

The Attractor Framework as a Formal Mapping of Taoist Dynamics

R. S. Galida

Attractor Framework Research Program

Application Paper – June 13, 2026

For open peer review

Abstract

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

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

All claims are conditional on successful model-recovery

validation. The framework is therefore a coupled system of theory, measurement, simulation, and intervention logic.

1. Introduction

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

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

- **Conditional κ :** update gain under low-conflict uncertainty
- **B_depth:** energetic stability of an attractor
- **B_transition:** switching cost between attractors
- **E:** metabolic/computational effort per update (derived unless independently identified)

However, this paper does not assume these variables are empirically separable. Instead, it advances a **stability-plasticity axis hypothesis**, where all observed structure may collapse onto a single latent dimension. Whether κ , B_depth, and B_transition are separable constructs or projections of one axis is treated as an empirical identifiability problem.

2. Formal Hypothesis Mapping

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

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

3. Dimensionality Assumption: Stability–Plasticity Axis Hypothesis

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

Under this hypothesis:

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

derivative)

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

4. Expected Correlation Structure and Model Constraints

Under a single-axis model:

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

Under a multi-factor model:

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

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

5. Temporal Model Constraint

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

$-\beta_k t)$

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

6. Simulation-Based Identifiability Analysis

6.1 Generative Null Model (Single Axis)

A latent variable $z_t \sim \mathcal{N}(0,1)$ generates all observables: $k_t = a_1 z_t + \epsilon_k$

$B_{\text{depth},t} = a_2 (-z_t) + \epsilon_{B_d}$

$B_{\text{transition},t} = a_3 (-z_t) + \epsilon_{B_t}$

$E_t = a_4 (-z_t) + \epsilon_E$

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

6.2 Competing Models

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

6.3 Recovery Conditions

Validity of measurement inference requires:

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

6.4 Rotation Stability Test

All solutions are evaluated under:

- geomin rotation
- promax rotation

Instability is defined by:

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

6.5 Decision Rule

Empirical interpretation is valid only if simulation confirms:

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

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

7. Asymmetry of Convergence

Three regimes are distinguished:

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

8. Control Philosophy: Coercive Perturbation vs. Incremental Attractor Shaping (NEW)

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

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

8.1 Coercive perturbation

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

These often produce:

- rebound effects
- attractor deepening
- increased hysteresis

8.2 Incremental attractor shaping

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

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

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

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

Importantly, this distinction is not cultural or civilizational. It refers to two classes of control strategy over nonlinear systems:

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

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

- high identity coupling
- strong hysteresis
- long memory effects

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

9. Testable Predictions (Pre-Registered)

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

10. Limitations

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

11. Conclusion

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

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

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

References

- Galida, R. S. (2026a). *How to measure corrective permeability κ in a human belief system*. Attractor Framework Research Program.
- Galida, R. S. (2026b). *A multi-timescale latent variable model for attractor dynamics in belief systems*.
- Galida, R. S. (2026c). *Simulation-based identifiability analysis of attractor dimensionality*.
- Swann et al. (2009). Identity fusion and extreme group behavior.