

The Coastline of Predictability: Coherent Multi-Scale Measurement of Surveillance Power

Larsen James Close

February 2026

Abstract

The predictive power that a data collection gives an observer over a target has no rigorous multi-scale measure — existing scalars (mutual information, fractal dimension, information dimension) depend on the observer’s prediction target Z , not only on the system. We introduce the *predictability coastline* $C(\epsilon) = I(Z; \mathcal{F}_\epsilon)$, which traces how predictive capacity scales with data resolution via an information-theoretic filtration, and show that $C(\epsilon)$ is a property of the (system, observer, Z) triple. To isolate system-intrinsic information, we define the *coherent coastline* $C_{\text{coh}}(\epsilon) = \min_k \text{NMI}^*(Z_k; \mathcal{F}_\epsilon)$ (NMI^* defined in Section 9.5), a min-envelope over diverse prediction targets that strips measurement artifacts. Across six systems — Lorenz, Hénon, SPY, TLT, GLD, and a thermostat baseline — the coherent coastline produces a clean three-tier separation: chaotic attractors ($C_{\text{coh}}^{\text{max}} \approx 0.28\text{--}0.40$), financial markets ($0.04\text{--}0.14$), and structureless noise (0.03). The *fragility ratio* $\mathcal{R}_{\text{frag}}$ quantifies how much predictability vanishes when the prediction target changes: the Lorenz attractor retains $\approx 41\%$ ($\mathcal{R}_{\text{frag}} = 1.42$); the thermostat loses 97% ($\mathcal{R}_{\text{frag}} = 29.6$). Block-bootstrap resampling distinguishes attractor-stable coastlines from correlation-fragile ones. We formalize the *capture threshold* — the resolution at which an observer’s model exceeds the target’s self-model — and show it arises from kernel asymmetry (the observer retaining variables the self-model projects away), not from resolution depth. Bridge-connectivity-targeted data removal is $7.7\times$ more effective than uniform data minimization. The framework’s value lies in the coastline’s shape, not in any extracted scalar.

Contents

1	The Coastline Paradox for Prediction	2
2	Formal Framework	3
2.1	Base Objects	3
2.2	Data Channels and the Measuring Stick	3
2.3	The Predictability Coastline	3
2.4	Fractal Dimension as Scaling Exponent	4
2.5	Phase Transitions	4
3	Topological Unlock Events	5
3.1	From Filtration to Topology	5
3.2	Betti Number Interpretation	5
3.3	Topological Unlock Criterion	5
4	The Capture Threshold	6
4.1	Definition	6
4.2	Kernel Asymmetry as the Mechanism of Capture	6
4.3	Structural Identity with Control Inversion	7
4.4	Structural Identity with Sub-Threshold Grammatical Capture	7
5	Bidirectional Temporal Prediction	7
5.1	The Temporal Symmetry	7
5.2	Retroactive Inference	7

5.3	Temporal Asymmetry in the Coastline	8
6	Defense Topology	8
6.1	The Kernel Principle	8
6.2	Bridge Connectivity and Data Hygiene	8
6.3	Defense Scaling	9
7	Domain Instantiations	9
7.1	Individual Surveillance	9
7.2	Financial Markets	9
7.3	Climate Systems	10
8	Connection to the Bottleneck Primitive	10
9	Empirical Validation	11
9.1	Systems Tested	11
9.2	Experiment 1: Defense Topology	11
9.3	Experiment 2: Self-Model and Capture Threshold	12
9.4	Experiment 3: Information Dimension Across Domains	12
9.5	Experiment 4: Coherent Coastline (CMC Resolution)	13
9.6	Summary of Empirical Status	16
10	The Flagship Theorem	16
11	Scope and Limitations	17
11.1	What the Framework Does Not Cover	17
11.2	Z-Dependence and the Scalar Reduction Problem	17
11.3	Ethical Considerations	18
	References	18

1 The Coastline Paradox for Prediction

Mandelbrot observed that the measured length of Britain’s coastline depends on the length of the measuring stick: shorter sticks reveal finer indentations, and the total length diverges rather than converging to a fixed value. The scaling relationship between stick length and measured length defines the fractal dimension of the coast — a non-integer quantity encoding how much new structure appears at each scale.

We observe the same phenomenon in predictive modeling. A threat actor’s capacity to model a target’s behavior depends on data resolution — the granularity and diversity of available telemetry. As resolution increases, the “perimeter” of modelable behavior does not converge. It expands, revealing qualitatively new behavioral structures rather than merely refining existing ones. GPS data reveals spatial routines. Adding purchase records reveals lifestyle patterns — a new axis, not a refinement of the spatial one. Adding communication metadata reveals social topology. Adding content analysis reveals belief structure. Adding biometrics reveals physiological state. Each data type does not sharpen existing predictions; it opens entirely new dimensions of prediction.

This is not a metaphor. We will formalize it as a fractal scaling law with measurable dimension, phase transitions at critical data thresholds, and a computable criterion for when the observer’s model exceeds the target’s capacity for self-knowledge.

The framework rests on a foundation established in a companion paper (Close 2026a, “The Bottleneck Primitive”): every epistemic act is a projection through an information bottleneck, and every epistemic result is jointly a property of the data and the projection geometry. The present paper asks: *how does the projection’s resolving power scale with aperture width, and what happens when the projector’s resolution exceeds the target’s self-model resolution?*

2 Formal Framework

2.1 Base Objects

Let the target be a stochastic process X_t on a Polish metric space $(\mathcal{X}, d)$, representing the full state of a system (an individual, a market, a climate field) evolving in time.

Fix a task variable Z representing what the observer seeks to predict or reconstruct:

- **Forward prediction:** $Z^+ = g(X_{t+\tau})$ for some horizon $\tau > 0$.
- **Backward reconstruction:** $Z^- = h(X_{t-\tau})$ — inferring past states from present data.
- **Structural inference:** Z^s — latent variables such as intent, belief state, social role, or causal structure.

The task variable remains fixed throughout. The framework applies to any Z ; changing Z changes the coastline but not the formalism.

2.2 Data Channels and the Measuring Stick

Let channels be indexed by $k \in \{1, \dots, K\}$. Each channel is a measurement map with quantization and noise:

$$Y^{(k)} = Q_\varepsilon(m_k(X) + \eta_k)$$

where $m_k : \mathcal{X} \rightarrow \mathbb{R}^{p_k}$ is the channel's measurement function, Q_ε is a resolution operator (binning, rounding, quantization) at scale ε , and η_k is channel noise.

The information available at resolution ε is the sigma-algebra generated by all accessible channels at that resolution:

$$\mathcal{F}_\varepsilon = \sigma(\{Y_{t_i}^{(k)} : k \in \mathcal{K}(\varepsilon), t_i \in \mathcal{T}(k, \varepsilon)\})$$

Filtration axiom: if $\varepsilon' < \varepsilon$ (finer measuring stick), then $\mathcal{F}_{\varepsilon'} \subseteq \mathcal{F}_\varepsilon$. More telemetry at higher resolution generates a monotone refinement of available information.

Note that ε parameterizes two distinct effects simultaneously: (i) finer quantization of existing channels and (ii) activation of new channels $\mathcal{K}(\varepsilon)$ at sufficient resolution. The latter is critical — many data types become available only past certain technological or legal thresholds, creating discrete jumps in the filtration that continuous models miss.

2.3 The Predictability Coastline

Define the **predictability coastline** as the mutual information (Shannon, 1949) between the task variable and the available information at resolution ε :

$$C(\varepsilon) = I(Z; \mathcal{F}_\varepsilon) = H(Z) - H(Z | \mathcal{F}_\varepsilon)$$

This functional is:

- **Monotone:** $\varepsilon' < \varepsilon \implies C(\varepsilon') \geq C(\varepsilon)$ (more information never hurts, by the data processing inequality).
- **Bounded:** $0 \leq C(\varepsilon) \leq H(Z)$ (cannot predict more than the entropy of the task variable; assuming Z is discrete or discretized so that $H(Z)$ is finite Shannon entropy).
- **Domain-agnostic:** comparable across surveillance, weather, markets, or any other domain.
- **Decision-theoretically grounded:** directly tied to optimal prediction performance via Fano's inequality.

$C(\varepsilon)$ IS the coastline. Its shape — how it grows as ε shrinks — encodes the entire story of how data resolution translates to predictive power.

The Z-dependence problem. A critical subtlety: $C(\varepsilon)$ is a property of the (system, observer, prediction-target) triple, not of the system alone. Changing Z — even for the same system and the same observer — changes the coastline's shape, its phase transitions, and any scalar summary extracted from it. Empirically, the Hénon map produces a sharp

curvature spike with binary Z and no detectable spikes with multinomial Z (Section 9). This means claims of the form “system X has information dimension D_I ” are ill-formed without specifying Z .

This is not a flaw — it is precisely what makes the framework relevant to surveillance analysis, where the measurement setup IS the question. But it requires separating two objects: (i) the *coastline of a surveillance configuration* — $C(\epsilon)$ for a fixed triple, which is computable and operationally meaningful; and (ii) *system-intrinsic predictive structure* — information that is invariant under changes to Z . Object (i) is well-defined by construction. Object (ii) requires a coherence criterion, which we develop in Section 9.5.

Naming convention. To make the epistemic status of each quantity explicit, we adopt the following notation throughout: $C_Z(\epsilon) = I(Z; \mathcal{F}_\epsilon)$ for the *raw coastline* conditioned on a specific prediction target Z ; $C_{\text{coh}}(\epsilon) = \min_k \text{NMI}^*(Z_k; \mathcal{F}_\epsilon)$ for the *coherent coastline* — the system-intrinsic component stable across all admissible prediction targets; and $C_{\text{frag}}(\epsilon) = C_{q90}(\epsilon) - C_{\text{coh}}(\epsilon)$ for the *fragile band* — the Z -dependent residual that the coherence filter strips away. The raw coastline C_Z is what an observer computes given a specific analytic goal. The coherent coastline C_{coh} is what survives adversarial re-questioning. The fragile band C_{frag} is measurement artifact.

2.4 Fractal Dimension as Scaling Exponent

Define the scale variable $s = \log(1/\epsilon)$, so increasing s corresponds to increasing resolution.

The **local scaling exponent** is:

$$D_{\text{loc}}(s) = \frac{d}{ds} C(e^{-s})$$

This is directly interpretable: *how many new bits of predictive capacity per octave of resolution*. It is the “fractal dimension” of the coastline — the rate at which new predictable structure appears as the measuring stick shrinks (cf. the correlation dimension of Grassberger & Procaccia, 1983, which characterizes geometric scaling of strange attractors; the information dimension D_I below is its information-theoretic analogue).

For a smooth system, D_{loc} is constant (linear scaling, integer effective dimension). For a fractal system, D_{loc} varies with s — different resolution regimes reveal structure at different rates. The non-constancy of D_{loc} IS the “fractal” signature: it means the coastline has structure at every scale, with some scales being richer than others.

The **global information dimension** (when the limit exists):

$$D_I = \lim_{s \rightarrow \infty} \frac{C(e^{-s})}{s}$$

characterizes the asymptotic scaling. A high D_I means predictive complexity never saturates — there is structure all the way down. A low D_I means early plateau — the system’s behavior is low-dimensional regardless of measurement precision. However, D_I inherits the Z -dependence of $C(\epsilon)$: it is a property of the surveillance configuration, not of the system. No single scalar extracted from $D_{\text{loc}}(s)$ correctly orders systems across domains (Section 9.4).

2.5 Phase Transitions

Define the **phase transition density** as the curvature of the coastline:

$$\kappa(s) = \frac{d^2}{ds^2} C(e^{-s})$$

Large positive values of κ identify resolution thresholds where small increases in data produce disproportionate jumps in predictive capacity. These are the points where qualitatively new predictive axes unlock. Note that κ -spikes, like $C(\epsilon)$ itself, depend on the prediction target Z — the same system can show sharp spikes or none depending on how Z is defined (Section 9.4). The κ -profile is a signature of the surveillance configuration, not of the system alone.

Equivalently, using Fisher information geometry: let the threat model maintain a parametric posterior $p_\theta(z | \mathcal{F}_\epsilon)$. The Fisher information matrix at scale ϵ ,

$$\mathcal{J}(\varepsilon) = \mathbb{E} \left[\nabla_{\theta} \log p_{\theta}(Z \mid \mathcal{F}_{\varepsilon}) \nabla_{\theta} \log p_{\theta}(Z \mid \mathcal{F}_{\varepsilon})^{\top} \right]$$

concentrates its eigenvalue growth at exactly the phase transition points. Peaks in $\text{tr}(\mathcal{J})$ or $\log \det(\mathcal{J})$ as functions of s identify the critical thresholds from a differential-geometric perspective.

The two characterizations (coastline curvature and Fisher information peaks) detect the same events from different mathematical angles: the coastline curvature is model-free, while the Fisher formulation reveals the *direction* of the new predictive axis in parameter space.

3 Topological Unlock Events

3.1 From Filtration to Topology

The fractal dimension and phase transition density characterize *how much* new predictive capacity appears. Persistent homology characterizes *what kind* of structure is appearing — the qualitative shape of the behavioral space as revealed at each resolution.

Let $\Phi : \mathcal{X}^L \rightarrow \mathbb{R}^m$ be a behavioral embedding (Takens delay map, learned representation, or domain-specific feature extraction) that maps length- L windows of the target process into a point cloud.

At each ε , consider the posterior distribution over behavioral windows:

$$\Pi_{\varepsilon}(\cdot) = \mathbb{P}(X_{t:t+L} \in \cdot \mid \mathcal{F}_{\varepsilon})$$

Sample N windows from Π_{ε} , embed via Φ , and obtain a point cloud $P_{\varepsilon} \subset \mathbb{R}^m$. Build the Vietoris-Rips filtration $\text{VR}(P_{\varepsilon}, r)$ over neighborhood radius r and compute persistent homology:

$$\text{PH}_{\varepsilon} = \text{PH}(\text{VR}(P_{\varepsilon}, r))$$

3.2 Betti Number Interpretation

The Betti numbers of the persistence diagram have direct behavioral interpretations:

- β_0 : **Connected components** — clusters of distinguishable behavioral modes. At coarse resolution, one blob (“a person”). At finer resolution, work-mode separates from home-mode separates from transit-mode.
- β_1 : **Loops** — cyclical behavioral patterns. Daily routines, weekly patterns, seasonal variations. These are the “Groundhog Day” structures that become visible only with sufficient temporal resolution.
- β_2 : **Voids** — regions of behavioral space the target systematically avoids. These are often the most revealing features: what someone *never does* is structural information about their constraints, values, or fears.

3.3 Topological Unlock Criterion

An ε -decrease causes a **topological unlock** if the persistent homology undergoes a macroscopic change:

- Appearance of a long-lived H_k class (a stable topological feature that persists across a wide range of r).
- Significant increase in total persistence $\sum_i (d_i - b_i)$ where b_i, d_i are birth/death times.
- Jump in persistence entropy $E_{\text{pers}} = -\sum_i p_i \log p_i$ where $p_i = (d_i - b_i) / \sum_j (d_j - b_j)$.
- Emergence of new connected-component structure stable over wide r -range.

The Coastline Theorem (informal): The total persistence of topological features scales fractally with data density. Curvature spikes in the predictability coastline $C(\varepsilon)$ correspond to regimes of rapid topological change in PH_{ε} — the moments where the observer’s model gains qualitatively new structural features, not merely quantitative refinement of existing ones. The precise nature of this correspondence — whether it is exact coincidence or a looser coupling at related scales — is characterized empirically in Section 9.

This is the formal bridge: *data resolution* $\rightarrow$ *predictability scaling* $\rightarrow$ *phase transitions* $\rightarrow$ *topological structure*.

4 The Capture Threshold

4.1 Definition

Every agent — human, institutional, computational — has a self-model: a representation of its own behavioral patterns, tendencies, and states accessible through introspection, memory, and local feedback.

Let the target’s self-accessible information at scale ε be $\mathcal{S}_\varepsilon$ (comprising memory, introspection, local logs, conscious self-monitoring — whatever self-knowledge the target can access). Define:

$$C_{\text{threat}}(\varepsilon) = I(Z; \mathcal{F}_\varepsilon), \quad C_{\text{self}}(\varepsilon) = I(Z; \mathcal{S}_\varepsilon)$$

The **capture threshold** is crossed when:

$$C_{\text{threat}}(\varepsilon) \geq C_{\text{self}}(\varepsilon) + \Delta$$

over a relevant scale band, where $\Delta > 0$ is the minimum gap required for operational exploitation (the observer must not merely match but exceed the target’s self-knowledge by enough to act on the difference).

4.2 Kernel Asymmetry as the Mechanism of Capture

The critical insight, confirmed by experiments on deterministic (Hénon) and partially observable (Lorenz) systems (Section 9), is that capture is not primarily about *resolution depth* — observing the same variables at finer granularity — but about *kernel asymmetry*: observing variables that the target’s self-model projects away entirely.

Let $\pi_{\text{self}} : \mathcal{X} \rightarrow \mathcal{S}$ and $\pi_{\text{obs}} : \mathcal{X} \rightarrow \mathcal{Y}$ be the observation maps of the self-model and observer respectively. Capture exists when:

$$\exists x, x' \in \mathcal{X} : \pi_{\text{self}}(x) = \pi_{\text{self}}(x') \text{ and } \pi_{\text{obs}}(x) \neq \pi_{\text{obs}}(x')$$

That is, the observer distinguishes states that the self-model collapses — it resolves distinctions within the self-model’s fibers. The capture threshold is the resolution at which this non-factorization becomes resolvable.

This reframes the surveillance problem. A target with full state access (Hénon, clean) has injective π_{self} — trivially no capture, since there are no non-trivial self-fibers. A target observing only $x(t)$ of a Lorenz system collapses the (y, z) plane; when the observer sees all three coordinates, the self-model’s fibers contain the wing-distinguishing information, and capture is possible even at *coarse* resolutions. Adding noise to a deterministic system creates non-trivial self-fibers (the noise dimensions the self-observer can’t average out), enabling capture via sample-size advantage.

Below the capture threshold, the observer has a *model* of the target — useful for aggregate prediction but unable to anticipate behavior the target doesn’t anticipate in itself. The target retains epistemic sovereignty.

Above the capture threshold, the relationship changes in kind:

- The observer can predict behaviors the target has not yet decided on, because it detects patterns in the target’s decision process that fall in the self-model’s kernel.
- The observer can reconstruct past states the target has forgotten or never consciously registered.
- The observer can identify influence vectors — inputs that will predictably shift the target’s behavior — that the target cannot detect because they operate in the self-model’s kernel.

This is not a continuous intensification of surveillance. It is a **phase transition** — a qualitative change in the relationship between observer and target. Below threshold: asymmetric information. Above threshold: asymmetric agency.

4.3 Structural Identity with Control Inversion

The non-factorization condition — that π_{obs} resolves distinctions within π_{self} 's fibers — is structurally identical to control inversion in Ashby's cybernetic framework: the observer's regulatory variety exceeds the target's self-regulatory variety *in the specific subspace that the self-model projects away*. The target becomes *substrate* — a system whose dynamics are governed by an external attractor rather than its own regulatory structure.

This connects directly to the excitability topology (Close 2026c): the capture threshold IS control inversion at the epistemic level. The population's regulatory variety (self-knowledge, privacy practices, institutional oversight) is exceeded by the modeler's dynamics in the kernel dimensions. The population becomes substrate in the precise formal sense — its completability class is forced from graceful to terminal by the variety inversion.

4.4 Structural Identity with Sub-Threshold Grammatical Capture

The capture threshold also formalizes what cognitive security research identifies as sub-threshold grammatical capture: influence that operates below the target's parse-resolution. When $C_{\text{threat}} > C_{\text{self}}$, the observer can construct inputs that are optimized for the target's actual response function (which the observer models at resolution the target lacks) while appearing neutral or benign to the target's coarser self-model.

The gap $C_{\text{threat}} - C_{\text{self}}$ IS the capture surface — the space of influence vectors invisible to the target. This is the same morphism whether the substrate is:

- **Computational:** AI safety classifier bypass (generative model dimensionality exceeds monitoring dimensionality).
- **Cognitive:** Targeted manipulation (behavioral model resolution exceeds self-model resolution).
- **Institutional:** Regulatory capture (entity's model of regulatory behavior exceeds the regulator's self-model of its own decision process).

5 Bidirectional Temporal Prediction

5.1 The Temporal Symmetry

A critical and underexplored consequence of the framework: the same data resolution that enables forward prediction enables backward reconstruction. Consequence chains are traversable in both directions from observed data points.

Define forward and backward coastlines:

$$C^+(\varepsilon) = I(Z^+; \mathcal{F}_\varepsilon), \quad C^-(\varepsilon) = I(Z^-; \mathcal{F}_\varepsilon)$$

where $Z^+ = g(X_{t+\tau})$ and $Z^- = h(X_{t-\tau})$.

In general, $C^+ \neq C^-$ — the past and future are not equally inferable from the present. But the *scaling* of both with ε follows the same fractal law, because the topological features that become visible at each resolution (the loops, clusters, and voids in behavioral space) constrain inference in both temporal directions simultaneously.

5.2 Retroactive Inference

The backward coastline $C^-(\varepsilon)$ formalizes something with profound legal and ethical implications: the capacity to reconstruct past states that were never directly observed. When data resolution crosses a topological unlock threshold — when a new persistent homological feature appears — the observer gains the ability to infer not only future behavior but past behavior that occurred before the relevant data was collected.

This is possible because the topological features revealed by current high-resolution data constrain the space of *histories consistent with the present*. A persistent H_1 class (a cyclical pattern) in current data implies the pattern existed before observation began. A void (β_2 feature) implies the avoidance behavior is structural, not situational, and can be projected backward.

The retroactive inference capacity scales with the same fractal dimension as forward prediction. An observer who crosses the capture threshold gains predictive dominance in both temporal directions simultaneously.

5.3 Temporal Asymmetry in the Coastline

The difference $C^+(\varepsilon) - C^-(\varepsilon)$ as a function of ε is itself an informative quantity: it characterizes the target system’s *temporal irreversibility* at each resolution. Systems with high forward-backward asymmetry (creative processes, rare decisions) are harder to reconstruct than to predict. Systems with low asymmetry (habitual behavior, institutional routines) are equally inferable in both directions.

The phase transitions in C^+ and C^- need not coincide. A data type that unlocks forward prediction may not unlock backward reconstruction (real-time biometrics reveal future stress responses but not past ones unless logged). Conversely, a data type that unlocks backward reconstruction may not help forward prediction (archived communications reveal past intent but may not predict future shifts). The *joint* phase transition structure — where both coastlines jump simultaneously — identifies data types that are maximally powerful in both directions. These are the data types most worthy of protection.

6 Defense Topology

6.1 The Kernel Principle

From the bottleneck framework (Close 2026a): every projection has a kernel — a subspace mapped to zero. The observer’s filtration $\mathcal{F}_\varepsilon$ is a projection of the target’s full state X_t . What survives the projection is what the observer can model. What falls in the kernel is invisible.

Defense is not about reducing data volume uniformly. It is about understanding the observer’s filtration topology and disrupting the specific data points that cause topological unlock events.

6.2 Bridge Connectivity and Data Hygiene

Theorem (Defense Non-Uniformity): Removing a data element that bridges otherwise disjoint behavioral clusters in the observer’s model is substantially more effective than removing data that merely increases point density within an existing cluster. Empirically, dispersed removal damages predictive capacity $\sim 7.7\times$ more than density-matched removal on the Hénon attractor (Section 9).

The mechanism is a counting argument: scattered points populate unique (quantized state, next state) bins in the joint distribution used to compute $C(\varepsilon)$; removing them empties bins and destroys mutual information. Clustered points share bins with their neighbors; removing them barely changes bin occupancy. This is a graph-connectivity statement, not a persistent homology statement — it does not require the removed data to correspond to specific topological features.

The stronger claim — that removing the birth edge of a persistent H_1 class is disproportionately effective — was tested directly (Section 9.2). Birth-edge removal produced $> 3\times$ amplification for 2 of 5 persistent features but not consistently across all features. The homological defense thesis remains undemonstrated; the bridge-connectivity thesis is robust.

This means:

- **Volume-based data minimization** (deleting old emails, reducing social media posts) has limited defensive value if the remaining data preserves the observer’s connectivity structure.
- **Bridge-targeted disruption** (eliminating the data that connects otherwise disjoint behavioral clusters — the bridge between work-self and home-self, the link between stated beliefs and revealed preferences) is the high-leverage intervention.
- **The most valuable data to protect** is not the most “sensitive” by conventional categories but the most *structurally critical* — the data whose presence or absence changes the connectivity of the observer’s behavioral model. Graph-theoretic measures (betweenness centrality on the ε -neighborhood graph) are the natural tool; persistent homology may identify relevant features in some cases but is not the right abstraction in general.

6.3 Defense Scaling

The fractal nature of the coastline implies a fundamental asymmetry between attack and defense:

- **Attack scales continuously:** more data monotonically increases $C(\epsilon)$. The observer never loses capability by gaining data.
- **Defense scales discretely:** removing data has large effects only at topological transition points. Between transitions, data removal barely affects the observer’s model.

This asymmetry suggests that *continuous* privacy practices (always minimizing data) are less effective than *topologically informed* practices (identifying and protecting specific critical data elements). The persistence diagram of one’s own data exposure — if computable — would be the optimal guide to defensive data hygiene.

7 Domain Instantiations

7.1 Individual Surveillance

Channels, ordered by typical resolution threshold:

Channel k	Measurement m_k	Typical ϵ_k	Topological unlock
Demographics	Name, age, zip	Coarsest	β_0 : population cluster membership
Purchase history	Transaction records	ϵ_1	β_1 : consumption cycles, lifestyle loops
Communication metadata	Call/message graph (no content)	ϵ_2	New β_0 at relational level: social graph topology
Location trace	GPS/WiFi continuous	ϵ_3	β_1 refinement: spatial routines, β_2 : avoided locations
Content analysis	Messages, posts, search queries	ϵ_4	β_2 : voids in belief/value space
Biometrics	Heart rate, gait, voice	ϵ_5	Phase transition: physiological→behavioral coupling visible
Cross-referencing	Correlation with others’ data at ϵ_5	ϵ_6	$C_{\text{threat}} > C_{\text{self}}$: relational dynamics invisible to either party

The capture threshold typically falls between ϵ_4 and ϵ_5 for individual targets — when the observer’s model integrates enough channels to detect patterns that operate below conscious self-monitoring. The exact location depends on the target’s metacognitive resolution C_{self} .

The channel orderings and topological unlock predictions in the domain tables are conceptual instantiations, not empirical results. Testing them requires datasets with known multi-channel structure at varying resolutions — e.g., smartphone sensor suites (accelerometer, GPS, app usage, communication logs) with ground-truth behavioral labels. The present empirical validation (Section 9) uses dynamical systems with known properties precisely because ground-truth attractor structure enables falsifiable predictions. The framework’s surveillance application depends on the formal machinery (kernel asymmetry, bridge connectivity, coherent coastline) being correct on controlled systems; the domain tables illustrate how the machinery would be deployed, not where it has been deployed.

7.2 Financial Markets

Channel	Measurement	Topological unlock
Daily close prices	End-of-day summary	β_0 : regime classification (bull/bear/sideways)
Intraday prices	Tick-level time series	β_1 : intraday cycles, market maker patterns
Order book	Full depth-of-book	β_2 : strategic voids (price levels systematically avoided)
Cross-market data	Correlations, arbitrage	New β_0 : market clustering by information flow
Alternative data	Satellite, social, supply chain	Higher-order topological features

Financial markets exhibit rich multi-scale structure, with κ -spikes corresponding to distinct predictive phase transitions across resolutions (Section 9). No single scalar metric correctly orders financial systems against dynamical systems — $C_{\max}$ correctly places deterministic chaotic systems above markets, while $D_{\text{structured}}$ and D_I fail. The coherent coastline (Section 9.5) resolves this: financial systems have low C_{coh} (SPY: 0.135, TLT: 0.051, GLD: 0.036) and high fragility ratios (2.4 to 13.4), indicating that most market predictability is prediction-target-dependent rather than system-intrinsic. This is consistent with the efficient-market hypothesis: genuine structural information in prices is scarce, and most apparent predictability reflects the observer’s choice of what to predict.

7.3 Climate Systems

Channel	Measurement	Topological unlock
Station network	Synoptic observations	β_0 : pressure systems, fronts
Satellite + radar	Mesoscale fields	β_1 : squall lines, convective cycles
Dense surface sensors	Boundary layer profiles	β_2 : microclimatic voids, urban heat islands
IoT distributed	Building-scale measurements	Higher-order coupling between scales

Climate has intermediate D_I — enormous complexity but governed by conservation laws that constrain the topology. The phase transitions correspond to the well-known scale gap between synoptic and mesoscale meteorology.

8 Connection to the Bottleneck Primitive

The bottleneck primitive (Close 2026a) formalizes the observation that every inference act projects high-dimensional data through a lower-dimensional representation, destroying some information (the kernel) and preserving some (the image). The predictability coastline is the specific case where the bottleneck is parameterized by resolution ε : at each ε , the filtration $\mathcal{F}_\varepsilon$ preserves the projection’s image and destroys its kernel. The coastline $C(\varepsilon)$ traces how the image grows as the bottleneck widens. The capture threshold is the resolution at which the observer resolves distinctions within self-fibers that the self-model collapses — i.e., the observer preserves what the target’s self-observation destroys.

The predictability coastline is thus a specific instantiation of the bottleneck primitive. The filtration $\mathcal{F}_\varepsilon$ IS the bottleneck, parameterized by resolution. $C(\varepsilon)$ measures what survives the projection. $H(Z | \mathcal{F}_\varepsilon)$ measures what is destroyed. The fractal dimension D_{loc} characterizes the *scaling law of the bottleneck’s resolving power as a function of aperture width*.

The connection deepens in the adversarial case. The bottleneck paper establishes that every defensive system is a statistical apparatus with a characteristic kernel — a subspace invisible to its projection. The present paper asks: *how does that kernel’s structure scale?* The answer — fractally, with topological unlock events — provides the formal bridge between the general epistemological point (all inference is projection) and the specific strategic point (how data translates to power).

The capture threshold $C_{\text{threat}} \geq C_{\text{self}} + \Delta$ is the bottleneck paper’s “apparatus measures itself” pathology turned outward: the observer’s bottleneck resolves structure that the target’s self-model bottleneck destroys. The target’s self-model IS

a bottleneck with its own kernel. When the observer can see into that kernel, the observer models the target better than the target models itself.

The fractal dimension of a threat model’s data filtration is therefore the formal measure of achievable Shannon entropy collapse in adversarial settings — the connection to encryption as entropy-relative boundary (Close 2026b). An observer with D_I approaching the target’s intrinsic behavioral complexity approaches the regime where the target’s communications are inferable from context regardless of cryptographic protection, because the observer’s world model collapses the plaintext entropy without breaking the cipher.

The coherent coastline (Section 9.5) makes this bottleneck connection explicit. The bottleneck paper defines the kernel of a single projection. The coherent coastline defines the kernel of the *system itself* by taking the intersection across all projection targets: $C_{\text{coh}}(\epsilon) = \min_k \text{NMI}^*(Z_k; \mathcal{F}_\epsilon)$. What survives this intersection is system-intrinsic information — the structure that any observer with access to $\mathcal{F}_\epsilon$ can extract regardless of what they choose to predict. What doesn’t survive is measurement artifact. This is the coherence primitive (Close 2026d) applied to predictive capacity: information is “real” if and only if it is conserved across independent verification paths.

9 Empirical Validation

We test the framework’s predictions on dynamical systems with known properties (Hénon map, Lorenz attractor) and on financial time series (SPY, TLT, GLD), with a thermostat process as a structureless baseline. Two rounds of experiments were conducted: the first identified methodological problems, and the second corrected them. We report both rounds because the failures are as informative as the successes.

9.1 Systems Tested

Hénon map ($a = 1.4$, $b = 0.3$): A two-dimensional discrete map with a strange attractor of fractal dimension ≈ 1.26 . Used as the primary test system for defense topology and self-model experiments. Fully deterministic given state access.

Lorenz system ($\sigma = 10$, $\rho = 28$, $\beta = 8/3$): A three-dimensional continuous flow with a strange attractor of fractal dimension ≈ 2.06 . Used for partial-observability capture experiments (target sees $x(t)$ only; observer sees (x, y, z)) and as an intermediate-complexity benchmark for the information dimension.

Financial time series (SPY, TLT, GLD; 2015-2025 daily): Delay-embedded into $\mathbb{R}^5$ with $\tau = 1$. Represent real-world multi-scale systems where the framework should apply.

Thermostat: A stochastic control process (Ornstein-Uhlenbeck around a setpoint). Represents the structureless baseline — high noise, low intrinsic complexity.

9.2 Experiment 1: Defense Topology

Round 1 (global removal). We identified topologically critical points via Ripser’s cocycle representatives for the top- k H_1 bars of the Hénon attractor and compared the mutual information loss from targeted removal vs. random removal vs. density-matched removal. The cocycle method selected 590 of 2000 sampled points (29.5%) as “critical” — far too many for a targeted test. At that fraction, targeted and random removal produced nearly identical MI loss (ratio 0.81 $\times$). However, density-matched removal (removing the same number of points from the densest region) caused 7.7 $\times$ less damage than dispersed removal. This confirmed that *spatial connectivity* matters for defense but did not test the specifically homological claim.

Round 2 (local birth-edge removal). We restricted to the 2 birth-edge vertices for each of the top-5 H_1 bars (10 critical points total) and computed local MI within a neighborhood of radius $R = 2 \times \text{death_value}$ around each feature. For each feature, we compared birth-edge removal to 50 trials of random-in-neighborhood removal. The amplification ratio exceeded 3 $\times$ for 2 of 5 features (F3: 9.6 $\times$, F4: 9.4 $\times$) but fell below 3 $\times$ for 3 of 5 (F0: 2.1 $\times$, F1: $-6.6\times$, F2: 2.4 $\times$). The 3-of-5 threshold was not met.

Assessment: The data supports the spatial/connectivity version of the defense thesis (bridge data is more valuable than redundant data) but does not yet demonstrate that H_1 birth edges specifically control local predictability. The defense claim in Section 6 is revised accordingly.

9.3 Experiment 2: Self-Model and Capture Threshold

Round 1 (Hénon, full state access). The naive self-model (predict $x_{t+1} = x_t$) achieved $C_{\text{self}} = 3.481$ bits, exceeding $C_{\text{threat}} = 3.444$ bits. The Hénon map is deterministic, so a target with full state access has complete self-knowledge — no capture threshold exists. The VAR(1) self-model’s gap (0.86 bits) reflects model-class mismatch (linear vs. nonlinear), not genuine information asymmetry. Noise-augmented variants ($\sigma = 0.01, 0.05, 0.1$) showed capture thresholds emerging via the observer’s sample-size averaging advantage over the noisy self-observer.

Round 2 (Lorenz, partial observability). With the target observing only $x(t)$ and the observer seeing (x, y, z) , the framework’s predictions partially confirmed:

- **Capture exists:** $C_{\text{threat}} > C_{\text{self}}$ at fine resolutions (mean gap 0.038 bits for delay-embedding self-model; 1.12 bits for naive self-model).
- **Self-model ordering:** naive ($\bar{C} = 1.44$) < delay-embedding ($\bar{C} = 2.52$) < full-state threat ($\bar{C} = 2.56$). This confirms that partial observability creates genuine self-model limitations.
- **Topological asymmetry at capture:** At the capture resolution, the observer’s point cloud had 4 persistent H_1 features with total persistence 52.4, while the self-model’s point cloud had 0 persistent features. This supports Theorem Part (3): capture IS accompanied by a topological gap.
- **κ -spike misalignment:** The κ -spikes occurred at $s = 0.38, -1.06, -2.01$, while capture thresholds fell at $s \approx -3.45$. All coincidence tests returned false ($|s_{\text{capture}} - s_{\kappa}| > 0.5$). The phase transitions in the coastline’s curvature do not coincide with the capture threshold on this system.

Assessment: Capture is real and involves topological asymmetry (the observer resolves structure the self-model cannot). But the precise claim that capture *coincides with* a κ -spike is not supported — capture depends on the self-model’s quality relative to the observer’s, which is a separate axis from where the coastline has its sharpest curvature.

9.4 Experiment 3: Information Dimension Across Domains

Round 1 (binary Z , mean D_{loc}). With binary Z (next-state sign), the MI ceiling is 1 bit, and $D_I = \text{mean}(D_{\text{loc}})$ produced the ranking: Thermostat (0.390) > TLT (0.357) > GLD (0.334) > SPY (0.287) > Hénon (0.190). A thermostat topping the ranking demonstrates the metric is broken: $\text{mean}(D_{\text{loc}})$ rewards uniformity of information gain, and the thermostat — having no structure — gains information uniformly across all scales.

Round 2 (multinomial Z , composite metric). With 16-bin Z (lifting the MI ceiling to $\log_2 16 = 4$ bits) and $D_{\text{structured}} = \max(D_{\text{loc}}) \times (1 - H_{\text{spectral}}/H_{\text{max}})$, the ranking becomes:

$$\text{GLD} (0.056) > \text{Thermostat} (0.050) > \text{TLT} (0.034) > \text{Lorenz} (0.029) > \text{Hénon} (0.020) > \text{SPY} (0.014)$$

The thermostat remains above Hénon under $D_{\text{structured}}$, failing to fix the critical ordering. However, Hénon < Lorenz tracks the known dimensional difference ($1.26 < 2.06$), confirming the metric captures intra-class dimensional ordering. The cross-class failure (thermostat vs. dynamical systems) arises from a curse-of-dimensionality confound: the thermostat’s 5D delay embedding with ~ 2500 points creates a narrow “usable ε ” window where the steep transition from bin-explosion to useful binning mimics concentrated information gain.

The information ceiling C_{max} — which the multinomial Z was designed to differentiate — succeeds cleanly: Hénon (2.67) > Lorenz (2.59) > Thermostat (2.24) > GLD (1.46) > TLT (1.33) > SPY (1.06). This reflects the intrinsic predictive complexity of each system. Deterministic chaotic systems carry more information about their own next state than stochastic or externally driven systems, exactly as expected.

The κ -spike structure differentiates domains: financial systems show 1-2 spikes each, the thermostat shows 2, while Hénon and Lorenz show 0 above the signal threshold — meaning the dynamical systems’ information gain is smooth across scales while financial and stochastic systems have discrete transition points. (This κ -spike count is Z -dependent; the coherent coastline in Section 9.5 reveals 8 κ -spikes for Hénon and 1 for Lorenz across the full admissible target family.)

Assessment: No single scalar (D_I , $D_{\text{structured}}$, C_{max} , κ -count) captures the full picture. Different scalars excel at different comparisons: C_{max} correctly orders by intrinsic complexity, $D_{\text{structured}}$ captures intra-class dimensional differences, and

κ -spike structure distinguishes domains qualitatively. The framework’s value is the multi-dimensional characterization (coastline shape, κ -structure, persistence landscape), not any single extracted number. This is itself a substantive finding: the coastline IS the object of interest, and reducing it to a scalar destroys the signal.

9.5 Experiment 4: Coherent Coastline (CMC Resolution)

The Z-dependence problem identified in Sections 9.3-9.4 — that $C(\varepsilon)$ and its derivatives depend on the prediction target Z , not just the system — motivates a coherence filter. Following the principle that system-intrinsic information should be stable across all ways of interrogating the system, we define the **coherent coastline**:

$$C_{\text{coh}}(\varepsilon) = \min_k \text{NMI}^*(Z_k; \mathcal{F}_\varepsilon)$$

where NMI^* is a null-corrected normalized mutual information defined as:

$$\text{NMI}(Z_k; \mathcal{F}_\varepsilon) = \frac{I(Z_k; \mathcal{F}_\varepsilon)}{H(Z_k)}, \quad \text{NMI}_{\text{null}} = \frac{\bar{I}_{\text{null}}(Z_k; \mathcal{F}_\varepsilon)}{H(Z_k)}$$

$$\text{NMI}^*(Z_k; \mathcal{F}_\varepsilon) = \max(0, \text{NMI}(Z_k; \mathcal{F}_\varepsilon) - \text{NMI}_{\text{null}})$$

Both the raw and null mutual information are normalized by $H(Z_k)$ before subtraction, ensuring dimensional consistency (both terms are dimensionless ratios in $[0, 1]$). The minimum is taken over K diverse prediction targets constrained to have entropy $H(Z_k) \in [1.4, 5.8]$ bits, which excludes near-degenerate and near-uniform discretizations. The normalization by $H(Z_k)$ prevents high-entropy targets from dominating; the null correction (permutation baseline, $n = 50$) kills noise-floor artifacts.

For each of the six systems, we generated 10-15 diverse prediction targets (coordinate projections, lag variations, amplitude/angle observables, quantization sweeps) and computed the coherent coastline over 40 logarithmically-spaced ε values.

Results. The coherent coastline produces a clear ordering:

System	$C_{\text{coh}}^{\text{max}}$	C_{q10}^{max}	Fragility ratio	κ -spikes	$n_{\text{admissible}}$
Lorenz	0.398	0.505	1.42	1	13
Hénon	0.282	0.658	3.95	8	11
SPY	0.135	0.201	2.40	1	7
TLT	0.051	0.147	7.55	1	8
GLD	0.036	0.147	13.4	1	9
Thermostat	0.025	0.034	29.6	5	10

The **fragility ratio** $\mathcal{R}_{\text{frag}} = \bar{C}_{\text{frag}}/\bar{C}_{\text{coh}}$ measures how much predictability is Z-dependent vs. system-intrinsic. The thermostat’s fragility ratio of 29.6 confirms the prediction: almost all its apparent predictive structure is noise-floor artifact that collapses under the coherence filter. The Lorenz system’s ratio of 1.42 — the lowest — means nearly all its predictability is intrinsic: the attractor’s two-wing structure is visible to every prediction target.

κ -significance caveat. The thermostat shows 5 κ -spikes on a coherent floor of $C_{\text{coh}}^{\text{max}} = 0.025$. These are not meaningful phase transitions — they are curvature fluctuations in a near-zero signal. A κ -spike is significant only when its amplitude is large relative to the coherent signal it modulates; spikes on a near-zero coherent baseline are noise artifacts, not structural unlocks. The Hénon system’s 8 κ -spikes on a coherent floor of 0.282 represent genuine multi-scale structure; the thermostat’s 5 spikes on 0.025 do not.

The fragility ratio has a direct adversarial interpretation: it quantifies *how much predictability disappears when the target re-asks the question*. A surveillance configuration with $\mathcal{R}_{\text{frag}} \approx 1$ (Lorenz) provides robust intelligence — the observer’s predictions hold regardless of what aspect of the target they focus on. A configuration with $\mathcal{R}_{\text{frag}} \gg 1$ (thermostat:

29.6, GLD: 13.4) provides brittle intelligence — the observer’s apparent predictive power depends critically on asking exactly the right question, and shifts in the prediction target collapse it.

Success criteria (4 of 5 pass):

1. *Thermostat collapse*: $C_{\text{coh}}^{\text{max}} = 0.025$, lowest of all systems. **Pass.**
2. *Hénon > Lorenz*: $C_{\text{coh}}^{\text{max}}(\text{Hénon}) = 0.282 < 0.398 = C_{\text{coh}}^{\text{max}}(\text{Lorenz})$. **Fail.** The Lorenz attractor’s two-wing structure produces more system-intrinsic information than the Hénon map’s more uniform chaos. This is interpretable: the Lorenz system’s low fragility ratio (1.42 vs. Hénon’s 3.95) indicates its structure is more robust to changes in what one predicts.
3. *Financial < chaotic*: All financial systems have C_{coh} less than both chaotic systems. **Pass.**
4. *Hénon κ -stability*: 8 κ -spikes across the coherent coastline, indicating rich multi-scale coherent structure. **Pass.**
5. *Thermostat more fragile*: Fragility ratio $29.6 \gg 1.42$ (Lorenz). **Pass.**

The one failure (Hénon > Lorenz) is itself informative and reveals a fundamental distinction between two axes of predictability:

System	C_{max} (best-Z)	Rank	$C_{\text{coh}}^{\text{max}}$ (all-Z)	Rank	$\mathcal{R}_{\text{frag}}$
Hénon	2.67	1	0.282	2	3.95
Lorenz	2.59	2	0.398	1	1.42
Thermostat	2.24	3	0.025	6	29.6
SPY	1.06	6	0.135	3	2.40

The Hénon map has higher *maximum extractability* (C_{max}) — under the best prediction target, more information is available. But the Lorenz system has higher *coherent extractability* (C_{coh}) — its information is more democratically available across all prediction targets. The inversion names a distinction: a system can be highly predictable under optimal questioning while being less robustly predictable than a system with lower peak capacity. For surveillance assessment, coherent extractability is the more relevant axis — an adversary rarely knows in advance which question to ask, and a target can shift what matters.

Estimator robustness. To verify that binning artifacts do not drive the results, we recomputed the full Lorenz coherent coastline using a k -nearest-neighbor mutual information estimator ($k = 5$, Ross 2014) instead of bin-counting. The kNN estimator produced $C_{\text{coh}}^{\text{max}} = 0.442$ vs. the binned estimate of 0.398 — 11% higher, consistent with the known downward bias of binned MI estimation. The ordering across all prediction targets was identical at every ε value tested (100% agreement rate). The qualitative conclusions are not sensitive to the choice of MI estimator.

A limitation of the present financial analysis is that daily data ($N \approx 2500$) in 5D delay embedding approaches the regime where binned MI estimation is unreliable for fine ε . The null correction partially compensates (by subtracting permutation-baseline bias), and the qualitative ordering (financial < chaotic) is robust to estimator choice on the systems where both estimators were tested. Extending the financial analysis to intraday data ($N > 10^5$) would provide a stronger test of the framework’s financial predictions and is a priority for future work.

Block-resampling diagnostic. Block-bootstrap resampling (100 iterations, block size 50) serves as a property test distinguishing attractor-derived from temporal-correlation-derived predictability. For the deterministic chaotic systems, the metrics are stable under temporal resampling: Lorenz $C_{\text{coh}}^{\text{max}} = 0.397$ (resampling mean 0.360 ± 0.010), Hénon 0.283 (mean 0.242 ± 0.005). The point estimates fall slightly above the resampling distributions because block boundaries mildly disrupt deterministic dynamics — but the attractor structure survives.

For the thermostat and financial systems, the point estimates fall far outside the resampling distributions: thermostat $C_{\text{coh}}^{\text{max}} = 0.027$ vs. resampling range $[0.131, 0.181]$; SPY 0.134 vs. $[0.199, 0.276]$. Block resampling destroys the precise temporal correlations that produce these systems’ extreme values. The thermostat’s very low coherent information (0.027) depends on specific temporal anti-correlations in the noise process; resampling scrambles these, making the resampled thermostat look like a moderate-complexity system. The financial systems’ low coherence depends on specific serial dependence structure; resampling inflates it.

We define the resampling sensitivity index $S_{\text{resample}} = |C_{\text{point}} - \mu_{\text{boot}}|/\sigma_{\text{boot}}$ to quantify this divergence. Chaotic systems have $S_{\text{resample}} < 5$ (point estimates near the resampling distribution); stochastic and financial systems have $S_{\text{resample}} \gg 10$

(point estimates far from it). This asymmetry is itself a finding: the coherent coastline of deterministic chaotic systems is an attractor property (stable under temporal disruption), while the coherent coastline of stochastic/financial systems is a temporal-correlation property (fragile under resampling). The framework detects this distinction automatically — the same systems with high fragility ratios (thermostat: 29.6, GLD: 13.4) are the ones whose resampling distributions diverge most from their point estimates.

Admissible family sensitivity. The coherent coastline is Z -independent only relative to the chosen admissible family $\{Z_k\}$. Our construction uses: (i) quantization sweeps at 4, 8, 16, 32, 64 bins on the primary observable; (ii) lag variations at $\tau = 1, 2, 5, 10$; (iii) coordinate projections where the system is multi-dimensional; and (iv) functional observables (norm, first difference). Targets are constrained to the entropy band $H(Z_k) \in [1.4, 5.8]$ bits, which excludes near-degenerate and near-uniform discretizations. The min-envelope is taken over $K = 7$ –13 admissible targets per system.

The coherent coastline is operationally defined: it quantifies the predictive information that survives adversarial re-questioning within a stated interrogation protocol. Expanding the protocol (adding prediction targets) can only decrease C_{coh} (the min-envelope is monotone non-increasing in family size). The question of whether C_{coh} converges to a protocol-independent limit — whether a “true” system-intrinsic predictive capacity exists — is a theoretical question we leave open. The empirical observation that $K = 7$ –13 targets suffice to separate six qualitatively different systems suggests convergence is rapid, but we do not claim it. We report the family construction explicitly so that the result is reproducible and the scope of the coherence claim is clear.

Connection to the Z -dependence problem. The coherent coastline resolves the conflation between objects (i) and (ii) identified in Section 2.3. $C_{\text{coh}}(\epsilon)$ is a property of the (system, observer) pair — the prediction-target dependence has been robustified away via a worst-case min-envelope over an admissible family. It is the closest we can get to a system-intrinsic information measure using empirical means. The residual — $C_{\text{frag}}(\epsilon) = C_{q90}(\epsilon) - C_{\text{coh}}(\epsilon)$ — is the measurement artifact that the coherence filter strips away.

Three-axis diagnostic. The paper’s central contribution is a three-axis characterization that replaces the single-scalar reduction that Sections 9.3–9.4 show to be inadequate:

System	Coherent capacity $C_{\text{coh}}^{\text{max}}$	Fragility ratio $\mathcal{R}_{\text{frag}}$	Resampling sensitivity S_{resample}	Interpretation
Lorenz	0.398	1.42	Low	Robust attractor; high intrinsic structure
Hénon	0.282	3.95	Low	Rich attractor; moderate Z -dependence
SPY	0.135	2.40	High	Moderate coherence; correlation- dependent
TLT	0.051	7.55	High	Low coherence; mostly Z -artifact
GLD	0.036	13.4	High	Low coherence; highly fragile
Thermostat	0.025	29.6	High	Near-zero coherence; noise floor

Axis 1 ($C_{\text{coh}}^{\text{max}}$): how much system-intrinsic predictive information exists. Axis 2 ($\mathcal{R}_{\text{frag}}$): how much of the apparent predictability is measurement artifact. Axis 3 (S_{resample}): whether the predictability derives from an attractor (stable) or from temporal correlations (fragile). No single axis suffices; together they give a complete surveillance-power fingerprint.

9.6 Summary of Empirical Status

Claim	Status	Evidence
Fractal scaling (D_{loc} non-constant)	Supported	All systems tested
Phase transitions (κ -spikes) differ across domains	Supported	Distinct κ -profiles for each system
κ -spikes coincide with topological unlocks	Open	Not directly tested at matching scales
Capture threshold exists under partial observability	Supported	Lorenz with x -only self-model
Capture arises from kernel asymmetry	Supported	Hénon clean (no capture), Lorenz partial (capture)
Capture accompanied by topological asymmetry	Supported	4 vs. 0 H_1 features at capture
Capture coincides with κ -spike	Not supported	Lorenz: $ s_c - s_\kappa > 1.4$
Birth-edge removal disproportionately damages prediction	Partial	2/5 features exceed 3 $\times$; not consistent
Dispersed removal $\gg$ dense removal	Supported	7.7 $\times$ ratio on Hénon
$D_I = \text{mean}(D_{\text{loc}})$ orders systems correctly	Not supported	Thermostat tops ranking
$D_{\text{structured}}$ orders systems correctly	Partial	Hénon < Lorenz confirmed; thermostat still misordered
C_{max} orders systems correctly	Supported	Hénon > Lorenz > Thermostat > financial
C_{coh} resolves Z-dependence	Supported	Thermostat collapses; chaotic > financial (4/5 criteria)
Thermostat coherent information ≈ 0	Supported	$C_{\text{coh}}^{\text{max}} = 0.025$, fragility 29.6
Financial systems less coherent than chaotic	Supported	All financial $C_{\text{coh}} <$ both chaotic systems

10 The Flagship Theorem

Theorem (Coastline Theorem): For a target process X_t observed through a filtration $\mathcal{F}_\varepsilon$ with prediction target Z , the predictability coastline $C(\varepsilon) = I(Z; \mathcal{F}_\varepsilon)$ satisfies:

1. **Fractal scaling:** the local scaling exponent $D_{\text{loc}}(s) = dC/ds$ is non-constant in s for generic targets with multi-scale behavioral structure. *Empirically confirmed across all systems tested (Section 9.6).*
2. **Observer-dependent phase transitions:** the coastline curvature $\kappa(s) = d^2C/ds^2$ exhibits sharp peaks at critical resolution thresholds, but the number, location, and magnitude of these peaks depend on the prediction target Z . The κ -profile is a domain-specific signature of the (system, observer, Z) triple, not of the system alone. *Confirmed: the same system (Hénon) shows sharp κ -spikes with binary Z and none with multinomial Z (Section 9.4).* The stronger claim that κ -peaks precisely coincide with topological unlock events remains conjectural; empirical tests show κ -peaks and topological changes at related but distinct scales.
3. **Capture via kernel asymmetry:** if π_{obs} does not factor through π_{self} — i.e., $\exists x, x'$ with $\pi_{\text{self}}(x) = \pi_{\text{self}}(x')$ and $\pi_{\text{obs}}(x) \neq \pi_{\text{obs}}(x')$ — then there exists a critical ε_c where $C_{\text{threat}}(\varepsilon) \geq C_{\text{self}}(\varepsilon) + \Delta$. At this crossing, the observer's model contains structural features absent from the target's self-model. *Confirmed on Lorenz system with partial observability: 4 persistent H_1 features in observer's model vs. 0 in the self-model at the capture resolution. Confirmed negative on Hénon with full state access: no capture when kernel is trivial.* Capture depends on which variables are observed, not on resolution depth per se.

4. **Defense via bridge connectivity:** the marginal defensive value of different data elements is highly non-uniform. Data that bridges otherwise disconnected behavioral clusters causes $\sim 7.7\times$ more damage when removed than density-matched data within existing clusters. *Confirmed on Hénon attractor across two experimental rounds.* The stronger claim that this non-uniformity is proportional to the persistence of specific homological features is not reliably supported (2/5 features show $> 3\times$ amplification from birth-edge removal).
5. **Coherent coastline:** the system-intrinsic predictive information $C_{\text{coh}}(\epsilon) = \min_k \text{NMI}^*(Z_k; \mathcal{F}_\epsilon)$ across diverse prediction targets resolves the Z-dependence problem. C_{coh} strips measurement artifacts: structureless systems (thermostat, $C_{\text{coh}}^{\text{max}} = 0.025$) collapse under the coherence filter while systems with genuine dynamical structure retain high coherent information (Lorenz: 0.398, Hénon: 0.282). The fragility ratio $\mathcal{R}_{\text{frag}} = \tilde{C}_{\text{frag}}/\tilde{C}_{\text{coh}}$ quantifies how much predictability is Z-dependent artifact — equivalently, how much surveillance intelligence disappears when the target shifts what it cares about. *Confirmed across 6 systems; 4 of 5 success criteria pass (Section 9.5).*

Parts (1) and (5) are the core mathematical contributions. Part (2) is a characterization result (the κ -profile signatures domains). Part (3) is the strategic application (capture requires kernel asymmetry). Part (4) is the defensive corollary (bridge data, not volume).

The proof strategy for (1)-(2) requires showing that the mutual information functional’s derivative with respect to the resolution parameter has discontinuities that align with changes in persistent homology for a fixed Z . The key technical step is establishing that the posterior point cloud P_ϵ (Section 3.1) undergoes topological bifurcations at the same ϵ values where the conditional entropy $H(Z | \mathcal{F}_\epsilon)$ has discontinuous derivatives. This is related to but distinct from known results on the stability of persistent homology under Wasserstein perturbation (Cohen-Steiner et al. 2007) and the information-theoretic characterization of topological features (Rucco et al. 2016). The precise coincidence likely requires conditions on the target system’s regularity and on the embedding dimension relative to the attractor’s intrinsic dimension; characterizing these conditions is an open problem.

Part (5) connects the coastline framework to the coherence primitive (Close 2026d): system-intrinsic information (relative to the admissible family) equals information conserved across independent verification paths. The coherent coastline is the information-theoretic version of this principle applied to predictive capacity — what survives interrogation from all angles is structural; what doesn’t is artifact.

11 Scope and Limitations

11.1 What the Framework Does Not Cover

The predictability coastline formalizes *inferential* predictive power — the capacity to predict a target’s behavior from observed data through statistical modeling. Following the scope limitation established in the companion bottleneck paper (Close 2026a, Section 7.1), the framework does not cover:

- **Genuinely novel behavior:** actions that arise from the target’s encounter with structure that did not previously exist in any model. The coastline measures predictability of behavior generated by the target’s existing attractor; it cannot predict the moment the target exits an attractor entirely. This is the Gödelian limit — genuinely novel content survives total prediction.
- **Reflexive defense:** a target who understands the framework and acts on it changes their behavioral topology, invalidating the observer’s model in a way the model cannot anticipate without modeling the target’s modeling of the model (infinite regress). This is the observer-effect problem that the Anthropic Sabotage Risk Report (2026) identifies as “evaluation awareness.”
- **Aggregate vs. individual:** the framework treats individual targets. Population-level application requires additional structure (distributional persistent homology, or treating the population as a single higher-dimensional process).

11.2 Z-Dependence and the Scalar Reduction Problem

Two related lessons emerge from the empirical work:

Z-dependence. $C(\epsilon)$ is not a system invariant — it depends on the prediction target Z . The Hénon map shows a sharp κ -spike with binary Z and none with multinomial Z . No scalar extracted from a single- Z coastline (D_I , $D_{\text{structured}}$, C_{max} , κ -count) produces the expected ordering across all six systems. Two rounds of experiments with six metrics failed to

find one. The coherent coastline (Section 9.5) partially resolves this by taking a worst-case envelope over Z , but the fundamental lesson is that phase transitions in predictive capacity are properties of measurement setups, not of systems.

Scalar reduction. The framework’s fundamental object is the *shape* of $C(\epsilon)$ — a function, not a number. The coastline shape classes (sharp S-curve for Hénon, smooth concave rise for Lorenz, waviness for thermostat, lower plateaus for financial) are qualitatively distinguishable and carry genuine information about the system-observer configuration. Reducing this to a scalar is information-destroying, and the fact that the framework is *about* information destruction at bottlenecks makes this a particularly self-referential failure mode — the coastline-to-scalar projection IS a bottleneck, and it destroys exactly the multi-scale information the framework claims to detect.

The right “numbers” depend on the question: *Does capture exist?* → check kernel asymmetry (binary). *How vulnerable is this data to surveillance?* → compute C_{coh} and $\mathcal{R}_{\text{frag}}$ (continuous). *What data should I protect?* → compute bridge centrality (per-element). These are specific questions with specific scalar answers extracted from the coastline for a specific purpose. The paper’s original error was trying to define ONE number (D_I) that characterizes everything.

Future work should develop distance metrics on coastline shapes (e.g., functional data analysis, Fréchet means on the space of monotone functions) and formalize the coherent coastline’s theoretical properties (conditions under which C_{coh} converges, its relationship to ergodic invariants of the dynamical system).

Whether the coherent coastline converges under family expansion — whether there exists a limit $C_{\text{coh}}^\infty(\epsilon) = \lim_{K \rightarrow \infty} \min_{k \leq K} \text{NMI}^*(Z_k; \mathcal{F}_\epsilon)$ that is independent of the interrogation protocol — is an open question with connections to ergodic theory (the relationship between time-average MI and the system’s invariant measure) and to the information bottleneck method (the minimum sufficient statistic as the family-independent limit).

11.3 Ethical Considerations

The framework is dual-use by construction. The same mathematics that enables an authoritarian state to identify capture thresholds over its population enables a privacy researcher to identify topologically critical data for protection, or a democratic institution to audit whether surveillance programs have crossed the capture threshold.

We note without resolving that the capture threshold $C_{\text{threat}} \geq C_{\text{self}} + \Delta$ defines a bright line that current legal frameworks do not recognize. Existing privacy law focuses on data categories (PII, health data, financial records) and consent mechanisms. The coastline framework suggests that the relevant distinction is not *what kind* of data is collected but *whether the integrated model has crossed the capture threshold for the target or population* — a structural criterion that current law is not equipped to evaluate.

References

- Close, L.J. (2026a). The Bottleneck Primitive: Statistics as the Study of Information Compression. *Zenodo*. <https://doi.org/10.5281/zenodo.18667644>
- Close, L.J. (2026b). Completeness. *Zenodo*. <https://doi.org/10.5281/zenodo.18512735>
- Close, L.J. (2026c). Excitability: A Post-Seizure Cybernetics of Control Inversion Across Substrates of Intelligence. *Zenodo*. <https://doi.org/10.5281/zenodo.18627253>
- Close, L.J. (2026d). Coherence and the Ground of Morality. *Zenodo*. <https://doi.org/10.5281/zenodo.18502434>
- Cohen-Steiner, D., Edelsbrunner, H., & Harer, J. (2007). Stability of persistence diagrams. *Discrete & Computational Geometry*, 37(1), 103-120.
- Grassberger, P., & Procaccia, I. (1983). Characterization of strange attractors. *Physical Review Letters*, 50(5), 346.
- Mandelbrot, B. (1967). How long is the coast of Britain? Statistical self-similarity and fractional dimension. *Science*, 156(3775), 636-638.
- Rucco, M., et al. (2016). Characterisation of the idiotypic immune network through persistent entropy. *Proceedings of ECCS 2014*, 117-128.
- Shannon, C.E. (1949). Communication theory of secrecy systems. *Bell System Technical Journal*, 28(4), 656-715.

- Ross, B.C. (2014). Mutual information between discrete and continuous data sets. *PLoS ONE*, 9(2), e87357.
- Anthropic. (2026). Sabotage Risk Report: Claude Opus 4.6. <https://anthropic.com/claude-opus-4-6-risk-report>
- Tishby, N., Pereira, F., & Bialek, W. (2000). The Information Bottleneck Method. *Proceedings of the 37th Allerton Conference*.