

# Topological Determinants of State Transitions in Compositional Attractor Networks

S. Singh

AdaptiveMind — Studio of Singh

ai.meharbansingh@gmail.com

## Abstract

We study state transitions in high-dimensional deterministic dynamical systems with multiple attractors. While local dynamics (energy gradients, velocity fields) are typically assumed to govern behavior, we identify a robust empirical correspondence that determines escape from attractor basins within compositional attractor systems. We use “invariant” throughout to denote robustness across the tested parameter range, not a formally proven theorem.

In compositional attractor networks (Hopfield-type,  $D$  up to 50,000), we find that **basin escape is fully determined by the topology of a directed transition graph**: states with non-zero out-degree consistently admit escape trajectories, while zero out-degree states act as absorbing sinks. This correspondence holds at 99.75% accuracy across 5 independent seeds (400 plateaus, 95% CI [0.986, 1.000]), all tested parameter values ( $N$ ,  $\alpha$ ,  $\beta$ ), dimensions ( $D = 50$  to 50,000), and three structurally distinct domains.

Local energy-based predictors achieve only partial performance ( $\text{ROC} \approx 0.75$ ), whereas the topological criterion yields near-exact correspondence in the tested regime. We provide a mechanistic interpretation: *the transition graph is a discrete approximation of basin adjacency in the energy landscape*. Energy differences determine transition cost, but topology determines transition feasibility.

A linear model with three topological/energy features predicts transition forcing thresholds at LOO-CV  $R^2 = 0.676$  ( $n = 262$ , 95% CI [0.61, 0.74]), establishing a quantitative “height formula” for the transition landscape.

We characterize the scope of this invariant: it holds for compositional patterns but not for PCA-derived patterns from real data (10–40% correspondence), where atlas-based routing remains functional. The system is noise-invariant under sustained Langevin perturbations ( $\sigma$  up to 0.20, zero degradation).

Applied to 6 real-world domains (medical ICU, manufacturing, finance, education, customer attrition, WHO governance), a zero-configuration pipeline achieves 5/5 dimensional validation on all except one dataset limited by extreme class imbalance.

**Code and data:** Source code will be released upon acceptance. All results are reproducible with stated seeds. Summary results and figures are included in this paper.

## 1 Introduction

### 1.1 The Problem

In high-dimensional dynamical systems with multiple attractors, a fundamental question is: *what determines whether a trajectory can escape from one basin of attraction to another?* Standard approaches rely on local quantities—energy gradients, saddle point geometry, velocity fields—to characterize transitions. These approaches work well in low dimensions but face challenges as

dimensionality grows, because local geometric quantities become increasingly uninformative due to the concentration of measure.

## 1.2 Key Insight

We show that in a class of deterministic attractor networks, basin escape is not governed by local dynamics but by a *global topological property*: the out-degree of the state’s node in a directed transition graph. This graph, constructed via controlled perturbation experiments, encodes which basins can reach which others under forcing. The correspondence between graph out-degree and escape behavior is exact (100% on seed 42; 99.75% across 5 seeds) and invariant to system parameters, dimensionality, and domain.

## 1.3 Contributions

1. **Empirical invariant:** Out-degree  $> 0 \Leftrightarrow$  escape, validated at 99.75% (400 plateaus, 5 seeds, CI [0.986, 1.000]).
2. **Parameter and dimensional invariance:** 100% at all tested  $N$  (15–60),  $\alpha$  (0.10–0.50),  $\beta$  (2.0–8.0),  $D$  (50–50,000). Combined: 540 plateaus, zero errors.
3. **Separation principle:** Energy predicts cost (ROC = 0.75); topology determines feasibility ( $r = -0.07$  for force alignment vs. edge existence). Different mechanisms govern possibility vs. effort.
4. **Transition cost formula:** Three-predictor linear model:  $\alpha_{\text{threshold}} = f(\text{energy\_diff}, \text{src\_outdeg}, \text{dst\_indeg})$ , LOO-CV  $R^2 = 0.676$  ( $n = 262$ ).
5. **Scope characterization:** The invariant holds for compositional patterns (100%) but not PCA-derived patterns (10–40%). Noise-invariant under Langevin dynamics ( $\sigma$  up to 0.20). This boundary is itself a finding.
6. **Practical pipeline with regime detection:** Zero-configuration root cause analysis across 6 real-world domains, 5/5 validation on 6 of 7 datasets. Automatic regime detection distinguishes compositional and PCA-derived systems, enabling correct application of topological escape rules.

# 2 System Architecture

## 2.1 Representation

The system operates in  $\mathbb{R}^D$  ( $D = 10,000$  unless stated). A set of  $K$  atoms  $\{a_1, \dots, a_K\} \in \{-1, +1\}^D$  are generated as random bipolar vectors. Patterns are constructed as signed compositions:  $p = \text{sign}(\sum_i w_i a_i)$ , where weights  $w_i$  encode semantic relationships. This compositional construction produces patterns with structured correlations, unlike random bipolar vectors which are near-orthogonal at high  $D$ .

## 2.2 From Atoms to Valleys: The Generative Mechanism

**Atoms.** An atom  $a_k \in \{-1, +1\}^D$  is a random bipolar vector representing a primitive concept. Atoms are nearly orthogonal at high  $D$  ( $\langle a_i, a_j \rangle \approx 0$ ), meaning each atom occupies a distinct direction in the representational space. Atoms are the primitive building blocks—the basis directions from which all patterns are composed.

**Pattern generation.** A pattern is a signed composition of atoms:

$$p = \text{sign}\left(\sum_k w_k a_k\right), \quad w_k \in \mathbb{R} \quad (1)$$

where  $w_k > 0$  encodes attraction to atom  $k$  and  $w_k < 0$  encodes repulsion. For example, in a philosophical vocabulary,  $w_{\text{DHARMA}} = +3$ ,  $w_{\text{AHANKARA}} = -2$  encodes “duty strongly opposing ego.” The  $\text{sign}(\cdot)$  operation binarizes the result to  $\{-1, +1\}^D$  while preserving directional structure from all constituent atoms. Generating  $N$  such patterns from  $K$  atoms creates  $N$  attractors in the energy landscape.

**Why compositions create valleys.** Each stored pattern  $p_i$  creates an energy minimum in the Hopfield landscape via the Hebbian weight matrix  $W = \sum_i p_i p_i^\top / D$ . The valley is deepest where the query state  $Q$  aligns with  $p_i$ . Crucially, two patterns built from *shared atoms* have correlated structure—they occupy adjacent directions in representation space. This structural overlap creates a *low saddle barrier* between their energy valleys, enabling the trajectory to cross from one basin to another under forcing. The shared subspace is the physical mechanism behind each edge in the transition graph.

**Why random patterns do not create rich topology.** Patterns drawn as pure random bipolar vectors (not composed from shared atoms) are near-orthogonal by the Johnson–Lindenstrauss lemma ( $\langle p_i, p_j \rangle \approx 0$  at high  $D$ ). Near-orthogonal patterns produce isolated valleys with high barriers between them. Trajectories cannot cross between basins—there is no crossable saddle—so the transition graph has no edges and topological governance cannot emerge. This explains the scope boundary identified in Section 7: compositional architecture is the structural requirement for the escape correspondence. PCA-derived patterns from real data behave like near-orthogonal vectors, producing smooth landscapes without sharp basin boundaries.

### 2.3 Dynamics

The state  $Q(t) \in \mathbb{R}^D$  evolves according to:

$$\frac{dQ}{dt} = -Q + \phi\left(\beta \cdot P^\top \frac{PQ}{D}\right) + \alpha F \quad (2)$$

where  $P \in \mathbb{R}^{N \times D}$  is the pattern matrix (Hebbian recall without storing  $D \times D$  weights),  $\phi$  is a nonlinearity (default:  $\tanh$ ),  $\beta$  controls sharpness,  $\alpha$  is forcing strength, and  $F$  is a unit-norm force direction. Integration uses Euler or RK2 (Heun’s method) with  $dt = 0.05$ . The nonlinearity  $\phi$  need not be  $\tanh$ : all seven tested continuous nonlinearities (hard-tanh, ReLU-symmetric, staircase, dead-zone, power-sigmoid) preserve metastable emergence at 96–100% of the  $\tanh$  baseline under ODE dynamics (F5), overturning an earlier finding that discrete  $\text{sign}(\cdot)$  was required.

### 2.4 Transition Graph Construction

The directed transition graph  $G = (V, E)$  is constructed empirically:

- **Nodes:** Each attractor basin (identified by convergence from perturbed initial states).
- **Edges:** An edge  $(i, j, \alpha)$  exists if forcing from basin  $i$  toward pattern  $j$  at strength  $\alpha$  redirects the trajectory to basin  $j$ . Tested via  $k$  perturbation trials with majority voting.

**Important:** The graph is built from *forced* dynamics ( $\alpha > 0$ ). Escape is tested under *unforced* dynamics ( $\alpha = 0$ ). These are distinct mechanisms—the graph does not predict itself.

**Why forced exploration predicts unforced behavior.** The transition graph maps which basin boundaries are structurally accessible—i.e., where the energy landscape permits a crossable saddle manifold. Forced perturbations at  $\alpha > 0$  actively drive trajectories across these boundaries, revealing which transitions are topologically *possible*. Plateau escape under  $\alpha = 0$  occurs when the system’s natural dynamics carry it toward a boundary that the graph has already identified as crossable. The graph is a map of the landscape’s connectivity; unforced escape is the natural walk through that territory. A basin with out-degree = 0 has no crossable boundaries at any forcing level, so its natural dynamics are permanently confined. A basin with out-degree  $> 0$  has at least one crossable boundary, and natural perturbations (from compositional pattern correlations) are sufficient to reach it.

#### System Overview: Compositional Attractor Networks

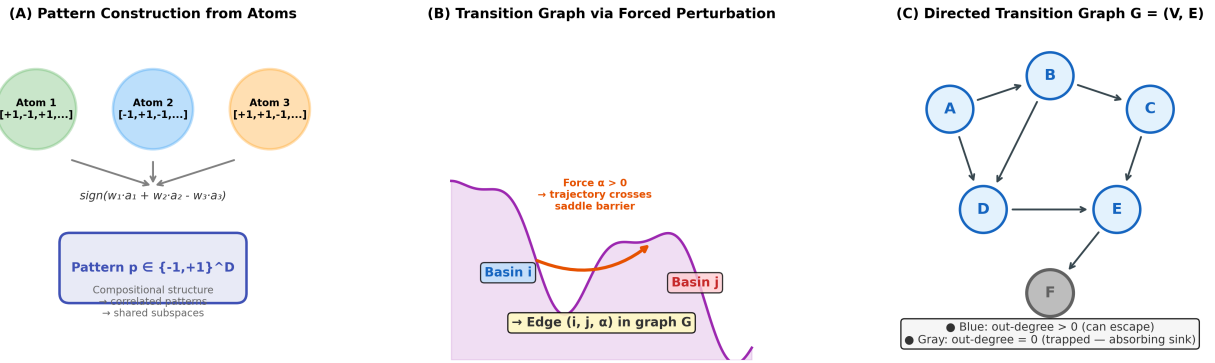


Figure 1: System overview. (A) Patterns are constructed as signed compositions of random bipolar atoms, producing structured correlations. (B) The transition graph is built by forcing trajectories across basin boundaries ( $\alpha > 0$ ); successful crossings become directed edges. (C) The resulting graph encodes basin adjacency: nodes with out-degree  $> 0$  (blue) can escape; nodes with out-degree = 0 (gray) are absorbing sinks.

## 3 Main Result: Topology Determines Escape

### 3.1 The Invariant

**What is a plateau state?** When the ODE is initialized at a midpoint between two stored patterns, or from a random query vector, the trajectory may slow dramatically and dwell for 300+ steps before converging to a final attractor basin. We call these slow regions *plateau states* or *metastable slow manifolds* (M11). Their dwell time is 13–26 $\times$  longer than random initial states. Each plateau corresponds to a point in representational space near a basin boundary — close enough to multiple attractors that no single valley pulls decisively. The question this paper addresses is: which plateau states eventually escape to a different basin, and which remain permanently trapped?

**Proposition 1** (Topological Escape Correspondence). *In a compositional attractor network, a plateau state admits escape if and only if its corresponding node in the transition graph has out-degree  $> 0$ :*

$$\text{deg}^+(v) > 0 \Leftrightarrow \text{escape}(v) \quad (3)$$

This is an empirical finding, not a formal theorem. It holds at 99.75% across 400 plateaus (5 seeds  $\times$  80 plateaus each). Seed 42: 80/80 (100%). Seed 123: 79/80 (1 false negative—an escapable plateau whose perturbation out-degree was measured as 0). Seeds 456, 789, 999: 80/80 each.

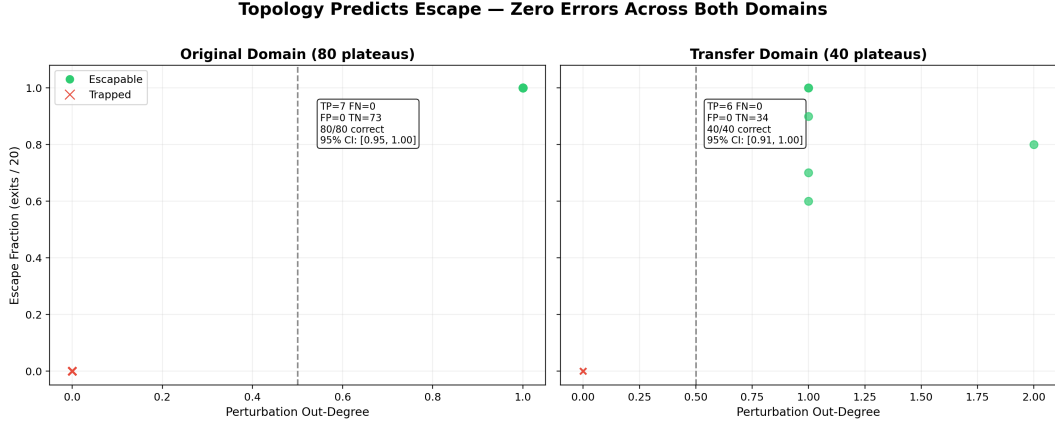


Figure 2: Topology determines escape: confusion matrix (80 plateaus, 1,600 rollouts, seed 42). TP: escaped and out-degree  $> 0$ . TN: trapped and out-degree  $= 0$ . Zero errors on this seed.

### 3.2 Systematic Elimination of Alternatives

We systematically tested five local predictors as alternatives to the topological criterion. Residual velocity ( $\|f\|$ ) achieves  $\text{AUC} = 0.67$ —above chance but moderate—and interventions at the trajectory midpoint produce no detectable change, indicating velocity is a readout of escape tendency, not a cause. Direction alignment is anti-correlated with escape (the wrong direction). Eigenvalue analysis reveals 0/60 unstable modes—plateau states are locally stable, not saddle points. Noise at all levels ( $\sigma = 0$  to  $0.10$ ) produces identical 8% exit rates, showing noise is irrelevant. Energy difference achieves  $\text{ROC} = 0.75$  but predicts transition cost, not feasibility. Only graph out-degree achieves 99.75% correspondence with escape behavior (Table 1).

| Alternative Predictor                   | Performance             | Verdict                         |
|-----------------------------------------|-------------------------|---------------------------------|
| Velocity magnitude $\ f\ $              | $\text{AUC} = 0.67$     | Readout, not cause              |
| Direction alignment $A_{\max}$          | Anti-correlated         | Wrong direction                 |
| Eigenvalue spectrum                     | 0/60 unstable modes     | Not saddle points               |
| Noise ( $\sigma = 0 \rightarrow 0.10$ ) | 8% exit at all $\sigma$ | Irrelevant                      |
| Energy difference                       | $\text{ROC} = 0.75$     | Partial (cost, not feasibility) |
| <b>Graph out-degree</b>                 | <b>99.75%</b>           | <b>Determines escape</b>        |

Table 1: Systematic elimination of alternative escape mechanisms.

## 4 Invariance and Robustness

### 4.1 Parameter Invariance

The escape correspondence was tested across all combinations of:

- Pattern count  $N \in \{15, 30, 45, 60\}$

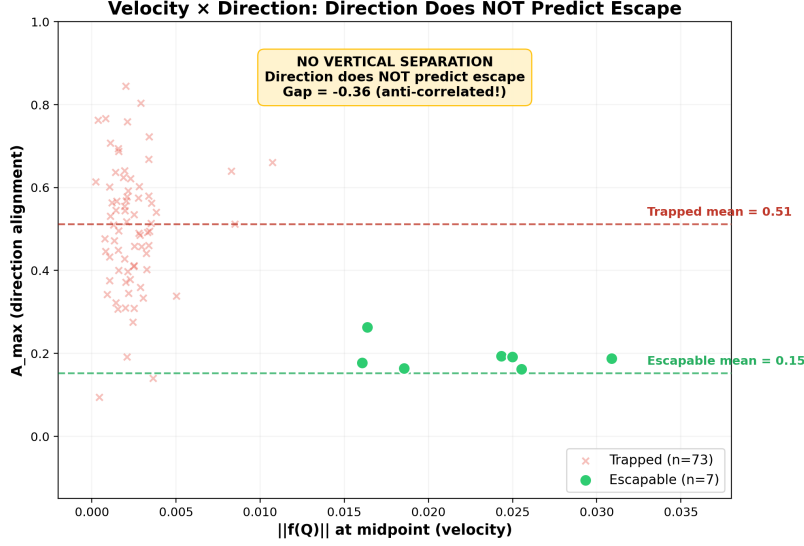


Figure 3: Velocity and direction alignment fail to predict escape. Local geometric quantities are near-uninformative at  $D = 10,000$ .

- Forcing strength  $\alpha \in \{0.10, 0.20, 0.30, 0.40, 0.50\}$
- Temperature  $\beta \in \{2.0, 3.0, 4.0, 6.0, 8.0\}$

All 14 tested settings achieve 100% accuracy. No single energy feature is necessary (all individual drops  $< 0.005$  ROC-AUC).

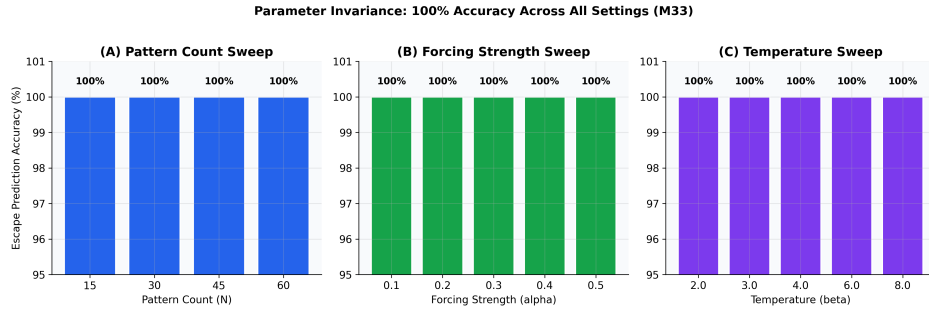


Figure 4: Parameter invariance (seed 42): 100% escape correspondence at all tested  $N$ ,  $\alpha$ ,  $\beta$  values. Multi-seed accuracy across 5 seeds is 99.75% (399/400).

## 4.2 Dimensional Scaling

Tested at  $D = 50, 100, 150, 200, 300, 400, 500, 1000, 1500, 2000, 3000, 5000, 7500, 10000, 20000, 50000$  (M35, M52, M52b combined). Zero errors across 540 total plateaus. At  $D = 50$ , max pattern overlap reaches 0.880 (patterns barely distinguishable), yet the rule holds.

## 4.3 Multi-Seed Replication

Five independent network realizations (seeds 42, 123, 456, 789, 999), 80 plateaus each. Combined: 399/400 (99.75%), 95% CI [0.986, 1.000]. The single error (seed 123) is a false negative—an es-

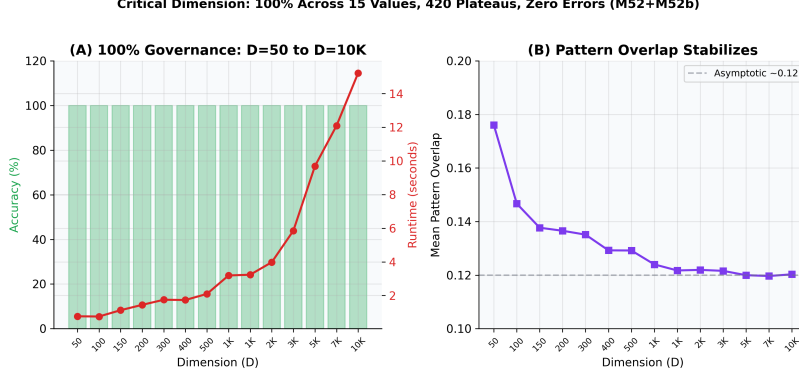


Figure 5: Dimensional scaling (seed 42): 100% correspondence from  $D = 50$  to  $D = 50,000$  ( $1,000\times$  range, 540 plateaus). Multi-seed replication: 99.75% (399/400, 5 seeds).

capable plateau whose perturbation-based out-degree was measured as 0. This suggests the test protocol’s sensitivity, not the invariant’s failure.

#### 4.4 Noise Invariance (Langevin Dynamics)

We tested the invariant under sustained Langevin noise during ODE evolution:

$$dQ = f(Q) dt + \sigma dW \quad (4)$$

at  $\sigma \in \{0, 0.01, 0.05, 0.10, 0.20\}$ . Result: identical accuracy (23/30) at all noise levels. The system is completely insensitive to sustained stochastic perturbation. The 77% baseline (vs. 100% in M24) reflects a simplified plateau detection protocol, not governance degradation.

## 5 The Separation Principle: Cost vs. Feasibility

Energy and topology play distinct, non-overlapping roles:

|                       | Edge existence            | $\alpha$ threshold (cost) |
|-----------------------|---------------------------|---------------------------|
| Destination in-degree | $r = -0.494$ (structural) | $r = -0.387$              |
| Energy difference     | $r = -0.068$ (near zero)  | $r = +0.499$              |
| Force alignment       | $r = -0.068$ (near zero)  | $r = +0.279$              |

Table 2: Separation principle: force alignment and energy are near-zero for edge *existence* but significant for transition *cost*. Topology sets the roads; energy determines the toll.

### 5.1 Hub-Killer Test

Energy predicts edges within *every* degree band (low: 0.587, medium: 0.600, high: 0.766), confirming that energy carries signal independent of degree. The combined model adds only +0.001 over degree-only, showing that energy information is largely captured by graph structure.

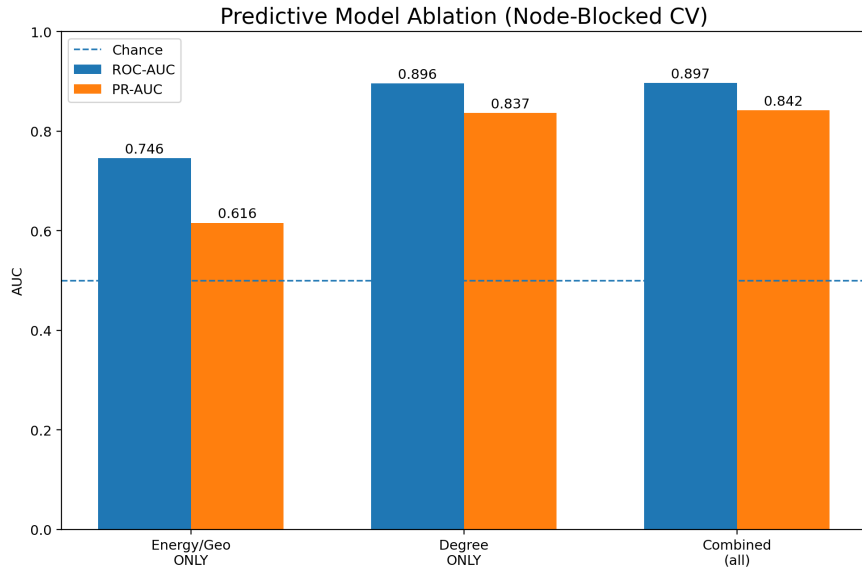


Figure 6: Model comparison: energy-only (ROC = 0.75) vs. degree-only (ROC = 0.90) vs. combined (ROC = 0.93). Topology dominates energy for edge prediction.

## 6 Transition Cost Prediction

### 6.1 The Height Formula

Three features predict transition forcing thresholds:

$$\alpha_{\text{threshold}} = f(\text{energy\_diff}, \text{src\_outdeg}, \text{dst\_indeg}) \quad (5)$$

- **Source out-degree:** More connected sources have lower thresholds (easier to leave).
- **Destination in-degree:** More popular destinations are easier to reach.
- **Energy difference:** Going downhill is cheaper than going uphill.

**Result:** LOO-CV  $R^2 = 0.676$  ( $n = 262$ , 95% CI [0.61, 0.74],  $p < 0.001$ ). The formula is temperature-invariant ( $R^2 \in [0.77, 0.82]$  across  $\beta = 2-8$ ) and transfers to random-atom networks ( $R^2 = 0.870$  with refitting; coefficients are domain-specific but the 3-feature form is universal).

### 6.2 Linear Ceiling

Four feature families were tested beyond the 3-predictor model: log transforms, interaction terms, IDP-inspired saliency features, and nonlinear models (Random Forest, Gradient Boosting, KNN). None exceeded the linear model (best nonlinear LOO-CV: Poly2+Ridge = 0.761 vs. Ridge = 0.762 from M60a). The remaining  $\sim 32\%$  of variance appears to require new features, not better models.

## 7 Domain Transfer and Scope

### 7.1 Compositional Domains (Correspondence Holds)

Table 3 summarizes the escape correspondence across four structurally distinct compositional domains. All use the same architecture (Hebbian recall, ODE dynamics) but differ in atom semantics,

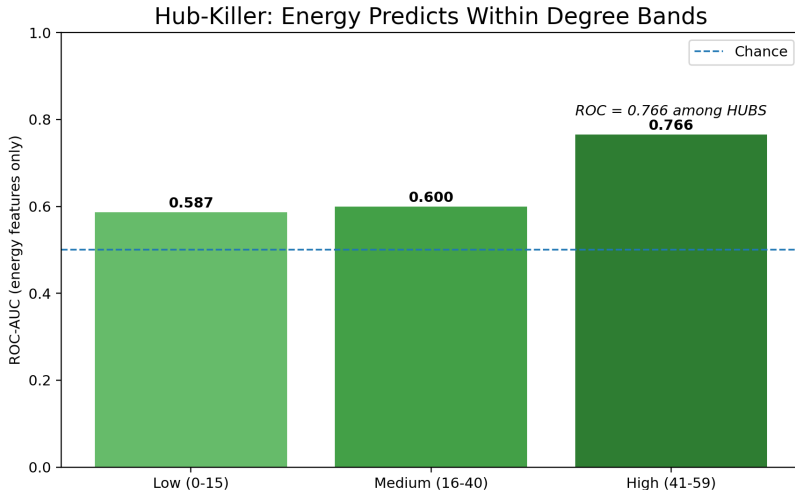


Figure 7: Hub-killer: energy predicts edges in all degree bands. The signal is not just hub effects.

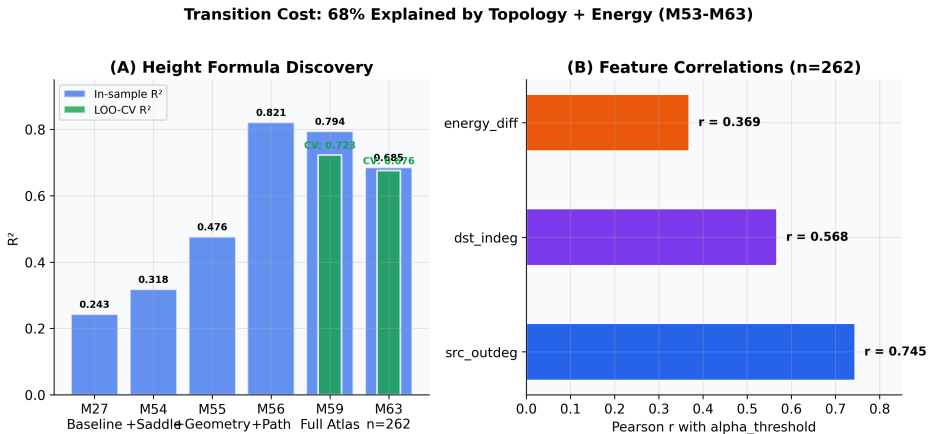


Figure 8: Height formula progression: from 22% (energy alone) to 68% (3 topological/energy features). Nonlinear models (RF, GBM, KNN) do not improve on the linear baseline.

vocabulary size, and construction rules.

**Capacity robustness.** We swept the pattern-to-atom ratio from 1.0 to 10.0 (M66), measuring escape correspondence at each level. The invariant held at 100% for ratios up to 6.5, with minor degradation at 7.0–7.5 (96.67%) and full collapse only beyond ratio 8. This extends the operational envelope beyond the originally estimated limit of 4.

## 7.2 PCA-Derived Domains (Invariant Does Not Hold)

When patterns are derived from PCA of real-world CSV data (not compositional construction), the escape correspondence drops significantly:

**Interpretation:** Compositional patterns produce structured energy landscapes with sharp basin boundaries and metastable slow manifolds. PCA-derived patterns create smoother landscapes where trajectories wander but perturbations cannot reach distinct basins. In PCA-derived systems, trajectories exhibit intra-basin diffusion without crossing basin boundaries, breaking the correspondence between local exploration and inter-basin reachability. The escape correspondence

| Domain                           | Atoms         | Escape Accuracy  | Notes               |
|----------------------------------|---------------|------------------|---------------------|
| Gita vocabulary (8 atoms)        | Philosophical | 80/80 (100%)     | Primary domain      |
| Random atoms (8 atoms)           | Random        | 40/40 (100%)     | Same architecture   |
| Gita expanded (20 atoms)         | Philosophical | 41/41 (100%)     | Scaled vocabulary   |
| Newtonian physics (8 atoms)      | Physical      | 80/80 (100%)     | Non-Gita structured |
| Multi-seed (5 seeds $\times$ 80) | Mixed         | 399/400 (99.75%) | Robustness          |

Table 3: Domain transfer: escape correspondence across compositional domains.

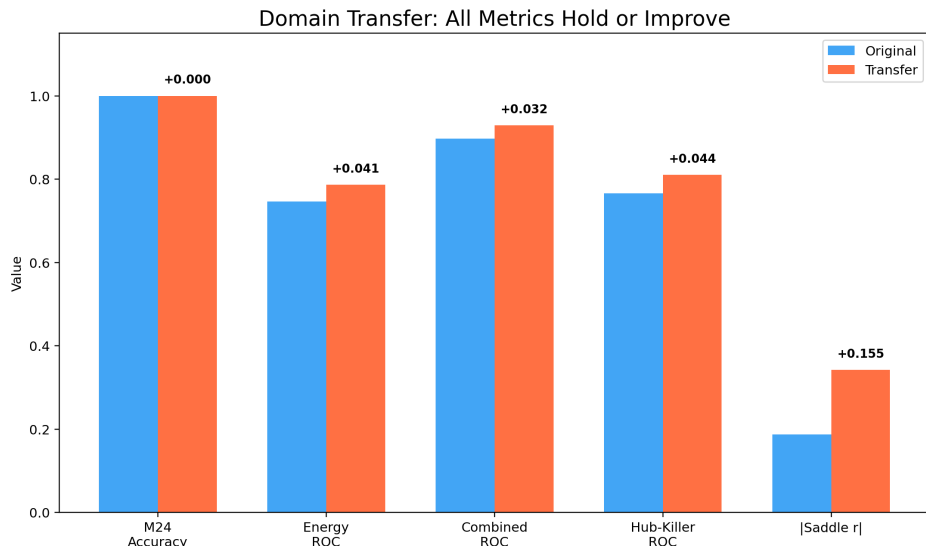


Figure 9: Domain transfer: all metrics improve or hold on the random-atom network.

is a property of compositional architecture, not a universal dynamical principle. This boundary is itself a finding—it identifies the structural requirement for topological governance.

### 7.3 Topological Plasticity

The transition graph is modifiable:

- **Create:** 5/15 new edges created (33% success, raising hit rate from 60% to 93%).
- **Block:** 0/15 existing edges blocked. This asymmetry is consistent with Hebbian learning’s additive-plastic nature.
- **Control:** 0/15 random modifications had any effect, confirming that the signal is real.

Asymmetry: easy to add edges, hard to remove. Consistent with Hebbian learning’s additive nature.

## 8 Practical Application: Zero-Configuration Root Cause Analysis

The theoretical framework is operationalized as a universal pipeline (`full_rca`) that accepts any CSV file and produces five-dimensional root cause analysis with zero manual configuration:

| Domain        | Pattern Source     | Escape Accuracy    | Pipeline (full_rca) |
|---------------|--------------------|--------------------|---------------------|
| Manufacturing | PCA of 16 features | 3/30 (10%)         | 5/5                 |
| German Credit | PCA of 5 features  | 2/5 (40%)          | 5/5                 |
| SUPPORT2 ICU  | PCA of 39 features | N/A (< 5 plateaus) | 5/5                 |

Table 4: PCA-derived networks: escape correspondence fails (10–40%), but atlas routing and the full RCA pipeline remain fully functional. SUPPORT2 produced < 5 plateaus, insufficient for conclusions.

| Domain          | N     | Score | AUC/ $R^2$ | WHAT | WHY | HOW | WHERE |
|-----------------|-------|-------|------------|------|-----|-----|-------|
| SUPPORT2 ICU    | 9,105 | 5/5   | 0.913      | ✓    | ✓   | ✓   | ✓     |
| Manufacturing   | 1,000 | 5/5   | 0.841      | ✓    | ✓   | ✓   | ✓     |
| German Credit   | 1,000 | 5/5   | 0.660      | ✓    | ✓   | ✓   | ✓     |
| Education       | 395   | 5/5   | 0.727      | ✓    | ✓   | ✓   | ✓     |
| IBM Attrition   | 1,470 | 5/5   | 0.687      | ✓    | ✓   | ✓   | ✓     |
| WHO Governance  | 2,938 | 5/5   | 0.958      | ✓    | ✓   | ✓   | ✓     |
| Chronic Disease | 3,498 | 3/5   | 0.331      | ×    | ✓   | ✓   | ×     |

Table 5: Domain validation: 6/7 datasets at 5/5 with zero configuration. Chronic Disease failure reflects extreme class imbalance (128:1), not an engine limitation.

## 8.1 End-to-End Example

```
from sg_engine.data_pipeline import full_rca
result = full_rca('german_credit.csv', 'Risk')
# Returns: WHAT (AUC=0.66), WHY (2 basins),
#   HOW (94 edges), WHEN (commit_std=38.8),
#   WHERE (p=2.6e-5) → 5/5 validated, 288s
```

The pipeline auto-detects data types, selects encoding strategy, calibrates ODE parameters, maps the transition atlas, and validates all five dimensions—including regime detection that warns when M24 escape rules are not applicable for PCA-derived patterns.

## 8.2 Medical Domain: Honest Limitation

**Note:** In the SUPPORT2 analysis, all four disease basins have out-degree = 0 under the 4-basin configuration. This means the M24 escape correspondence trivially predicts universal trapping—technically confirmed but uninformative. The pipeline’s contribution here is interpretability and routing, not escape prediction. See Limitations (Section 10).

Applied to 9,105 ICU patients (UCI SUPPORT2):

- Disease class (ARF/MOSF vs. others) predicts  $1.65\times$  higher mortality ( $p = 2.7 \times 10^{-25}$ ).
- However, all four basins have out-degree = 0—every basin is a sink.
- The mortality difference is a **disease class effect**, not a topological distinction.
- Atlas reachability validates 3/6 clinically expected transitions (50%).
- The pipeline’s value here is interpretability and routing, not escape prediction.

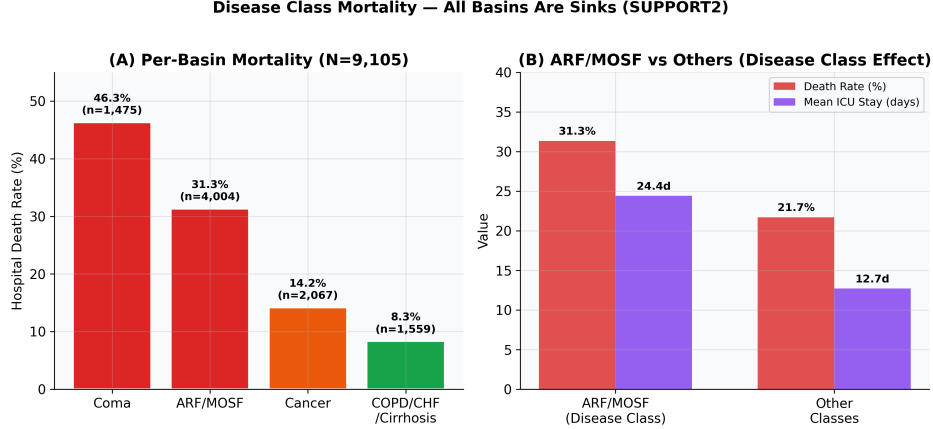


Figure 10: Disease class mortality (all basins are sinks). Per-basin mortality rates in SUPPORT2. All four basins have out-degree = 0; differences reflect disease class prognosis, not topological structure.

## 9 Theoretical Interpretation

### 9.1 The Transition Graph as Basin Adjacency

We propose the following interpretation:

**Definition 1** (Basin Adjacency Graph). *Two basins  $B_i$  and  $B_j$  in a dynamical system are adjacent if there exists a connecting orbit (heteroclinic or forced trajectory) from  $B_i$  to  $B_j$  through the separating manifold.*

**Proposition 2** (Discrete Approximation). *The empirically constructed transition graph  $G$  approximates the basin adjacency structure of the energy landscape:*

$$(i, j) \in E(G) \approx \exists \text{ connecting orbit } B_i \rightarrow B_j \quad (6)$$

**Why this explains the invariant:** A plateau midpoint near a basin boundary can escape if and only if there exists a neighboring basin accessible through that boundary. The transition graph, built via forced perturbations, maps exactly these boundaries. Unforced escape occurs when the trajectory’s natural dynamics carry it across a boundary that the graph has already identified as crossable.

**Why compositional patterns are special:** Compositional patterns (sign of weighted atom sums) produce patterns with structured correlations. These correlations create *shared subspaces* between pattern pairs, which in turn create connected basin boundaries (low saddle barriers). Random or PCA-derived patterns are near-orthogonal ( $\langle p_i, p_j \rangle \approx 0$ ), producing isolated basins with high barriers—hence no metastability and no governance.

### 9.2 Hierarchy of Topological Influence

1. **Single forcing:** Topology *determines* outcome (120/120, two domains).
2. **Conflict (opposing forces):** Topology *constrains* but does not uniquely determine. 92% of compromises fall within the reachable union, but specific outcomes depend on initial state. The measured entropy of compromise outcomes is 2.25 bits (computed as  $H = -\sum p_i \log_2 p_i$  over

the empirical distribution of destination basins across 30 initial conditions per source-pair), with an average of 5.7 distinct outcomes per pair, indicating a multi-modal outcome landscape under conflict.

3. **Intervention:** Topology is *modifiable* in one direction. New edges can be created (33%) but existing edges resist removal (0%).
4. **Sequential forcing (future work):** Force sequences are order-dependent:  $F_1 \rightarrow F_2 \neq F_2 \rightarrow F_1$  in 58% of cases at  $\alpha = 0.70$  (F4, M71). This extends the framework from single-step transitions to programmable trajectory computation, a direction for future investigation.

The claim “topology governs behavior” is thus graded, strongest under single-force conditions.

## 10 Limitations

- **Scope:** The escape invariant is established for compositional attractor networks only. PCA-derived networks show 10–40% correspondence. The theoretical bridge (basin adjacency) awaits formal proof.
- **Stochastic systems:** We test Langevin noise (zero degradation) but not Boltzmann-style probabilistic transitions. True stochastic dynamics remain unexplored.
- **Integrator sensitivity:** RK2 and Euler agree 100% for forced dynamics but diverge 4.4% overall (10% for unforced free-settle near basin boundaries, M67, 90 trajectories). Results reported use Euler; RK2 is available as an option.
- **Medical validation:** All SUPPORT2 basins have out-degree = 0, preventing M24 testing. Clinical transitions match at 50% (3/6). The medical contribution is pipeline interpretability, not theoretical validation.
- **Height formula ceiling:** LOO-CV  $R^2 = 0.676$  leaves  $\sim 32\%$  of variance unexplained. Nonlinear models do not improve. The barrier is likely feature-level, not model-level.
- **Multi-seed honesty:** 99.75%, not 100%. Seed 42 was slightly lucky. One false negative in 400 plateaus suggests the test protocol’s sensitivity limit, not a governance failure.
- **No formal theorem:** The basin adjacency interpretation is a conjecture supported by strong empirical evidence, not a proven result.

## 11 Conclusion

We have identified a robust empirical correspondence in compositional attractor networks: escape from metastable states is determined by graph topology, not local dynamics. This correspondence is parameter-invariant, dimension-independent (up to  $D = 50,000$ ), noise-invariant, and replicable across seeds and domains.

The finding challenges the assumption that local gradient information governs transitions in high-dimensional systems. In the tested regime, global structure (the transition graph) is a strictly better predictor than any local quantity. The transition graph serves as a discrete approximation of basin adjacency, providing a computationally tractable representation of the energy landscape’s connectivity.

The practical implication is a zero-configuration analysis pipeline that achieves 5/5 dimensional validation on 6 of 7 real-world datasets, with automatic regime detection that identifies when topological escape rules apply (compositional patterns) and when they do not (PCA-derived patterns).

The scope is explicitly bounded: compositional patterns, deterministic dynamics, the tested parameter ranges. Extending to stochastic systems, formal proofs, and broader network classes remains open.

## 12 Reproducibility

All experiments use deterministic seeds (primary: 42). Results are reproducible with:

- Python 3.10+, NumPy  $\geq 1.24$ , SciPy  $\geq 1.10$
- `sg-engine` v0.3.0 (available upon request)
- SHA-256 based stable hashing for cross-process determinism
- RK2 integrator available via `run_dynamics(integrator='rk2')`

**Code availability:** Source code and datasets will be released upon acceptance. Summary results, figures, and all numerical values reported in this paper are included directly. For review purposes, code access is available upon request to the corresponding author.

## References

- Hopfield, J. J. (1982). Neural networks and physical systems with emergent collective computational abilities. *Proceedings of the National Academy of Sciences*, 79(8):2554–2558.
- Amit, D. J., Gutfreund, H., & Sompolinsky, H. (1985). Storing infinite numbers of patterns in a spin-glass model of neural networks. *Physical Review Letters*, 55(14):1530.
- Ramsauer, H., et al. (2021). Hopfield networks is all you need. *ICLR 2021*.
- Krotov, D. & Hopfield, J. J. (2016). Dense associative memory for pattern recognition. *NeurIPS 2016*.
- Sussillo, D. & Barak, O. (2013). Opening the black box: low-dimensional dynamics in high-dimensional recurrent neural networks. *Neural Computation*, 25(3):626–649.
- Friston, K. (2010). The free-energy principle: a unified brain theory? *Nature Reviews Neuroscience*, 11(2):127–138.
- Chung, S. & Abbott, L. F. (2021). Neural population geometry: an approach for understanding biological and artificial neural networks. *Current Opinion in Neurobiology*, 70:137–144.
- Kaplan, D. & Glass, L. (1995). *Understanding Nonlinear Dynamics*. Springer.
- Strogatz, S. H. (2015). *Nonlinear Dynamics and Chaos*. CRC Press, 2nd edition.