* with independently validated measures of cardiac resilience or arrhythmia susceptibility in a larger, independent cohort, the k proxy interpretation would lose its empirical anchor.
- **Disconfirming observation 2:** If the SPAR attractor's classification accuracy for the *Scn5a*^{+/-} mutation were shown to derive primarily from features unrelated to attractor geometry (e.g., heart rate alone or predominantly heart rate), the attractor interpretation would be substantially weakened.
- **Disconfirming observation 3:** If alternative signal processing methods with no attractor reconstruction component achieved equal or higher classification accuracy using the same data, the attractor interpretation would not be uniquely supported.

Affirmative predictions:

- **Primary prediction:** If the provisional κ -proxy mapping in Section 4.2 captures genuine components of corrective permeability, then pharmacological interventions that improve cardiac ion channel function (e.g., sodium channel modulators) should produce measurable shifts in specific SPAR features—density, symmetry, arm structure—toward the wild-type baseline. Conversely, interventions that degrade ion channel function should shift these features away from the baseline. This prediction is testable using pre- and post-intervention ECG recordings with the same SPAR methodology.
 - **Secondary prediction:** If attractor-derived features are more sensitive to κ -relevant perturbations than point-based measures, then SPAR features should show *greater* sensitivity to these pharmacological interventions than standard ECG intervals and amplitudes. This secondary claim is more speculative; failure of the secondary prediction while the primary prediction holds would suggest that SPAR features track relevant physiological changes without uniquely capturing κ as distinct from other measures.
-

7. Conclusion

The SPAR method developed by Bonet-Luz et al. (2020) generates a mathematically defined attractor from ECG signals that encodes diagnostically relevant information about cardiac stability. A provisional mapping between SPAR features and proposed components of corrective permeability (κ) has been offered, along with primary and secondary affirmative predictions. The 96% classification accuracy for a disease-causing mutation demonstrates that attractor-based features capture information about system integrity that standard point-based measures miss. These parallels are structural

analogies, not independent validation. The framework remains a self-published, preliminary research program. This mapping is a contribution to its ongoing development.

References

- Bonet-Luz, E., Lyle, J. V., Huang, C. L.-H., Zhang, Y., Nandi, M., Jeevaratnam, K., & Aston, P. J. (2020). Symmetric Projection Attractor Reconstruction analysis of murine electrocardiograms: Retrospective prediction of Scn5a+/- genetic mutation attributable to Brugada syndrome. *Heart Rhythm* 02, 1(5), 368–375. <https://doi.org/10.1016/j.hroo.2020.08.007>
- Galida, R. (2026a). *Persistence Under Perturbation: The Eternal Skeleton and the Transient Dance*. Fantasy Attractor. Published May 2026.