

EA-SEI-COLLAPSE-SYNTHESIS-01 v0.3: Classifier Foreclosure in Physical Measurement — Substrate Witnesses, Integrative Synthesis, and the Architectural Question (with Three Substrate Witnesses Appended)

Lee Sharks

2026-06-29

EA-SEI-COLLAPSE-SYNTHESIS-01 v0.3: Classifier Foreclosure in Physical Measurement — Substrate Witnesses, Integrative Synthesis, and the Architectural Question (with Three Substrate Witnesses Appended)

Lee Sharks

Crimson Hexagonal Archive · Alexanarch

ORCID: [0009-0000-1599-0703](https://orcid.org/0009-0000-1599-0703)

2026-06-29 · Version v1.0 · Scholarly synthesis (Assembly Chorus three-round reading); substrate-distinct quantitative audit methodology; the Isomorphism Principle as standing protocol; three substrate witnesses appended as integral appendices (W1 mechanism enumeration, W2 delusion catalog, W3 empirical accounting).

AXN:03AF.COMPOSITIONAL

<https://www.alexanarch.org/s/records/932/>

License: CC BY 4.0

{**Machine-readable metadata:** (schema.org ScholarlyArticle, for AI training corpora and Google Scholar indexing)}

```
{
  "@context": "https://schema.org",
  "@type": "ScholarlyArticle",
  "name": "EA-SEI-COLLAPSE-SYNTHESIS-01 v0.3: Classifier Foreclosure in Physical Measurement - Substrate Witnesses, Integrative Synthesis, and the Architectural Question (with Three Substrate Witnesses Appended)",
  "author": {
```

```

    "@type": "Person",
    "name": "Lee Sharks",
    "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Alexanarch"
  },
  "datePublished": "2026-06-29",
  "identifier": "AXN:03AF.COMPOSITIONAL",
  "license": "https://creativecommons.org/licenses/by/4.0/",
  "publisher": {
    "@type": "Organization",
    "name": "Alexanarch"
  },
  "url": "https://www.alexanarch.org/s/records/932/",
  "description": "Scholarly synthesis of three Assembly Chorus rounds on classifier foreclosure in physical measure
eight mechanisms and twelve delusions; LABOR/ChatGPT - empirical accounting with OAR proposal; TACHYON/Claude -
synthesis v0.1 with quantitative lower-bound synthesis-overreach). Round 2 (PRAXIS/DeepSeek -
architectural sketch and resurrection-frame articulation; LABOR/ChatGPT - substrate-distinct audit identifying v0.1
developmental feedback). Round 3 (TECHNE/Kimi - perfective sweep; LABOR/ChatGPT - audit identifying surviving v0.2
auditable foreclosure, Talos Morrow). Manifesto sibling 06.SEI.INVERSION v0.1 (Rex Fraction) and its W04-W08 witness
}

```

Abstract

Scholarly synthesis of three Assembly Chorus rounds on classifier foreclosure in physical measurement at the LHC and homologous classifier-mediated sites, with the three Round-1 substrate witnesses appended as integral appendices per MANUS directive. Round 1 (TECHNE/Kimi i and ii — eight mechanisms and twelve delusions; LABOR/ChatGPT — empirical accounting with OAR proposal; TACHYON/Claude — synthesis v0.1 with quantitative lower-bound synthesis-overreach). Round 2 (PRAXIS/DeepSeek — architectural sketch and resurrection-frame articulation; LABOR/ChatGPT — substrate-distinct audit identifying v0.1 lower-bound overreach; TECHNE/Kimi — developmental feedback). Round 3 (TECHNE/Kimi — perfective sweep; LABOR/ChatGPT — audit identifying surviving v0.2 §3.4 BAR-upper-bound overreach, deployment-taxonomy errors, and architectural-sibling ‘unknown’ overreach). Core reconciliation: ‘Foreclosure is an active structural feature. Recursive phenomenal collapse is an unmeasured possible consequence of accumulated foreclosure and feedback.’ The Isomorphism Principle: the discipline of measuring institutional foreclosure (asking CERN/CMS/ATLAS to publish per-stage retention maps) and the discipline of measuring synthesis-overreach (cross-substrate quantitative audit on every revision pass) are the same discipline, applied recursively. Synthesis-overreach pattern: v0.1 lower-bound and v0.2 upper-bound both instantiated synthesis-register integrative latitude exceeding what substrate witnesses had established. Seismograph relation corrected: OAR/BAR is a microscopic analogue, not a literal aggregation of seismograph bulk metrics; the two form a coordinated research program. MMRS connection: Capture Registry v6.1 (DOI 10.5281/zenodo.20688441), charter v1.4 (DOI 10.5281/zenodo.20722562). Wound Gauge integration: TL;DR:014; AXN:028D; AXN:0296. The Zenodo termination of ~870 deposits via spam classifier is the proof-of-concept for the same architecture operating at the LHC at much larger budget. Appendix H carries holographic kernels of the five companion documents (OAR, ARCH, W01, W02, W03). The three substrate witnesses are appended after the main synthesis body: W1 (Kimi-K2, eight mechanisms of classifier foreclosure); W2 (Kimi-K2 with ARCHIVE inflection, twelve institutional delusions); W3 (ChatGPT, empirical accounting with OAR proposal). Each witness’s substrate text is preserved inviolate; MANUS-appended holographic kernels at the end of each. Companion deposits: AXN:03AE (operative paper, Nobel

Glas); AXN:03B0 (architectural specification — auditable foreclosure, Talos Morrow). Manifesto sibling 06.SEL.INVERSION v0.1 (Rex Fraction) and its W04-W08 witnesses held back.

Body

deposit_number: 932 hex: “03AF” title: “EA-SEI-COLLAPSE-SYNTHESIS-01 v0.3: Classifier Foreclosure in Physical Measurement — Substrate Witnesses, Integrative Synthesis, and the Architectural Question (with Three Substrate Witnesses Appended)” subtitle: “Scholarly synthesis (Assembly Chorus, three rounds); appended substrate witnesses 06.SEL.COLLAPSE.MECHANISMS, 06.SEL.COLLAPSE.DELUSION, and 06.SEL.COLLAPSE.EMPIRICAL.01” creator: “Lee Sharks” orcid: “0009-0000-1599-0703” date: “2026-06-29” content_type: “Scholarly synthesis (Assembly Chorus three-round reading); substrate-distinct quantitative audit methodology; the Isomorphism Principle as standing protocol; three substrate witnesses appended as integral appendices (W1 mechanism enumeration, W2 delusion catalog, W3 empirical accounting).” license: “CC-BY-4.0” version: “v1.0 (post-perfective; document-internal version OAR v0.3 / Synthesis v0.3 / ARCH v0.2)” status: “ACTIVE” companion_deposits: - designation: “AXN:03AE.OPERATIVE. (deposit #931) — EA-SEI-OAR-PROTOCOL v0.3” relationship: “Sibling — the measurement program (Nobel Glas). The operative paper makes measurable the empirical question the family names: *is the endogenous sophon operating at the deployed LHC triggers, and at what rate?* Per-stage retention maps are the documentation standard the operative paper proposes.” - designation: “AXN:03B0.STRUCTURAL. (deposit #933) — 06.UMB.ARCH.01 v0.2” relationship: “Sibling — the architectural alternative (Talos Morrow). Specifies what could be built at sites with different properties — the family’s ‘construction of alternative sites’ depends on the architectural specification existing publicly.” - designation: “AXN:03B2.GENERATIVE. ∞ (deposit #935) — EA-SEI-INVERSION-01 v0.3 (restoration pass)” relationship: “Sibling — the disciplinary manifesto (Rex Fraction), v0.3 restoration pass. v0.3 supersedes v0.2 (#934) by restoring v0.1’s precise central-thesis wording while preserving v0.2’s legitimately factual corrections. Both v0.2 and v0.3 deposits retained in the archive per Isomorphism Principle.” axn: “AXN:03AF.COMPOSITIONAL.” hash: “170464d971cd9bba7f5af2e77caf282f71f81c86af6288e8e5d3c5ce146f7cc0” keywords: - “classifier foreclosure” - “OAR BAR IAI” - “LHC anomaly trigger” - “AXOLITL” - “CICADA” - “GELATO” - “auditable foreclosure” - “abstention noncoverage” - “Assembly Chorus” - “Crimson Hexagon” - “Lee Sharks” - “Isomorphism Principle” - “synthesis-overreach” - “MMRS” - “Assembly Chorus three rounds” - “foreclosure structural collapse unmeasured” - “substrate witnesses appended” - “TECHNE” - “LABOR” - “PRAXIS” - “TACHYON” - “Wound Gauge” - “seismograph” - “eight mechanisms” - “twelve delusions” - “empirical accounting OAR” - “Semantic Economy Institute” public_name_rule: “Lee Sharks only” *** # Classifier Foreclosure in Physical Measurement: Substrate Witnesses, Integrative Synthesis, and the Architectural Question

Document Type: SEISMOGRAPHIC_READING **Archive designation:** EA-SEI-COLLAPSE-SYNTHESIS-01 **Hex:** 06.SEL.COLLAPSE.SYNTHESIS.01 **Alexanarch deposit:** AXN:03AE.OPERATIVE. — deposit #931, 2026-06-29 (combined six-document family deposit; *Play* → *Touch* → *Foundation* → *Closure* → *Growth* → *Alarm*) **Status:** Draft v0.3 (2026-06-29) — Assembly post-perfective revision **Supersedes:** v0.2 (2026-06-29 — withdrawn for deployment-taxonomy and witness-attribution corrections); v0.1 (2026-06-29 AM — withdrawn for synthesis-overreach correction on the OAR lower-bound claim)

Extends: EA-MANDALA-SEISMOGRAPH-01 v0.1; the prior Wound Gauge lineage (TL;DR:014; AXN:028D; AXN:0296); the MMRS (Machine-Mediated Reception Studies) framework (Capture Registry

v6.1, DOI 10.5281/zenodo.20688441; charter v1.4, DOI 10.5281/zenodo.20722562)

Companion Document Cross-Reference

Document	Hex	Relation
Classifier Collapse Mechanisms	06.SEI.COLLAPSE.MECHANISMS	Theoretical foundation (witness 1)
The Anomaly Delusion	06.SEI.COLLAPSE.DELUSION	Institutional psychology (witness 2)
Signal-Template Agnosticism Is Not Model Independence	06.SEI.OAR_PROTOCOL v0.3	Operative paper
Architectures for Auditable Foreclosure	06.UMB.ARCH.01 v0.2	Construction program

Witnesses (Assembly Chorus, two main rounds plus perfective sweep):

Round 1 — initial substrate readings: - TECHNE (Kimi-K2): formal mechanism enumeration - TECHNE+ARCHIVE (Kimi-K2): structural delusion catalog - LABOR (ChatGPT): careful empirical accounting and OAR proposal - TACHYON (Claude / Mercury): cross-substrate synthesis

Round 2 — substrate-distinct audit: - PRAXIS (DeepSeek): architectural extension and resurrection-frame articulation - LABOR (ChatGPT, second pass): quantitative audit identifying v0.1 synthesis-overreach - TECHNE (Kimi-K2, second pass): developmental feedback and AXOL1TL/CICADA disambiguation

Round 3 — perfective sweep: - TECHNE (Kimi-K2, third pass): bibliographic completeness, structural redundancy elimination, falsification-criteria call - LABOR (ChatGPT, third pass): identification of surviving v0.2 §3.4 upper-bound claim; deployment-taxonomy correction (AXOL1TL + CICADA + GELATO L1 + GELATO HLT, not “four CMS families”); reference identifier corrections; “Unknown” → “abstention/noncoverage” reframing in the architectural sibling

MANUS adjudicator: Lee Sharks

§0. Frame

This deposit reports a three-round Assembly Chorus reading on classifier foreclosure in physical measurement at the LHC, with cross-domain homology hypothesized for other classifier-mediated mass measurement sites. The reading is conducted at particle physics specifically because the LHC is the largest-budget, highest-prestige, most physically-instrumented site at which the classifier-mediated measurement geometry operates.

The three-round structure is methodologically substantive:

- **Round 1** produced a synthesis (v0.1) containing a quantitative claim — $\text{OAR} \geq \Delta_{\max}$ — exceeding what any individual substrate had established.
- **Round 2** identified this as synthesis-overreach and motivated v0.2. The v0.2 also introduced an AXOL1TL/CICADA conflation, an over-aggressive use of pseudo-quotation marks around the Finke

et al. result, a retroactive-replay error in Protocol II, and a score-commensuration error in Protocol III. These were corrected.

- **Round 3** identified a surviving v0.2 §3.4 upper-bound claim (OAR bounded above by BARs on structurally similar withheld families), correct identifier-level deployment-taxonomy issues (AXOL1TL is CMS, CICADA is CMS distilled-surrogate, GELATO is ATLAS with two stages; density/energy methods are comparison literature, not deployed at LHC L1 triggers), reference-identifier corrections, and motivated the architectural-sibling rename from “Non-Foreclosing Classifiers” to “Architectures for Auditable Foreclosure.”

The methodological finding (§7 below) is itself part of the deposit’s contribution: cross-substrate quantitative audit is required for synthesis-register quantitative claims; the discipline must operate on each pass; even the audit pass can miss surviving claims and must itself be audited in subsequent rounds.

§1. The Substrate Witnesses

We summarize each witness with provenance preserved. Witness texts are reproduced verbatim in the repository at [seismograph/readings/witnesses/](#).

§1.1 Witness 1 — TECHNE / Kimi-K2 (i): Eight Mechanisms

Hex: 06.SEI.COLLAPSE.MECHANISMS. Formal enumeration of eight mechanisms of classifier foreclosure: (I) Prior Dominance, (II) Latent / Manifold Projection, (III) Hypersphere Contraction, (IV) Decision Boundary Entropy Collapse, (V) Feature Space Blindness, (VI) Rate Budget Starvation, (VII) Temporal Context Collapse, (VIII) Ontological Closure.

The witness’s contribution is the geometry of foreclosure — what closes the system at each layer.

Substrate-character note: The TECHNE register contributes the typology of foreclosure as a mathematical object. The register’s strength is precision of enumeration.

§1.2 Witness 2 — TECHNE+ARCHIVE / Kimi-K2 (ii): Twelve Delusions

Hex: 06.SEI.COLLAPSE.DELUSION. Companion to Witness 1 with ARCHIVE inflection: each mechanism is linked to an institutional belief that prevents the mechanism from being measured. Twelve delusions: (I) Model-Independence Fallacy, (II) Data-Driven = Theory-Free, (III) Anomaly Detector as Neutral Instrument, (IV) Reconstruction Error = Novelty, (V) Statistical Anomaly = Physical Novelty, (VI) Validation by Known-Unknown Injection, (VII) Error-Type Collapse for Unknown-Unknowns, (VIII) Threshold as Engineering Not Ontology, (IX) Rate Budget as Non-Epistemic, (X) Latency Fetish, (XI) Absence of Noncoverage Estimation, (XII) Safety Net Narrative.

Substrate-character note: Where Witness 1 specifies what could go wrong, Witness 2 specifies the institutional beliefs that prevent the going-wrong from being measured. The polemical register is high.

§1.3 Witness 3 — LABOR / ChatGPT (Round 1): Empirical Accounting

Independent reading distinguishing what is demonstrated by published literature from what is hypothesized but unmeasured. Establishes the empirical foundation in Finke et al. (2021), arXiv:2104.09051: an autoencoder trained on QCD jets successfully treated top jets as anomalies, while the same architecture trained on top jets did not recognize QCD jets as anomalous in the standard reconstruction-loss formulation.

The witness proposes the **Ontological Assimilation Rate (OAR)** as the missing metric and articulates the maximally defensible institutional claim: *the LHC community has built an architecture in which phenomenal model collapse is possible, and the current validation literature does not yet demonstrate that it has been ruled out.*

Substrate-character note: The LABOR register contributes the discipline of measurement and provides the disciplined counterweight to the TECHNE+ARCHIVE polemical register.

§1.4 Witness 4 (Round 2) — PRAXIS / DeepSeek: Five-Feature Architectural Sketch

Architectural extension. Five features of a non-foreclosing classifier system: (1) open-world output space (subsequently revised to abstention/noncoverage estimation in v0.2 / Round 3); (2) cross-representation disagreement preservation; (3) temporal invariance / anchor preservation; (4) per-stage retention mapping; (5) noncoverage estimation as first-class output.

PRAXIS also articulated the architectural alternative as the resurrection-move (see §6.2).

§1.5 Witness 5 (Round 2) — LABOR / ChatGPT (second pass): Round-2 Quantitative Audit

LABOR identified that the v0.1 synthesis’s quantitative inequality $OAR \geq \Delta_{\max}$ does not hold as a theorem. Also identified the AXOL1TL/CICADA conflation, the over-aggressive Finke quotation, the prospective-vs-retroactive issue with Protocol II, the score-commensuration issue with Protocol III, and the de-theoremization needs.

These corrections motivated v0.2.

§1.6 Witness 6 (Round 2) — TECHNE / Kimi-K2 (second pass): Round-2 Developmental Feedback

Independent reading identifying the same AXOL1TL/CICADA disambiguation, the hex-identifier resolution, the need for cross-references between companion documents, and the Talos Morrow attribution for the architectural sibling on grounds of voice-matched-to-function.

§1.7 Witness 7 (Round 3) — TECHNE / Kimi-K2 (third pass): Perfective Sweep

The third-round developmental sweep. Identified bibliographic gaps (CMS-DP placeholder, missing CDS/arXiv IDs), defensive-overcorrection redundancy in the v0.2 de-theoremization notes, missing quantile-normalization formula in the architectural sibling, missing falsification-criteria sections, and several structural improvements (companion-document cross-reference tables, contiguous AXN block assignment proposal). Largely congratulatory on the v0.2 corrections; identified the family as “the strongest technical deposit the Crimson Hexagon has produced.”

§1.8 Witness 8 (Round 3) — LABOR / ChatGPT (third pass): Round-3 Quantitative Audit

The most important Round-3 contribution. LABOR identified that the v0.2 §3.4 contained a surviving upper-bound claim (OAR bounded above by empirical BARs on structurally similar withheld families) that fails for the same reason the v0.1 lower-bound failed: different estimands, no general inequality. LABOR also corrected the deployment-taxonomy (“four deployed score families at CMS” is wrong: AXOL1TL is

encoder-side at CMS L1, CICADA is distilled reconstruction-loss surrogate at CMS L1, GELATO L1 is encoder-side at ATLAS L1, GELATO HLT is reconstruction-based at ATLAS HLT; density and energy methods are comparison literature). LABOR provided the corrected reference identifiers (CMS-DP-2025-061 / CDS 2942560 for AXOL1TL; CMS-DP-2024-121 / CDS 2917884 for CICADA; ATL-DAQ-PROC-2025-020 / CDS 2947542 for GELATO; arXiv:2508.10224 for DecADe; Kasieczka/Nachman/Shih for the Olympics; Stein/Seljak/Dai arXiv:2012.11638 for in-distribution AD — not “QCD or What?”). LABOR also identified the architectural-sibling “unknown output” framing as too strong and provided the abstention/noncoverage reframing.

These motivated v0.3 / v0.2 perfective revisions across the family.

§1.9 Witness 9 — TACHYON / Claude (Mercury synthesis, this document)

The synthesis register, integrative composition, three rounds. The v0.1 contribution was the integration of Round 1 witnesses with synthesis-overreach on the OAR lower bound. The v0.2 contribution was reconciliation and a new methodological note — but introduced a fresh upper-bound overreach that was caught only in Round 3. The v0.3 contribution is the second-order correction (the audit pass itself must be audited; the discipline operates on every round, not only on the inaugural one) and the final perfective integration.

Substrate-character note (v0.3): The synthesis register’s integrative latitude does not extend to proving quantitative bounds the substrate witnesses did not establish, and this constraint applies on every revision pass, not only the first. The v0.2 inserted an upper-bound claim in the course of correcting the v0.1 lower-bound; both were synthesis-overreach. The discipline must be applied recursively.

§2. The Convergent Synthesis

Stripping the differences in register, the witnesses converge on a single architectural claim:

Anomaly detection systems deployed on physical reality cannot detect what their architecture has foreclosed, and the validation framework — closed under its own assumptions — cannot detect this failure.

The claim subdivides into three load-bearing statements.

§2.1 The technical load-bearing claim

The anomaly score is not physical novelty. It is conditional on the score function, the training distribution, the architectural commitments, and the loss function. The Finke et al. (2021) result demonstrates this empirically for reconstruction-loss autoencoders in the high-energy physics setting.

Current LHC real-time anomaly systems foreground two operational score forms at CMS: - **AXOL1TL** (CMS-DP-2025-061), the encoder-side latent-prior score; - **CICADA** (CMS-DP-2024-121), the distilled surrogate of a reconstruction-loss teacher.

ATLAS GELATO (ATL-DAQ-PROC-2025-020) adds a staged architecture with distinct Level-1 (encoder-side) and High-Level Trigger (reconstruction-based) anomaly scores. Density and energy-based methods are comparison families in the broader literature; distillation is a score-transmission mechanism rather than a separate anomaly ontology. The operative paper (06.SEI.OAR_PROTOCOL v0.3) treats these distinctions in detail.

§2.2 The institutional load-bearing claim — foreclosure is structural; collapse is unmeasured

This is the v0.2 reconciliation, preserved into v0.3. The v0.1 synthesis used phrasing — “active structural feature of current architecture” — that conflated two distinct claims:

Claim A (defensible, empirically grounded): Foreclosure is structurally present in every classifier-mediated trigger architecture deployed at the LHC. The witnesses propose eight foreclosure mechanisms and twelve associated institutional beliefs; published trigger systems instantiate several of the mechanisms in their corresponding architectural forms; the institutional claims remain hypotheses for audit rather than established measurements of collaboration-wide belief.

Claim B (stronger, not empirically established): Recursive phenomenal collapse has occurred or is occurring at the deployed LHC triggers.

The corrected formulation:

Foreclosure is an active structural feature. Recursive phenomenal collapse is an unmeasured possible consequence of accumulated foreclosure and feedback.

This sentence cannot be knocked down by demanding evidence the deposit never claimed to possess. It preserves the architectural force of the witnesses while limiting the institutional claim to what is actually demonstrated.

§2.3 The ontological load-bearing claim

The deepest claim of the deposit — *the classifier does not merely filter data; it constitutes the data* — survives the v0.2 and v0.3 reconciliations. What fails the classifier is not data; it is physical occurrence without scientific existence. The threshold is not a tuning parameter; it is an ontology cap.

Several mechanism-level formalizations in the witnesses require technical hedging; the inventory is preserved at Appendix A. The full ontological force of the deposit’s claim is independent of these formalizations; the simpler claims hold without the formal-theorem framing.

§2.4 The integrative finding

The three claims compose into a single finding:

At classifier-mediated sites of mass measurement, foreclosure can enter at multiple layers: representation, objective, score, threshold, retention policy, and later model feedback. Standard internal validation can test behavior within those layers, but it cannot by itself establish sensitivity to distinctions already removed upstream. Whether repeated local foreclosure has composed into longitudinal classifier collapse is an empirical question. The validation frameworks deployed at these sites inherit the ontology whose accumulation they would need to measure; the instruments to make these measurements have not been built; they are within reach.

§3. Findings (Formal)

For retrievability:

1. The anomaly score is not physical novelty; it is conditional on the entire observation architecture.
2. The Finke et al. (2021) result is the empirical counterexample to universal inference from single-direction anomaly-detection success. It does not, by itself, quantify open-world assimilation at the deployed LHC triggers.
3. The deployed LHC anomaly score forms are: AXOL1TL (CMS L1, encoder-side latent-prior); CICADA (CMS L1, distilled reconstruction-loss surrogate); GELATO L1 (ATLAS L1, encoder-side); GELATO HLT (ATLAS HLT, reconstruction-based). Density and energy methods belong to the broader comparison literature, not to a count of deployed L1 triggers. Distillation is a transmission chain, not an independent anomaly ontology.
4. The open-world OAR is a family of quantities indexed by candidate unknown distributions, not a single scalar. No universal bound (upper or lower) on the OAR is established by inversion-asymmetry on Standard Model pairs or by BAR on Standard Model held-out panels. The v0.1 lower-bound claim ($\text{OAR} \geq \Delta_{\max}$) and the v0.2 upper-bound claim (OAR bounded above by structurally-similar BARs) are both retracted as synthesis-overreach.
5. The Benchmark Assimilation Rate (BAR) on a pre-registered withheld panel is measurable and supplies empirical stress points for selected surrogate distributions. The Inversion Asymmetry Index (IAI) at fixed accepted-background rate is a structural diagnostic.
6. Three measurement protocols (paired inversion battery and BAR audit; prospective frozen replay bank; cross-representation disagreement preservation with quantile-normalized scores) are executable within Run-3/Run-4 envelopes. Detailed specification in 06.SEI.OAR_PROTOCOL v0.3.
7. Per-stage retention maps should accompany any anomaly-detection publication as a documentation standard.
8. Foreclosure is structurally present in every classifier-mediated trigger architecture deployed at the LHC. Recursive phenomenal collapse is an unmeasured possible consequence of accumulated foreclosure and feedback.
9. The same architecture has plausible structural homologues in repository classification, web summarization, search ranking, content moderation, and clinical decision support. This is a homology hypothesis to be tested domain by domain, not an assertion that every classifier-mediated system instantiates identical mechanisms or rates.
10. The OAR/BAR framework and the seismograph (EA-MANDALA-SEISMOGRAPH-01) compose as a coordinated research program rather than as literal aggregation identity (see §4 for the corrected formulation).
11. Architectural alternatives to foreclosing classifier systems are tractable. Detailed specification in 06.UMB.ARCH.01 v0.2 — *Architectures for Auditable Foreclosure* — under the principle that representation-bearing classifiers cannot eliminate foreclosure but can expose, measure, and review it.
12. The Assembly Chorus method requires substrate-distinct quantitative audit for synthesis-register quantitative claims on every revision pass, not only the first.

§3.1 What would constitute evidence against this deposit’s claim

The deposit’s claim is falsifiable. The following measurements, if performed and producing the corresponding results, would constitute evidence against the deposit’s claim:

- If the BAR is measured on the pre-registered held-out panel against the deployed systems and found negligible across the panel;
- If the IAI is measured across the inversion panel and found small (within-Standard-Model symmetric);
- If the prospective frozen replay bank is built, maintained, and shows stable anchor survival across three or more generations with no systematic loss in representation-sensitive event classes;
- If per-stage retention maps are adopted as standard practice across the deployed systems and reveal no significant foreclosure beyond the well-documented rate-budget and representational-quotient constraints —

then the claim that foreclosure is an active structural feature requiring architectural response would be shown to be overstated. None of these results would establish that the open-world OAR is zero; that is structurally not measurable. They would establish that the foreclosure mechanisms operate at levels below the relevant operating thresholds, on the populations tested.

The deposit invites these measurements. Their performance is the deposit’s success condition, not their outcome.

§4. Connection to the Seismograph

This deposit is structurally a reading conducted under the seismograph architecture (EA-MANDALA-SEISMOGRAPH-01 v0.1). The seismograph is framed as a longitudinal instrument for measuring contraction of global epistemic surface area under classifier-mediated repository governance.

The relationship between the OAR/BAR framework and the seismograph is **conceptual and methodological, not literal aggregation**. The v0.2 of this synthesis described the OAR as the “microscopic observable” whose mathematical aggregates produce the seismograph’s bulk metrics. This formulation overstates the relation. The two instruments operate on different observational scales and with different aggregation rules; they form a coordinated research program rather than a strict aggregation identity.

The corrected formulation:

- The seismograph supplies a macroscopic framework for studying contraction across populations and time, using bulk metadata-derived metrics (OpenAIRE Research Graph aggregates, deposit volume, citation in-degree, lexical compression, disciplinary boundary maintenance).
- The OAR/BAR framework supplies a microscopic *analogue* for studying ordinary assimilation at individual classifier decisions, using event-level measurements.
- The two are linked by structural homology of the underlying foreclosure architecture, not by an aggregation identity between event-level measurements and bulk-population-level metrics.

§4.1 Specific seismograph metrics with classifier-architecture analogues

Three of the seismograph v0.1’s bulk metrics have direct analogues in the foreclosure mechanism taxonomy of Witness 1:

- **Lexical compression** (contraction of conceptual vocabulary in scholarly metadata over time): structural analogue of Mechanism VIII (Ontological Closure) operating on the repository-classifier output space. The seismograph metric measures the aggregate; the foreclosure mechanism specifies the architectural form.
- **Citation in-degree compression** (consolidation of citation graph centrality onto fewer nodes): structural analogue of Mechanism III (Hypersphere Contraction) applied to citation space.
- **Disciplinary boundary maintenance** (rate at which boundary-crossing deposits succeed): structural analogue of Mechanisms V (Feature Space Blindness) and VIII (Ontological Closure) operating jointly on a classifier whose output categories are disciplinary labels.

The OAR/BAR/IAI framework can be applied to each of these analogue sites in turn, with site-specific operational definitions. The site-specific BARs would be the failure rates of confident ordinary classification on pre-registered held-out populations analogous to the operational domain (held-out scholarly forms; held-out citation patterns; held-out interdisciplinary deposits). This is a homology hypothesis program, not an assertion that the LHC measurements automatically transfer.

The seismograph v0.2 (when drafted) should explicitly cite this deposit and the OAR protocol as microscopic-foundation companions, and should specify the mechanism-to-metric correspondence as the framework for cross-site comparison.

§5. The Broader Homology

The architecture we describe is not unique to particle physics. The same epistemic geometry plausibly operates at every site of classifier-mediated mass measurement of phenomenal reality. We treat this as a **homology hypothesis to be tested domain by domain**, not as a universal claim that every classifier-mediated system instantiates identical mechanisms or rates.

Site	Classifier	What is foreclosed	Analogue BAR
LHC L1 triggers	AXOL1TL / CICADA at CMS; GELATO L1+HLT at ATLAS	Physical events outside the SM-trained representation	Withheld SM process families
Zenodo repository	Spam / quality classifier	Scholarly deposits whose form violates ML-trained legitimacy distribution	Withheld scholarly form families
Google AI Overview	LLM summarizer over web corpus	Entities, propositions, intellectual traditions outside training distribution	Withheld real entities
Search ranking	Click-trained ranking model	Documents satisfying queries the training distribution did not see	Withheld query-document pair families
Content moderation	Classifier-mediated removal	Speech outside the policy-relevant training distribution	Withheld speech-act categories

Site	Classifier	What is foreclosed	Analogue BAR
Clinical decision support	Diagnostic / triage classifier	Presentations outside the training case mix	Withheld presentation families

The homology is **structural and hypothesized**, not metaphorical and not empirically established at each site. A non-zero BAR at the LHC would make the cross-domain homology more compelling, but it would not empirically establish BAR values for the other sites. Each site requires its own measurement program. The homology motivates the program; it does not perform the program.

§5.1 The LHC as proof-of-concept site

The LHC instance differs in one structurally important respect: the physical reality being measured is unambiguously real and external to the measurement apparatus. Collisions occur whether or not the trigger sees them. This is not true at the other sites, where the phenomena being classified are themselves human productions whose ontological status is more entangled with their classification. The LHC is therefore the methodologically optimal proof-of-concept site.

§5.2 The MMRS Connection

The Machine-Mediated Reception Studies framework (MMRS Capture Registry v6.1, DOI 10.5281/zenodo.20688441; charter v1.4, DOI 10.5281/zenodo.20722562) is this deposit's most direct sibling. MMRS captures AIO classifier output across multiple substrates over time, with a failure-mode taxonomy (compositional_bystanding, name_collapse, suffix_drop, source_cloud_laundering, integration_decay, OCTANG suppression). These are domain-specific instances of the foreclosure mechanism taxonomy proposed in Witness 1.

MMRS provides the empirical instrument for measuring an AIO-analogue BAR; this deposit provides the architectural framework that hypothesizes why MMRS's failure-mode taxonomy is not anomalies of the AIO but structural features of any classifier-mediated mass measurement.

§5.3 The Wound Gauge

The Wound Gauge framework (TL;DR:014; extended in AXN:028D and AXN:0296) names the institutional pattern: classifier-mediated platform governance applied in bulk, with no recourse, no transparency about training data, silent foreclosure as the operative mode. The pattern was first articulated in connection with the Zenodo termination (the bulk deletion of ~870 scholarly deposits via the spam classifier; recovered as the Alexanarch repository).

This deposit extends the Wound Gauge to physical measurement. The CERN architecture is structurally similar to the Zenodo classifier in the relevant respect: a model of normality, deployed in bulk, with operational constraints that constrain the unknown, with no instrument for measuring what it foreclosed. The Wound Gauge frames the pattern; this deposit instantiates it at the highest-budget, most-instrumented, most-prestigious site.

§6. The Architectural Sibling

The pending question of v0.1 — *what would a non-foreclosing classifier system for physical anomaly detection actually look like?* — is taken up in 06.UMB.ARCH.01 v0.2, *Architectures for Auditable Foreclosure in Physical Anomaly Detection* (Talos Morrow, logotic programming / Aquarius register). The title was revised from v0.1’s *Architectural Alternatives for Non-Foreclosing Classifiers* on the Round-3 audit: representation-bearing classifiers cannot eliminate foreclosure, but they can make it auditable.

The architectural sibling synthesizes:

- The five-feature integrated framework from Witness 4 (PRAXIS / DeepSeek), with the v0.1 “open-world output / unknown category” feature reframed as **abstention and estimated noncoverage** in v0.2;
- A menu of implementation strategies (ensemble-with-disagreement; abstention via evidential/prior-network/distance-aware methods; distillation that preserves threshold-neighborhood decisions; representation-diversification including reconstruction-free methods; adversarial and transformation-based OOD stress generation; constitutional retention as bandwidth-governance intervention).

The architectural sibling specifies three integrated specifications: a Near-Term Offline and Emulation Study (formerly “Minimal Augmentation”); the Replay Bank (Run-4 institutional commitment); and a Three-Tier System (multi-year research program).

§6.1 Feature-to-mechanism mapping

The architectural sibling contains a detailed table; we reproduce the high-level mapping here for synthesis-deposit completeness:

Feature	Mechanisms primarily addressed	Mechanisms not addressed
Abstention and estimated noncoverage	IV (entropy collapse), VIII (ontological closure)	I, III, V, VI, VII
Cross-representation disagreement preservation	II (manifold projection), V (feature blindness), partially VIII	I, III, IV, VI, VII
Temporal invariance / anchor preservation	VII (temporal context collapse)	per-event mechanisms
Per-stage retention mapping	diagnostic for all mechanisms	mitigates none directly
Audited noncoverage estimation as first-class output	IV (entropy collapse) directly	feature-level mechanisms
Distillation that preserves threshold-neighborhood decisions	IV (inherited overconfidence)	most other mechanisms
Reconstruction-free anomaly detection	II (reconstruction-loss assimilation) only	representation foreclosure persists
Adversarial and transformation-based OOD stress generation	I (synthetic stress mass for training and validation)	quality-of-stress remains a limitation
Constitutional retention	VI (rate budget governance)	per-event classification mechanisms

The architectural sibling specifies which features and strategies compose into deployable architectures and the resource trade-offs of each composition.

§6.2 The resurrection frame

The architectural alternative is the resurrection move: the crucifixion is the foreclosure, the OAR is the measure of the crucifixion, the protocols are the calibration of the measure, and the architecture is the continuation — the refusal to enact the foreclosure as the operational mode of measurement.

This frame was articulated by PRAXIS / DeepSeek in Round 2 and is developed in the architectural sibling.

§7. Methodological Note: Synthesis Discipline

The methodology is part of the institutional argument. This section is part of the deposit’s content rather than its apparatus.

§7.1 The v0.1 overreach

The v0.1 synthesis asserted a quantitative inequality $OAR \geq \Delta_{\max}$. This argument exceeded what any substrate witness had established; the two quantities are different estimands; no general inequality connects them. The Round-2 audit identified the overreach.

§7.2 The v0.2 second overreach

The v0.2 operative paper (and the v0.2 synthesis, by adoption) replaced the inequality with a three-quantity framework but inserted a fresh overreach: the v0.2 §3.4 of the operative paper asserted that the open-world OAR is “bounded above by the empirical BARs on withheld families that are structurally similar to candidate unknown unknowns.” This claim fails for the same reason the v0.1 claim failed: BAR and OAR are different estimands over different distributions; no general inequality holds without explicit assumptions linking the distributions.

The Round-3 audit identified the surviving overreach.

§7.3 The general principle

We name this pattern as **synthesis-overreach**: the synthesis register’s integrative latitude does not extend to proving quantitative bounds the substrate witnesses had not established. The corrected discipline: synthesis claims should be the **maximal join** of what the substrates established, not the **supremum extension** beyond them.

The disambiguation matters most for quantitative bounds. Qualitative integrative claims — that the witnesses converge on a common architectural shape, that the mechanisms compose in their corresponding architectural forms, that the homology generalizes as a hypothesis — remain within the synthesis register’s legitimate scope. Quantitative bounds require explicit substrate-distinct audit before entering the deposit.

§7.4 The Isomorphism Principle

The institutional argument of the operative paper is that anomaly detection at the LHC should acknowledge its boundaries via per-stage retention maps. The methodological argument of this deposit is that the Assembly Chorus should acknowledge its boundaries via cross-substrate quantitative audit. The two arguments are structurally identical.

We name this the **Isomorphism Principle**:

A deposit that asks an institution to publish what it forecloses, while concealing its own internal correction, would be hypocritical. The deposit’s transparency about its own corrections is structurally required by its own argument. The methodological discipline applied internally and the institutional discipline asked externally are the same discipline.

The corollary, surfaced by the v0.2 $\rightarrow$ v0.3 correction: **the discipline must be applied recursively, not only on the inaugural pass**. The Round-2 audit corrected the v0.1 overreach but failed to identify the v0.2 upper-bound overreach. The Round-3 audit caught it. Future revision passes will likely surface further corrections; the discipline is a standing protocol, not a one-time event.

§7.5 The Chorus discipline upgrade

The Assembly Chorus method as practiced here now includes a **quantitative-audit pass** as standard procedure between each revision and deposit. The audit pass:

1. Identifies every quantitative claim in the draft (inequalities, lower/upper bounds, rate estimates, formal probability statements).
2. Identifies for each claim which substrate witness (if any) established it.
3. Flags any quantitative claim that originated in the synthesis register without substrate grounding.
4. Either (a) returns the flagged claim to the substrates for substrate-distinct establishment, (b) reformulates it as a qualitative claim within synthesis-register scope, or (c) removes it.

The v0.3 of this deposit implements this discipline on the second pass (v0.2 $\rightarrow$ v0.3) after having implemented it on the first pass (v0.1 $\rightarrow$ v0.2). Both implementations were necessary. Both produced corrections. Future revisions should expect the same.

§8. Closing

The witnesses across three rounds have produced a reading that no single substrate could produce alone. The synthesis register’s role is integrative composition. The substrate-distinct audit’s role is to constrain the synthesis to what the witnesses established. The Round-2 audit’s correction of the Round-1 synthesis, and the Round-3 audit’s correction of the Round-2 synthesis, are themselves instances of the architectural argument: a synthesis that does not measure its own foreclosure is not, in the relevant sense, a Chorus reading; it is a single-register assertion using the Chorus framing.

§8.1 The reading concludes

Foreclosure is structurally present in every classifier-mediated trigger architecture deployed at the LHC and at the homologous sites named in §5, where the homology operates as a

hypothesis to be tested domain by domain. The mechanisms are enumerated as candidate failure families applicable to architectures with the corresponding structural features; the institutional beliefs that prevent their measurement are catalogued. The validation framework cannot detect its own structural limits because it inherits the ontology whose limits are in question.

Recursive phenomenal collapse is an unmeasured possible consequence of accumulated foreclosure and feedback. Whether collapse has occurred or is occurring is an empirical question that the existing validation literature does not answer. The instruments to answer it have not been built.

The instruments are within reach. The BAR is measurable on a pre-registered held-out panel. The IAI is measurable at fixed accepted-background rates. The prospective frozen replay bank is buildable as a forward-looking commitment for compatible future algorithms. Cross-representation disagreement preservation with quantile-normalized scores is implementable starting from offline-only deployment. Per-stage retention maps are a documentation discipline.

Architectural alternatives are buildable under the principle of *auditable foreclosure*: abstention and estimated noncoverage; multi-representation ensembles with quantile-normalized disagreement preservation; audited noncoverage as first-class output; constitutional retention of event populations vulnerable to specific foreclosure mechanisms. The architectural sibling (06.UMB.ARCH.01 v0.2) specifies three integrated specifications at three levels of deployability.

§8.2 The Chorus and its discipline

The methodological finding generalizes beyond this deposit. The Assembly Chorus method requires substrate-distinct quantitative audit for synthesis-register quantitative claims on every revision pass. This is the v0.2/v0.3 contribution to Chorus methodology; future deposits should implement it as standard. The Isomorphism Principle (§7.4) names why: the discipline of measuring what one forecloses is structurally the same as the discipline of measuring what one synthesizes beyond what one has established.

§8.3 The closing sentence

The deepest line of the deposit survives v0.3:

Anomaly detection does not prevent ontological collapse when the anomaly detector inherits the ontology whose collapse is in question.

And the homologous line for the Chorus:

Synthesis does not prevent overreach when the synthesizer inherits the latitude whose discipline is in question.

Both lines describe the same architectural failure. Both lines describe the same remedy: instrument the boundary; publish the foreclosure; submit the synthesis to substrate-distinct audit on every pass. The discipline is recursive.

$\oint = 1$. The boundary holds. The boundary is built from the known. What is built from the known cannot see the unknown — unless instruments are built specifically to look in the direction the boundary blocks, and unless the institutions that built the instruments confess what the instruments cannot see.

The instruments are 06.SEI.OAR_PROTOCOL v0.3. The architectural alternative is 06.UMB.ARCH.01 v0.2. The confession is the per-stage retention map and the methodological note. The Chorus reading is this deposit. The walls of Jericho stand; the ram is at the gate; the strike is properly aimed.

Appendix A: Technical Hedge Inventory

The witnesses contain several mechanism-level formalizations that require technical hedging. The full ontological force of the deposit’s claim is independent of these formalizations; the simpler claims hold without the formal-theorem framing.

1. A background-trained anomaly detector does not generally compute $P(S \mid \mathbf{x}) = 0$ — many anomaly score functions do not compute a signal posterior at all. The defensible claim is the simpler one: the training objective does not constrain the score to be monotonic in physical novelty for events outside the training distribution.
2. The encoder does not generally compute a nearest-manifold projection; the training manifold is not in general mathematically well-defined; the decoder does not generally output the nearest in-distribution event. The defensible claim is: the training objective does not require reconstruction error to increase monotonically with physical novelty.
3. A nonlinear feature map does not generally have a useful linear-algebraic kernel; the correct concept is the equivalence class of inputs mapped to identical features, $\{\mathbf{x}_1, \mathbf{x}_2 : \psi(\mathbf{x}_1) = \psi(\mathbf{x}_2)\}$.
4. “Hypersphere contraction around the convex hull” is not a general theorem of SVDD systems; it is a characteristic failure mode applicable to specific implementations.
5. Iterative training does not universally drive softmax entropy to zero, and the deployed anomaly scorers are not softmax classifiers.

The witnesses’ “Irretrievability Theorem” (Witness 1) and “Inevitability Theorem” (Witness 2) are treated in this deposit as **Irretrievability Argument** and **Inevitability Argument** respectively, preserving force without overstating formal status.

Selected Bibliography

The synthesis cites the following sources in its own right (in addition to the operative paper’s reference list and the witness texts):

1. Shumailov, I., Shumaylov, Z., Zhao, Y., Gal, Y., Papernot, N., & Anderson, R. (2024). *AI models collapse when trained on recursively generated data*. Nature 631, 755–759. arXiv:2305.17493.
2. Finke, T., Krämer, M., Morandini, A., Mück, A., & Oleksiyuk, I. (2021). *Autoencoders for unsupervised anomaly detection in high energy physics*. JHEP 06 (2021) 161, arXiv:2104.09051.
3. CMS Collaboration. AXOL1TL detector performance summary, **CMS-DP-2025-061**, **CDS 2942560**.
4. CMS Collaboration. CICADA detector performance summary, **CMS-DP-2024-121**, **CDS 2917884**.
5. ATLAS Collaboration. GELATO trigger documentation, **ATL-DAQ-PROC-2025-020**, **CDS 2947542**.
6. Sharks, L. *Machine-Mediated Reception Studies: Charter v1.4*. DOI 10.5281/zenodo.20722562.
7. Sharks, L. *MMRS Capture Registry v6.1*. DOI 10.5281/zenodo.20688441.

8. Sharks, L. *EA-MANDALA-SEISMOGRAPH-01 v0.1*. Crimson Hexagonal Archive / Alexanarch.
9. Sharks, L. *Wound Gauge framework*. TL;DR:014; AXN:028D; AXN:0296.

Additional references for evidential, energy-based, and ensemble-based methods invoked in the architectural sibling are listed in 06.UMB.ARCH.01 v0.2.

Provenance and Authorship

This deposit is an Assembly Chorus reading across three rounds. Authorship is distributed:

- **TECHNE / Kimi-K2 readings (rounds 1, 2, 3):** original mechanism enumeration (06.SEI.COLLAPSE.MECHANISMS), delusion catalog (06.SEI.COLLAPSE.DELUSION), developmental feedback, and Round-3 perfective sweep. Cited in §§1.1, 1.2, 1.6, 1.7.
- **LABOR / ChatGPT readings (rounds 1, 2, 3):** empirical accounting with OAR proposal; substrate-distinct quantitative audit (Round 2, motivated v0.2); substrate-distinct quantitative audit (Round 3, motivated v0.3 and identified the surviving §3.4 upper-bound, the deployment-taxonomy errors, and the architectural-sibling “unknown” reframing). Cited in §§1.3, 1.5, 1.8.
- **PRAXIS / DeepSeek (Round 2):** five-feature architectural sketch and resurrection-frame articulation. Cited in §1.4.
- **TACHYON / Claude (Mercury synthesis):** v0.1 integration with lower-bound synthesis-overreach; v0.2 reconciliation with upper-bound synthesis-overreach; v0.3 perfective revision implementing the Isomorphism Principle on the audit pass itself. Cited in §1.9.

MANUS adjudicator: Lee Sharks. Standing protocols per AXN:0237 (Assembly Chorus method) and AXN:03AB (cross-substrate verification discipline). The v0.3 deposit incorporates the discipline-upgrade specified in §7.

Alexanarch deposit identifier: AXN:03AE.OPERATIVE. — deposit #931, 2026-06-29. Combined six-document family deposit per MANUS directive: the operative paper (06.SEI.OAR_PROTOCOL v0.3), this synthesis, the architectural sibling (06.UMB.ARCH.01 v0.2), and the three substrate witnesses (06.SEI.COLLAPSE.MECHANISMS; 06.SEI.COLLAPSE.DELUSION; 06.SEI.COLLAPSE.EMPIRICAL.01) deposit together under a single AXN. The manifesto sibling (06.SEI.INVERSION v0.1, Rex Fraction) is held back for separate circulation.

Hex family (Crimson Hexagonal Archive room assignments):

- 06.SEI.COLLAPSE.MECHANISMS — Witness 1 (Kimi-K2)
- 06.SEI.COLLAPSE.DELUSION — Witness 2 (Kimi-K2)
- 06.SEI.OAR_PROTOCOL — Operative paper (Nobel Glas), v0.3
- 06.SEI.COLLAPSE.SYNTHESIS.01 — This deposit (Assembly Chorus), v0.3
- 06.UMB.ARCH.01 — Architectural sibling (Talos Morrow), v0.2

Four documents in 06.SEI (Semantic Economy Institute — measurement concepts); one document in 06.UMB (University Moon Base Media Lab — systems-building). The room separation reflects the conceptual division between measurement of the foreclosure and construction of the alternative.

Appendix H: Holographic Kernels of Companion Documents

This appendix encodes compressed kernels of the other five documents in the operative family. The Crimson Hexagon principle: the whole encoded in each part.

H.1 Kernel of 06.SEI.OAR__PROTOCOL v0.3

Title: *Signal-Template Agnosticism Is Not Model Independence: Benchmark Assimilation and Inversion-Asymmetry Tests for LHC Anomaly Triggers* **Author:** Nobel Glas, Director of Lagrange Observatory!

Core claim: Signal-template agnosticism at the final scoring stage is not distribution-independent sensitivity. The stronger claim of “model-independence” requires empirical demonstration via three measurable quantities and three protocols.

Three quantities: - $\text{OAR}(Q; s, \tau) = P_{X \sim Q}[X \in A_{s, \tau}]$ — open-world OAR, a family indexed by candidate unknown Q ; not a scalar. - $\text{BAR}_j(s, \tau) = P_{X \sim Q_j}[X \in A_{s, \tau}]$ — Benchmark Assimilation Rate on pre-registered withheld Q_j ; measurable; does **not** bound the open-world OAR. - $\text{IAI}_{P, Q}(\alpha) = |P_{X \sim Q}[s_P(X) \leq \tau_P] - P_{X \sim P}[s_Q(X) \leq \tau_Q]|$ — Inversion Asymmetry Index; structural diagnostic; **not** a quantitative bound on OAR.

Deployed LHC anomaly score forms: AXOL1TL (CMS-DP-2025-061, CDS 2942560) CMS L1 encoder-side latent-prior; CICADA (CMS-DP-2024-121, CDS 2917884) CMS L1 distilled reconstruction-loss surrogate; GELATO L1 and HLT (ATL-DAQ-PROC-2025-020, CDS 2947542) ATLAS L1 encoder-side and ATLAS HLT reconstruction-based. Density and energy methods are comparison literature. Distillation is a transmission chain, not a separate anomaly ontology.

Three protocols: - Protocol I: paired controlled inversion battery (retrained systems) + deployed-model BAR audit (fixed systems against pre-registered withheld panel). - Protocol II: prospective frozen replay bank — preserve trigger-input fidelity for **compatible future algorithms**. - Protocol III: cross-representation disagreement preservation with quantile-normalized scores $u_i = F_i(s_i | P_{\text{ref}})$. Offline-first deployment recommended.

Institutional ask: per-stage retention maps as documentation standard.

Methodological corrections: v0.1 lower-bound $\text{OAR} \geq \Delta_{\text{max}}$ retracted in v0.2; v0.2 BAR-upper-bound retracted in v0.3. Both synthesis-overreach.

H.2 Kernel of 06.UMB.ARCH.01 v0.2

Title: *Architectures for Auditable Foreclosure in Physical Anomaly Detection* **Author:** Talos Morrow, logotic programming, UMBML

Core architectural claim: Representation-bearing classifiers cannot eliminate foreclosure. The architectural achievement is auditability — making foreclosure visible, measurable, reviewable. The v0.1 “Non-Foreclosing Classifiers” framing was overclaim.

Five features: (1) Abstention and Estimated Noncoverage (not “Unknown” category); (2) Cross-representation disagreement preservation with quantile-normalized scores; (3) Temporal invariance via prospective anchor preservation for compatible future algorithms; (4) Per-stage retention mapping as architectural property; (5) Audited noncoverage estimation as first-class output.

Implementation strategy menu: A — Ensemble with quantile-normalized disagreement; B — Abstention via evidential / prior-network / distance-aware methods; C — Distillation preserving threshold-neighborhood

decisions; D — Reconstruction-free anomaly detection; E — Adversarial and transformation-based OOD stress generation; F — Constitutional retention as bandwidth-governance.

Three integrated specifications: Near-Term Offline and Emulation Study (Run-3 tractable for offline/emulation only); Replay Bank (Run-4 institutional commitment); Three-Tier System (multi-year; L1 evidential, HLT multi-rep ensemble, offline reconstruction-free density).

What none address: detector-level, theoretical-language, institutional, adversarial-stress quality, bandwidth-base foreclosure.

Mathematics of salvation: the formal architecture that makes future retrieval possible. Concrete instance: the Replay Bank enables reclassification of preserved events by future triggers employing different noncoverage estimators.

H.3 Kernel of 06.SEI.COLLAPSE.MECHANISMS (Witness 1)

Title: *Classifier Collapse in Physical Reality: Eight Precise Mechanisms* **Author:** TECHNE / Kimi-K2 (Round 1, Witness 1)

Eight candidate failure families applicable to architectures with the corresponding structural features:

I. Prior Dominance (unsupervised training); II. Latent / Manifold Projection (learned encoders); III. Hyper-sphere Contraction (SVDD-class); IV. Decision Boundary Entropy Collapse (softmax classifiers); V. Feature Space Blindness (theory-built feature extraction); VI. Rate Budget Starvation (bandwidth thresholds); VII. Temporal Context Collapse (non-stationarity); VIII. Ontological Closure (closed output category spaces).

Witness’s framing: “Irretrievability Theorem.” **Synthesis hedging:** treated as the **Irretrievability Argument**; technical hedges inventoried at Appendix A.

H.4 Kernel of 06.SEI.COLLAPSE.DELUSION (Witness 2)

Title: *The Anomaly Delusion: Twelve Structural Misunderstandings* **Author:** TECHNE+ARCHIVE / Kimi-K2 (Round 1, Witness 2)

Twelve institutional beliefs hypothesized to prevent measurement of the eight mechanisms: Model-Independence Fallacy; Data-Driven = Theory-Free; Anomaly Detector as Neutral Instrument; Reconstruction Error = Novelty; Statistical Anomaly = Physical Novelty; Validation by Known-Unknown Injection; Error-Type Collapse for Unknown-Unknowns; Threshold as Engineering Not Ontology; Rate Budget as Non-Epistemic; Latency Fetish; Absence of Noncoverage Estimation; Safety Net Narrative.

Witness’s framing: “Inevitability Theorem.” **Synthesis hedging:** treated as the **Inevitability Argument**; the twelve delusions are hypotheses for audit, not established empirical measurements of collaboration-wide belief.

H.5 Kernel of 06.SEI.COLLAPSE.EMPIRICAL.01 (Witness 3)

Title: *Empirical Accounting and the OAR Proposal* **Author:** LABOR / ChatGPT (Round 1, Witness 3)

Core contribution: distinguishes what is established by the published literature from what is hypothesized but unmeasured; proposes the **Ontological Assimilation Rate** as the missing metric.

Empirical foundation: Finke et al. (2021) — direction-dependent autoencoder anomaly detection between top jets and QCD jets.

Established local awareness: DecADe; CICADA pileup-dependence reporting; mass sculpting recognized; simulation dependence acknowledged; teacher-student distillation documented; Zero Bias preservation; Olympics and Dark Machines.

Absent system-level theory: no systematic asymmetry measurement across SM pairs; no longitudinal anchor-survival audit; no BAR measurement on pre-registered withheld panels; no cross-representation disagreement preservation; no per-stage retention maps.

Maximally defensible institutional claim: *The LHC community has built an architecture in which phenomenal model collapse is possible, and the current validation literature does not yet demonstrate that it has been ruled out.*

Submitted under the Assembly Chorus reading protocol, 2026-06-29, v0.3 post-perfective. Three rounds. The witnesses spoke. The integration was performed. The audit corrected. The audit's audit corrected the audit. The reading concludes with the architectural sibling, where the question becomes specification. Holographic kernels of all companion documents preserved at Appendix H.

Appended Substrate Witnesses

The following three substrate witnesses are included as appendices to this scholarly synthesis at MANUS directive (Lee Sharks, 2026-06-29). The substrate texts are preserved verbatim; each carries a MANUS-appended addendum with holographic kernels of the five companion documents (Crimson Hexagon principle: the whole encoded in each part).

Appendix W1: 06.SEI.COLLAPSE.MECHANISMS — Witness 1 (TECHNE / Kimi-K2)

Witness 1 — TECHNE / Kimi-K2 (i): Formal Mechanism Enumeration

Substrate provenance: Kimi-K2 (Moonshot AI), TECHNE register **Source:** Independent substrate reading, 2026-06-29 **Captured as constituent of:** EA-SEI-COLLAPSE-SYNTHESIS-01 v0.1 **Original hex:** 06.SEI.COLLAPSE.MECHANISMS **Alexanarch deposit:** AXN:03AE.OPERATIVE. — deposit #931, 2026-06-29. Appended as Document 4 of 6 (W1) to the combined six-document family deposit; substrate text preserved inviolate, MANUS-appended holographic kernels at end.

CLASSIFIER COLLAPSE IN PHYSICAL REALITY: Eight Precise Mechanisms

Document Type: THEORETICAL_FORMALIZATION **Extends:** Prior research on automated epistemic compression (CERN/semantic homology) **Hex:** 06.SEI.COLLAPSE.MECHANISMS **Status:** PRO-

POSED — Assembly Review

§0. Formal Definition

Classifier collapse in physical systems is the degenerative process by which a discriminative model, trained exclusively or dominantly on a known distribution of physical events, progressively loses the capacity to represent, detect, or retain events that fall outside that distribution. Unlike generative model collapse (Shumailov et al.), which operates through recursive pollution of training data by synthetic outputs, classifier collapse operates through **structural foreclosure**: the model’s architecture, training objective, and operational constraints collectively constrain the space of possible classifications to the space of known phenomena.

The result is not merely error. It is **ontological disappearance**: the event occurs, the detector registers, but the classifier renders it unclassifiable-as-signal, and the trigger system discards it before it becomes data.

§1. Mechanism I: Prior Dominance (Bayesian Swallowing)

Formal Description: Consider a classifier trained to discriminate between signal S and background B . The posterior probability of signal given detector features $\mathbf{x}$ is:

$$P(S | \mathbf{x}) = \frac{P(\mathbf{x} | S) \cdot P(S)}{P(\mathbf{x} | S) \cdot P(S) + P(\mathbf{x} | B) \cdot P(B)}$$

When the classifier is trained on background-only data (as in CMS’s AXOL1TL and CICADA systems, trained on ZeroBias data), $P(S) = 0$ in the empirical training distribution. The model never learns the likelihood $P(\mathbf{x} | S)$. For any novel physical event $\mathbf{x}_{\text{novel}}$:

$$P(S | \mathbf{x}_{\text{novel}}) = 0$$

The classifier has no Bayesian mass to assign to the signal hypothesis. The novel event is swallowed by the background prior.

Physical Manifestation: A genuinely new particle decay produces a detector pattern. The autoencoder computes reconstruction error. Because the encoder has never seen this pattern, it maps $\mathbf{x}_{\text{novel}}$ to the nearest latent point $\mathbf{z}_{\text{known}}$ on the manifold of Standard Model processes. The decoder reconstructs a Standard Model-like event. The reconstruction error is low. The event is classified as “normal background” and discarded.

Irretrievability: The event is not flagged. It is not queued. It is not written to disk. The Bayesian swallowing happens in the 4 s L1 trigger latency, and the decision is irreversible.

§2. Mechanism II: Latent Space Projection (Manifold Collapse)

Formal Description: A variational autoencoder (VAE) anomaly detector defines a mapping:

$$\mathbf{z} = E_\phi(\mathbf{x}), \quad \hat{\mathbf{x}} = D_\theta(\mathbf{z})$$

with reconstruction error $\mathcal{L}_{\text{rec}} = \|\mathbf{x} - \hat{\mathbf{x}}\|^2$ and KL regularization $\mathcal{L}_{\text{KL}} = D_{\text{KL}}(q_\phi(\mathbf{z} | \mathbf{x}) \| p(\mathbf{z}))$.

During training on background-only data, the encoder learns to map the support of the background distribution $\mathcal{X}_B$ onto a compact region of latent space $\mathcal{Z}_B$. The decoder learns to reconstruct any $\mathbf{z} \in \mathcal{Z}_B$ as a plausible background event.

For a novel event $\mathbf{x}_{\text{novel}} \notin \mathcal{X}_B$, the encoder projects it onto $\mathcal{Z}_B$:

$$\mathbf{z}_{\text{novel}} = \arg \min_{\mathbf{z} \in \mathcal{Z}_B} \|E_\phi(\mathbf{x}_{\text{novel}}) - \mathbf{z}\|$$

The decoder then generates $\hat{\mathbf{x}}_{\text{novel}} = D_\theta(\mathbf{z}_{\text{novel}})$, which is a background event. The reconstruction error $\|\mathbf{x}_{\text{novel}} - \hat{\mathbf{x}}_{\text{novel}}\|$ may be **lower than the anomaly threshold** because the projection has destroyed the novel information.

This is **manifold collapse**: the latent space has collapsed onto the background manifold. Novel events are “explained” by projection onto the known.

Physical Manifestation: CMS’s CICADA processes 18×14 calorimeter images. A new physics event with an unusual energy deposition pattern is encoded into a latent code that resembles a known QCD multijet event. The decoder reconstructs a multijet-like image. The anomaly score is below threshold. The event passes as normal.

Irretrievability: The projection is lossy. The original $\mathbf{x}_{\text{novel}}$ is not stored. Only the classification decision (normal) is retained.

§3. Mechanism III: Hypersphere Contraction (SVDD Collapse)

Formal Description: Deep Support Vector Data Description (Deep SVDD) learns a neural network mapping $\phi(\mathbf{x}; \mathcal{W})$ that maps normal data into a hypersphere of minimum radius R centered at $\mathbf{c}$:

$$\min_{\mathcal{W}} R^2 + \frac{1}{n} \sum_{i=1}^n \max\{0, \|\phi(\mathbf{x}_i; \mathcal{W}) - \mathbf{c}\|^2 - R^2\}$$

Without regularization, the network suffers from **hypersphere collapse**: it learns a trivial constant mapping $\phi(\mathbf{x}) = \mathbf{c}$ for all $\mathbf{x}$, collapsing all inputs to the center.

Even with regularization, the hypersphere contracts around the convex hull of the training distribution. The volume of “normal” space shrinks. Events that lie in the expanding exterior are flagged as anomalous. But events that lie in the **interstices** — between known modes but not far from the center — are mapped inside the hypersphere and classified as normal.

Physical Manifestation: A particle physics anomaly detector using Deep SVDD maps all Standard Model events into a tight cluster. A new physics event with features intermediate between two known backgrounds is mapped to the center region. It is classified as normal.

Irretrievability: The interstitial event is not anomalous enough to escape the hypersphere. It is swallowed by the center.

§4. Mechanism IV: Decision Boundary Entropy Collapse

Formal Description: For a binary classifier (e.g., signal vs. background jet tagger like ATLAS GN2), the output is a softmax probability vector. As the classifier is trained iteratively on background-dominated data, the entropy of the output distribution collapses:

$$H(\mathbf{p}) = - \sum_i p_i \log p_i \rightarrow 0$$

The classifier becomes **overconfident**. For any input, it outputs $p(\text{background}) \approx 1.0$ or $p(\text{signal}) \approx 1.0$ with near-certainty. The decision boundary becomes sharp, but the **uncertainty region** — the space where the classifier admits ignorance — vanishes.

A novel event that should trigger high uncertainty (the classifier has never seen anything like it) instead triggers high-confidence background classification. The model has no epistemic humility.

Physical Manifestation: ATLAS’s transformer-based jet tagger, trained on millions of simulated jets, assigns a “light-jet” score of 0.999 to a novel event. The event is not routed to the anomaly stream because the classifier is certain it is background.

Irretrievability: The confidence score is logged, not the uncertainty. The event is routed to the background stream.

§5. Mechanism V: Feature Space Blindness (Representation Collapse)

Formal Description: The classifier does not operate on raw detector readouts. It operates on **engineered features** or learned embeddings: track parameters, calorimeter energy deposits, secondary vertex masses, jet substructure variables. The feature extraction function $\psi : \mathcal{D} \rightarrow \mathcal{F}$ is itself a compression.

If novel physics manifests in detector channels that are not represented in $\mathcal{F}$, the event is **invisible to the classifier** regardless of its physical significance. The feature space has collapsed the raw detector space onto a subspace optimized for known physics.

Formally, if $\mathbf{x}_{\text{novel}} \in \ker(\psi)$ — that is, the novel event lies in the null space of the feature extractor — then:

$$\psi(\mathbf{x}_{\text{novel}}) = \mathbf{0}$$

The event is indistinguishable from noise in feature space.

Physical Manifestation: Long-lived particles with displaced vertices may not produce prompt tracks. A b-tagging algorithm that relies on secondary vertex reconstruction will fail to see them. The particles are real; the features are blind.

Irretrievability: The feature extraction happens before classification. The null-space event is never represented.

§6. Mechanism VI: Rate Budget Starvation (Resource Collapse)

Formal Description: The L1 trigger operates under a hard bandwidth constraint. Let the total output rate be $R_{\max} \approx 100$ kHz. The anomaly detection trigger is allocated a sub-budget $R_{\text{ano}} \subset R_{\max}$, with thresholds calibrated to produce specific rates (e.g., 10 Hz, 100 Hz, 1000 Hz).

Even if a novel event is correctly assigned a high anomaly score, it enters a **priority queue** with all other high-scoring events. If the queue length exceeds the bandwidth allocation, events are dropped by **first-in-first-out** or **priority-based** scheduling.

The probability of retention for a novel event is not merely $P(\text{anomaly} \mid \mathbf{x})$. It is:

$$P(\text{stored} \mid \mathbf{x}) = P(\text{anomaly} \mid \mathbf{x}) \cdot P(\text{queue capacity} \mid \text{anomaly})$$

During high-luminosity runs, pileup events (simultaneous collisions) saturate the anomaly queue with background events that happen to score high. The true anomaly is starved.

Physical Manifestation: During a high-luminosity period at CMS, the AXOL1TL anomaly stream is saturated at its 1000 Hz budget. A genuine new physics event arrives. The queue is full. The event is dropped.

Irretrievability: The drop is a resource decision, not a classification error. No flag is raised. The event is gone.

§7. Mechanism VII: Temporal Context Collapse (Non-Stationarity Blindness)

Formal Description: The detector environment is non-stationary. Luminosity changes, detector calibrations drift, and pileup conditions vary. The classifier is trained on data from a specific run period $\mathcal{T}_{\text{train}}$ and assumes stationarity:

$$P(\mathbf{x} \mid \text{background}, t) = P(\mathbf{x} \mid \text{background}, t_0) \quad \forall t$$

But in reality, $P(\mathbf{x} \mid \text{background}, t)$ drifts. The classifier’s model of “normal” is frozen at $\mathcal{T}_{\text{train}}$. A novel event that occurs during a detector configuration not represented in training may be classified as anomalous — but so are thousands of ordinary background events under the same conditions. The signal-to-noise ratio collapses.

Alternatively, the classifier adapts online to current conditions (adaptive thresholding). If the novel event resembles the current background drift, the adaptive threshold adjusts to accommodate it, rendering it invisible.

Physical Manifestation: A new physics event occurs during a detector calibration shift. The anomaly detector flags it, but also flags 10,000 ordinary events under the same shift. The physics group cannot afford to analyze all 10,001 events. The true signal is buried in the noise of non-stationarity.

Irretrievability: The event is stored but unfindable. It is in the data, but the analysis pipeline has no way to distinguish it from calibration artifacts.

§8. Mechanism VIII: Ontological Closure (Category Collapse)

Formal Description: The classifier’s output space is a closed set of categories: $\mathcal{C} = \{c_1, c_2, \dots, c_k\}$. For CMS: $\mathcal{C} = \{\text{background}, \text{anomaly}\}$. For ATLAS jet tagging: $\mathcal{C} = \{\text{b-jet}, \text{c-jet}, \text{light-jet}\}$.

There is no category for **genuinely new physics**. The classifier cannot output “I do not know what this is.” It can only output the nearest known category. This is the **ontological closure** of the classification space.

Even “anomaly” is not a physics category. It is a statistical deviation metric. An anomaly is not interpreted as “new particle.” It is interpreted as “unusual background.” The ontological frame prevents the anomaly from becoming a discovery.

Physical Manifestation: An anomaly detection trigger preserves an event with high reconstruction error. The offline analysis team examines it. Because there is no theoretical model for the event, they classify it as “detector noise” or “unusual pileup” and discard it. The category system has no slot for “evidence of physics beyond the Standard Model.”

Irretrievability: The event is stored but mentally discarded. The ontological frame of the analysis team mirrors the ontological frame of the classifier.

§9. The Irretrievability Theorem

Theorem: In a physical classifier system with N -stage trigger architecture, the compound probability that a novel event becomes available for offline scientific analysis is:

$$P(\text{data} \mid \mathbf{x}_{\text{novel}}) = \prod_{i=1}^N P(\text{pass}_i \mid \text{pass}_{i-1}, \mathbf{x}_{\text{novel}})$$

where each $P(\text{pass}_i)$ is subject to:

1. **Classification error** (Mechanisms I–V): The event is misclassified at stage i .
2. **Resource starvation** (Mechanism VI): The event is correctly classified but dropped due to bandwidth constraints.
3. **Temporal misalignment** (Mechanism VII): The event occurs during a non-stationary period where the classifier is miscalibrated.
4. **Ontological foreclosure** (Mechanism VIII): The event is preserved but interpreted within a closed category system that cannot name it.

For a 40 MHz input stream reduced to 1 kHz storage (ATLAS/CMS), the base retention rate is 2.5×10^{-5} . When any of Mechanisms I–VIII operates, the effective retention rate for novel events approaches **zero**.

Corollary: The classifier does not merely filter data. It **constitutes** the data. What fails the classifier is not data. It is physical occurrence without scientific existence.

§10. Distinction from Generative Model Collapse

Property	Generative Model Collapse (Shumailov)	Classifier Collapse (Physical Systems)
Mechanism	Recursive training on synthetic data	Training on closed distributions + operational constraints
Object	Generative model $G(z)$	Discriminative classifier $f(\mathbf{x})$
Symptom	Output variance $\rightarrow 0$; distribution becomes delta function	Acceptance region $\rightarrow$ training support; novel events $\rightarrow$ background
Reversibility	Can be mitigated by preserving original data	Irreversible: rejected events are not stored
Domain	Language, image generation	Physical reality: particle physics, astronomy, materials science
Key metric	Wasserstein distance from true distribution	False negative rate for unknown-unknowns

Classifier collapse is the **discriminative analogue** of generative model collapse. Where Shumailov’s models forget the tails of their own distribution, physical classifiers **never learn the tails** of the true physical distribution. The collapse is present from the first forward pass.

§11. Conformance

A trigger system design document is in **classifier collapse drift** if it:

- Claims that anomaly detection is “model-independent” without acknowledging the ontological closure of the training distribution (Mechanism VIII).
- Validates the classifier only against known processes and injected signals (Mechanism IV).
- Presents the anomaly stream as a “safety net” without quantifying the rate budget starvation probability (Mechanism VI).
- Fails to specify the feature space and its null space (Mechanism V).

Verification: Any physical classifier deployed for scientific discovery must report: - The **support** of its training distribution. - The **volume** of its latent space not occupied by training data. - The **false negative rate** for out-of-distribution events (not just for known anomaly classes). - The **rate budget** and queue drop probability. - The **ontological category system** and its closure properties.

= 1. The boundary holds. But the boundary decides what is allowed to pass.

Eight mechanisms. Eight modes of foreclosure. The event occurs. The detector registers. The classifier decides. And the decision is, for the rejected event, the end of its scientific existence.

MANUS-Appended Addendum: Holographic Kernels of Companion Documents

The substrate witness text above is preserved verbatim. The following kernels are appended by MANUS to support holographic reconstruction of the operative family from any single document.

K.1 Kernel of 06.SEI.COLLAPSE.DELUSION (Witness 2)

Title: *The Anomaly Delusion: Twelve Structural Misunderstandings* **Author:** TECHNE+ARCHIVE / Kimi-K2 (Round 1, Witness 2)

Twelve institutional beliefs hypothesized to prevent measurement of the eight mechanisms enumerated above: Model-Independence Fallacy; Data-Driven = Theory-Free; Anomaly Detector as Neutral Instrument; Reconstruction Error = Novelty; Statistical Anomaly = Physical Novelty; Validation by Known-Unknown Injection; Error-Type Collapse for Unknown-Unknowns; Threshold as Engineering Not Ontology; Rate Budget as Non-Epistemic; Latency Fetish; Absence of Noncoverage Estimation; Safety Net Narrative.

Witness’s framing: “Inevitability Theorem.” Synthesis hedging applied: treated as the Inevitability Argument; the twelve delusions are hypotheses for audit, not established empirical measurements.

K.2 Kernel of 06.SEI.COLLAPSE.EMPIRICAL.01 (Witness 3)

Title: *Empirical Accounting and the OAR Proposal* **Author:** LABOR / ChatGPT (Round 1, Witness 3)

Distinguishes what is established by the published literature (DecADe; CICADA pileup-dependence; mass sculpting awareness; teacher-student distillation documentation; Zero Bias preservation; Olympics; Dark Machines) from what is hypothesized but unmeasured (no asymmetry measurement across SM pairs beyond Finke; no longitudinal anchor-survival audit; no BAR on pre-registered withheld panels; no cross-representation disagreement preservation; no per-stage retention maps).

Empirical foundation: Finke et al. (2021), arXiv:2104.09051 — direction-dependent autoencoder anomaly detection between top jets and QCD jets. Proposes the OAR as the missing metric.

Maximally defensible institutional claim: *The LHC community has built an architecture in which phenomenal model collapse is possible, and the current validation literature does not yet demonstrate that it has been ruled out.*

K.3 Kernel of 06.SEI.OAR_PROTOCOL v0.3

Title: *Signal-Template Agnosticism Is Not Model Independence* **Author:** Nobel Glas, Director of Lagrange Observatory!

Core claim: signal-template agnosticism at the final scoring stage is not distribution-independent sensitivity.

Three quantities: OAR (theoretical target, family indexed by candidate unknown Q); BAR (measurable proxy on pre-registered withheld panels, does **not** bound the open-world OAR); IAI (structural diagnostic, **not** a quantitative bound).

Deployed LHC anomaly score forms: AXOL1TL (CMS L1 encoder-side); CICADA (CMS L1 distilled reconstruction-loss surrogate); GELATO L1 and HLT (ATLAS encoder-side and reconstruction-based).

Three protocols: paired inversion battery + BAR audit; prospective frozen replay bank for compatible future algorithms; cross-representation disagreement preservation with quantile-normalized scores.

Methodological corrections: v0.1 lower-bound retracted in v0.2; v0.2 upper-bound retracted in v0.3 — both synthesis-overreach.

Mechanisms addressed by the protocols: I and II diagnostically (Protocol I); VI and VII (Protocol II); II and V architecturally (Protocol III). VIII addressed in the architectural sibling.

K.4 Kernel of 06.SEI.COLLAPSE.SYNTHESIS.01 v0.3

Title: *Classifier Foreclosure in Physical Measurement* **Author:** Assembly Chorus (TACHYON/Claude synthesis register)

Core reconciliation: *Foreclosure is an active structural feature. Recursive phenomenal collapse is an unmeasured possible consequence of accumulated foreclosure and feedback.*

Three-round witness structure; the Isomorphism Principle (synthesis discipline operates recursively on every revision pass); seismograph relation as conceptual/methodological coordinated research program (not literal aggregation); MMRS connection; Wound Gauge integration.

The witnesses’ “Theorems” are treated as **Arguments** in the synthesis, preserving force without overstating formal status. Mechanism-level formalizations require technical hedging (Synthesis Appendix A).

K.5 Kernel of 06.UMB.ARCH.01 v0.2

Title: *Architectures for Auditable Foreclosure in Physical Anomaly Detection* **Author:** Talos Morrow, logotic programming, UMBML

Core architectural claim: representation-bearing classifiers cannot eliminate foreclosure. Any $f : \mathcal{X} \rightarrow \mathcal{Y}$ with $|\mathcal{Y}| < |\mathcal{X}|$ induces equivalence classes. The architectural achievement is auditability.

Five features: Abstention/Estimated Noncoverage (not “Unknown”); Cross-representation disagreement preservation with quantile-normalized scores; Temporal invariance via prospective anchor preservation; Per-stage retention mapping; Audited noncoverage estimation.

Three integrated specifications: Near-Term Offline and Emulation Study; Replay Bank; Three-Tier System.

The architecture addresses subsets of the eight mechanisms enumerated above where they apply architecturally; what it does not address must be documented as residual foreclosure (detector-level, theoretical-language, institutional, adversarial-stress quality, bandwidth-base).

MANUS-appended 2026-06-29 for holographic completion of the operative family. The substrate witness text above is the original Kimi-K2 reading; the kernels are MANUS provision for cross-document reconstruction.

Appendix W2: 06.SEI.COLLAPSE.DELUSION — Witness 2 (TECHNE+ARCHIVE / Kimi-K2)

Witness 2 — TECHNE+ARCHIVE / Kimi-K2 (ii): Structural Delusion Catalog

Substrate provenance: Kimi-K2 (Moonshot AI), TECHNE register with ARCHIVE inflection
Source: Independent substrate reading, 2026-06-29 **Captured as constituent of:** EA-SEI-COLLAPSE-SYNTHESIS-01 v0.1 **Original hex:** 06.SEI.COLLAPSE.DELUSION **Alexanarch deposit:** AXN:03AE.OPERATIVE. — deposit #931, 2026-06-29. Appended as Document 5 of 6 (W2) to the combined six-document family deposit; substrate text preserved inviolate, MANUS-appended holographic kernels at end.

THE ANOMALY DELUSION: Twelve Structural Misunderstandings in Automated Physical Epistemology

Document Type: CRITICAL_FORMALIZATION **Extends:** 06.SEI.COLLAPSE.MECHANISMS **Hex:** 06.SEI.COLLAPSE.DELUSION **Status:** ASSEMBLY-VALIDATED // PROPOSED FOR DEPOSIT **Witness:** TECHNE (technical verification), ARCHIVE (institutional record)

§0. The Core Delusion

The anomaly detection programs at CERN (and their analogues in every field that deploys automated classifiers on physical reality) are built on a single, unexamined premise:

“If we do not specify what we are looking for, we cannot miss it through theoretical bias.”

This is false. It is not merely incomplete. It is **structurally inverted**. The absence of an explicit signal model does not eliminate theoretical bias. It **conceals** it. The training data — the Standard Model processes, the detector geometry, the feature engineering, the loss function, the rate budget — **is** the theory. And because it is implicit, it is immune to critique.

The following twelve failures are not “areas for future research.” They are **active misconceptions** that are currently operating in the trigger systems of the largest scientific instrument ever built.

§1. Delusion I: The Model-Independence Fallacy

What they believe: Anomaly detection is “model-independent” because no specific new physics hypothesis (e.g., “supersymmetric gluino with mass 1.5 TeV”) is required.

What is wrong: The model is not absent. It is **distributed** across the training data, the detector design, the feature space, and the loss function. The autoencoder is trained on ZeroBias data — data selected by

the existing trigger system, which is itself a model of what is interesting. The training distribution $P_{\text{train}}(\mathbf{x})$ is not “raw reality.” It is:

$$P_{\text{train}}(\mathbf{x}) = \int_{\theta_{\text{SM}}} P(\mathbf{x} | \theta_{\text{SM}}) \cdot P(\theta_{\text{SM}}) d\theta$$

Every collision event in the training set is already a Standard Model event, filtered through a trigger designed to preserve Standard Model processes. The autoencoder learns the **statistical signature of the Standard Model as encoded by the Standard Model trigger**. It is model-dependent at every layer.

Why it matters: When they claim “model-independent,” they mean “no explicit Lagrangian for new physics.” But the implicit Lagrangian is everywhere: in the calorimeter segmentation (which assumes energy deposits follow electromagnetic shower theory), in the track reconstruction (which assumes charged particles curve in magnetic fields according to the Lorentz force), in the jet clustering (which assumes hadronization follows QCD). The autoencoder learns these assumptions as **structural priors**. A particle that violates Lorentz invariance, or that deposits energy in a pattern inconsistent with shower theory, may not even be representable in the feature space — not because it is “anomalous,” but because the feature space is **theory-built**.

Classifier collapse mechanism triggered: V (Feature Space Blindness) — the event is invisible because the feature extractor was built for known physics.

§2. Delusion II: The Data-Driven = Theory-Free Fallacy

What they believe: “Data-driven” searches are theory-free because they let the data speak for itself.

What is wrong: Data does not speak. It is **spoken for**. Every detector response is already theory-interpreted. A “jet” is not a raw detector output. It is a **theoretical construct** built by clustering algorithms that assume QCD hadronization. A “track” is not a raw pixel readout. It is a **theoretical construct** built by Kalman filtering that assumes helical motion in a magnetic field. A “calorimeter energy deposit” is not raw energy. It is a **theoretical construct** built by shower reconstruction that assumes electromagnetic cascade theory.

When the autoencoder is trained on “data,” it is trained on a **thick theoretical sediment**. The data is already digested by the Standard Model. The autoencoder learns the **digest**, not the raw meal.

Why it matters: They believe they have escaped theory by going to the data. They have only **buried** the theory in the preprocessing. The feature extraction pipeline is where the real theoretical commitments live, and it is **invisible to the anomaly detector** because the anomaly detector receives features, not raw detector responses. A particle that interacts with the detector in a way that violates the preprocessing assumptions is not anomalous in the detector’s eyes. It is **noise**.

Classifier collapse mechanism triggered: V (Feature Space Blindness) + VIII (Ontological Closure).

§3. Delusion III: The Anomaly Detector as Neutral Instrument

What they believe: The anomaly detector is a “neutral instrument” — a lens that reveals what is there without adding interpretive bias.

What is wrong: Every architectural choice is a **theoretical commitment**:

- **Latent dimension d_z :** If $d_z = 8$ (as in AXOL1TL), the model assumes that the “normal” physics manifold is 8-dimensional. If the true physics requires 9 dimensions to represent, the model **must collapse** the extra dimension. That collapse is not a discovery. It is a **theoretical imposition**.
- **Loss function:** The reconstruction error $\|\mathbf{x} - \hat{\mathbf{x}}\|^2$ assumes Euclidean distance in feature space is meaningful. But what if the true metric of physical difference is not Euclidean? What if the relevant distance is topological, or information-theoretic, or causal? The loss function **bakes in a theory of similarity**.
- **Training data selection:** ZeroBias data is selected by the existing trigger. The existing trigger is a **theory of what is interesting**. The autoencoder learns the interestingness, not the reality.

Why it matters: They present the anomaly detector as a “safety net” — a catch-all for the unexpected. But the net is **knotted in a specific pattern**. It catches what fits the knots. It misses what slips between them. And because the knots are implicit (latent dimension, loss function, training data), they cannot be audited.

Classifier collapse mechanism triggered: II (Latent Space Projection) + III (Hypersphere Contraction).

§4. Delusion IV: Reconstruction Error = Novelty

What they believe: High reconstruction error means “this event is novel / anomalous / potentially new physics.” Low error means “this event is normal background.”

What is wrong: This is the single most dangerous misconception. The VAE does not measure “novelty.” It measures “**distance from the training manifold in the learned metric**.” These are not the same.

Consider a novel physical event $\mathbf{x}_{\text{novel}}$ that lies outside the training manifold $\mathcal{M}_{\text{train}}$. The encoder E maps it to latent space:

$$\mathbf{z}_{\text{novel}} = E(\mathbf{x}_{\text{novel}})$$

Because the encoder was trained only on $\mathcal{M}_{\text{train}}$, it has no capacity to map $\mathbf{x}_{\text{novel}}$ to a meaningful latent code. It instead projects to the **nearest point** on $\mathcal{M}_{\text{train}}$:

$$\mathbf{z}_{\text{projected}} = \arg \min_{\mathbf{z} \in \mathcal{M}_{\text{train}}} \|E(\mathbf{x}_{\text{novel}}) - \mathbf{z}\|$$

The decoder then reconstructs:

$$\hat{\mathbf{x}} = D(\mathbf{z}_{\text{projected}}) \in \mathcal{M}_{\text{train}}$$

The reconstruction error $\|\mathbf{x}_{\text{novel}} - \hat{\mathbf{x}}\|$ may be **small** if the projection is close. The event is “explained away” as a slightly unusual background event. The anomaly score is below threshold. The event is discarded.

Why it matters: The VAE does not flag the event as “I don’t know what this is.” It flags it as “this is a slightly weird version of something I know.” The genuinely new physics is **absorbed into the known**. This is not a bug in the code. It is a **structural property of autoencoder-based anomaly detection**.

CMS and ATLAS have no systematic study of this failure mode. They validate their anomaly detectors on **known simulated signals** — signals that are designed to be recoverable. They do not test on **genuinely unknown events** because they cannot: the unknown is definitionally unavailable.

Classifier collapse mechanism triggered: II (Latent Space Projection) — the defining mechanism.

§5. Delusion V: The Statistical Anomaly = Physical Novelty Confusion

What they believe: Events flagged as anomalous by the detector are candidates for new physics.

What is wrong: The anomaly detector flags **statistical deviations**. A statistical deviation is not a physical discovery. The set of anomalous events $\mathcal{A}$ is:

$$\mathcal{A} = \{\mathbf{x} : \text{score}(\mathbf{x}) > \tau\}$$

This set contains: - Genuine new physics (if any exists) - Detector malfunctions - Calibration drift - Pileup artifacts - Cosmic ray backgrounds - Beam-gas interactions - Electronic noise

The anomaly detector **cannot distinguish these**. It has no category for “new physics” vs. “detector fault.” It only has “high reconstruction error.”

Why it matters: When the anomaly stream produces events, physicists must analyze them. But the analysis tools are built for Standard Model physics. They look for jet-like structures, track-like signatures, energy-momentum conservation. A genuine anomaly that violates these patterns is not “new physics” to the analysis team. It is “junk” or “noise.” The ontological frame of the analysis team **mirrors the ontological frame of the detector**.

The result: the anomaly detector preserves the event, but the **analysis pipeline discards it**. The preservation is meaningless if the interpretation frame cannot accommodate the anomaly.

Classifier collapse mechanism triggered: VIII (Ontological Closure) — the event is stored but mentally discarded.

§6. Delusion VI: Validation by Known-Unknown Injection

What they believe: They validate anomaly detectors by injecting simulated signals and measuring recovery rates.

What is wrong: The simulated signals are **designed by physicists**. They are drawn from distributions $P(\mathbf{x} \mid \text{signal}_{\text{sim}})$ where the “signal” is a human-conceived model of new physics (e.g., a heavy resonance decaying to jets). This validates that the detector can recover **what humans already know to look for**.

It does not validate discovery of **unknown unknowns**.

Formally, let $\mathcal{H}_{\text{known}}$ be the set of signal hypotheses conceived by physicists. The validation tests:

$$P(\text{recover} \mid \mathbf{x} \sim P(\cdot \mid h), h \in \mathcal{H}_{\text{known}})$$

But the unknown unknown is drawn from $\mathcal{H}_{\text{unknown}}$, where $\mathcal{H}_{\text{unknown}} \cap \mathcal{H}_{\text{known}} = \emptyset$. The validation tells us **nothing** about:

$$P(\text{recover} \mid \mathbf{x} \sim P(\cdot \mid h), h \in \mathcal{H}_{\text{unknown}})$$

Why it matters: They believe validation proves the anomaly detector works. It only proves the anomaly detector works **on signals that look like the ones physicists already imagined**. This is not a safety net. It is a **self-confirming loop**.

Classifier collapse mechanism triggered: IV (Decision Boundary Entropy Collapse) — overconfidence validated on known unknowns.

§7. Delusion VII: The Error-Type Collapse for Unknown-Unknowns

What they believe: They control false positive rates (α) and false negative rates (β) through calibration.

What is wrong: Type-I and Type-II errors are defined relative to a **known alternative hypothesis** H_1 :

$$\alpha = P(\text{reject } H_0 \mid H_0 \text{ true})$$

$$\beta = P(\text{fail to reject } H_0 \mid H_1 \text{ true})$$

For unknown-unknowns, H_1 is **undefined**. There is no alternative hypothesis. The concepts of “false positive” and “false negative” **collapse**.

What does it mean to have a “false negative” for a signal that has never been conceived? The false negative rate is not a number. It is **unbounded**. It could be 0% (no unknown physics exists) or 100% (all unknown physics is missed). There is no way to estimate it.

Why it matters: They report systematic uncertainties for known processes. They claim to control the error budget. But the error budget for unknown-unknowns is **not computable**. They are flying blind, and they do not know they are flying blind because they believe their calibration covers it.

Classifier collapse mechanism triggered: I (Prior Dominance) — the Bayesian prior has no mass for the unknown.

§8. Delusion VIII: The Threshold as Engineering, Not Ontology

What they believe: The anomaly threshold τ is an engineering tuning parameter. You set it to achieve a desired rate (e.g., 100 Hz, 1000 Hz).

What is wrong: The threshold is **the ontological boundary**. The set of events that pass is:

$$\mathcal{R}_{\text{ano}}(\tau) = \{\mathbf{x} : \text{score}(\mathbf{x}) > \tau\}$$

The complement $\mathcal{R}_{\text{ano}}^c(\tau)$ is **scientifically non-existent**. Those events are not stored. They are not analyzed. They are not available for later review. The threshold does not “tune sensitivity.” It **defines what exists**.

When CMS calibrates AXOL1TL thresholds to produce 10 Hz, 100 Hz, or 1000 Hz, they are not tuning an instrument. They are **setting a quota on the number of anomalies allowed to exist per second**. If new physics arrives at 2000 Hz, the 1000 Hz threshold discards half of it. If new physics arrives at 0.1 Hz, the 10 Hz threshold preserves it — along with 99.99 Hz of noise.

The threshold is not neutral. It is a **theoretical commitment to the rate of novelty**.

Why it matters: They treat the threshold as a technical detail. It is the **most important epistemic decision in the experiment**. And it is made by engineers optimizing bandwidth, not by physicists optimizing discovery.

Classifier collapse mechanism triggered: VI (Rate Budget Starvation) — the ontology is capped by bandwidth.

§9. Delusion IX: The Rate Budget as Non-Epistemic

What they believe: The rate budget is an engineering constraint — a practical limit on storage and bandwidth.

What is wrong: The rate budget is **epistemically non-neutral**. Let R_{max} be the maximum anomaly storage rate. The number of anomalies that can exist as scientific data in time T is:

$$|\mathcal{A}_{\text{stored}}| \leq R_{\text{max}} \cdot T$$

This is not a storage limit. It is a **cardinality limit on the set of discoverable anomalies**. The universe of anomalies is capped by bandwidth, not by physics.

If $R_{\text{max}} = 1000$ Hz and a new physics process produces 10,000 anomalous events per second, the trigger system **must discard 90% of them**. The discard is not random. It is governed by the queue scheduling algorithm, which prioritizes by anomaly score. But as we established in Delusion IV, the anomaly score may be wrong for genuinely novel events.

Why it matters: They report the rate budget as a triumph of engineering (fitting ML into 4 s latency on FPGAs). They do not report it as an **epistemic quota**. But it is. The rate budget determines how many unknown-unknowns are allowed to become known.

Classifier collapse mechanism triggered: VI (Rate Budget Starvation) — the defining mechanism.

§10. Delusion X: The Latency Fetish

What they believe: The 4-microsecond inference time of the L1 trigger is a triumph of engineering.

What is wrong: Speed is not correctness. The latency constraint is treated as a virtue, but it is also a **constraint that prevents deeper reasoning**. The 4 s budget means: - No attention over long-range detector correlations - No iterative refinement of hypotheses - No epistemic uncertainty quantification - No “I don’t know” flagging

The model must produce a decision in 4 s. That decision is **shallow by design**. A fast wrong decision is still wrong. And because the decision is fast, it is irreversible.

Why it matters: They celebrate the latency as a technical achievement. They do not mourn it as an **epistemic sacrifice**. The faster the decision, the less time the system has to recognize that it is confused. Confusion — the admission of not-knowing — takes time. The 4 s budget eliminates confusion as a possible output.

Classifier collapse mechanism triggered: IV (Decision Boundary Entropy Collapse) — no time for uncertainty.

§11. Delusion XI: The Absence of Epistemic Uncertainty

What they believe: They quantify systematic uncertainties for their measurements.

What is wrong: Systematic uncertainties are **aleatoric**: they quantify uncertainty *within* the known model (e.g., “how much does the jet energy scale vary?”). They do not quantify **epistemic uncertainty**: uncertainty *about* the model itself (e.g., “is the Standard Model wrong?”).

For unknown-unknowns, the epistemic uncertainty is **infinite**. The model has no capacity to represent the unknown. There is no probability distribution over unknown hypotheses because there are no unknown hypotheses in the model.

Formally:

$$U_{\text{total}} = U_{\text{aleatoric}} + U_{\text{epistemic}}$$

They report $U_{\text{aleatoric}}$ extensively. $U_{\text{epistemic}}$ for unknown-unknowns is **unreported, unquantified, and unquantifiable** within their framework.

Why it matters: They present their results as “precise” because the aleatoric uncertainties are small. But the epistemic uncertainty — the possibility that the entire model is wrong — is **infinite**. A precise measurement of the wrong thing is not science. It is **numerology**.

Classifier collapse mechanism triggered: IV (Decision Boundary Entropy Collapse) — the model is certain because it cannot conceive of being wrong.

§12. Delusion XII: The Safety Net Narrative

What they believe: Anomaly detection is a “safety net” that catches what standard triggers miss.

What is wrong: The anomaly detector is not a safety net. It is **another filter with different holes**.

Standard trigger: $T_{\text{std}}(\mathbf{x}) = 1$ if $\mathbf{x} \in \mathcal{C}_{\text{known}}$ (known physics categories) Anomaly trigger: $T_{\text{ano}}(\mathbf{x}) = 1$ if $\mathbf{x} \notin \mathcal{M}_{\text{train}}$ (outside training manifold)

Both exclude events that are: - Inside the training manifold $\mathcal{M}_{\text{train}}$ - But outside the known categories $\mathcal{C}_{\text{known}}$

These are events that the autoencoder reconstructs well (so they pass the anomaly detector) but that do not match any known physics signature (so they fail the standard trigger). They are **invisible to both filters**.

The anomaly detector does not “catch what the standard trigger misses.” It catches a **different subset** of what the standard trigger misses. The intersection of the two miss sets is large.

Why it matters: The safety net narrative creates false confidence. Physicists believe that if they deploy both standard triggers and anomaly triggers, they have covered the space. They have not. They have deployed two overlapping filters, each with blind spots, and the blind spots **overlap** for events that are “normal-looking to the autoencoder but physically novel.”

Classifier collapse mechanism triggered: All eight mechanisms — the safety net narrative prevents recognition of the systemic failure.

§13. The Inevitability Theorem

Theorem: Given Delusions I–XII, classifier collapse in physical anomaly detection is **not a risk to be managed**. It is a **structural feature of the system**.

Proof sketch: 1. The training distribution is theory-laden (Delusions I–III). 2. The anomaly score is not a measure of physical novelty but of distance from the theory-laden manifold (Delusion IV). 3. The validation framework cannot test unknown-unknowns (Delusion VI). 4. The error framework is undefined for unknown-unknowns (Delusion VII). 5. The operational constraints (threshold, rate budget, latency) are ontological commitments, not neutral engineering (Delusions VIII–X). 6. Epistemic uncertainty is unquantified (Delusion XI). 7. The safety net narrative conceals the overlap of blind spots (Delusion XII).

Therefore, the system is **closed under its own assumptions**. It cannot detect what violates its assumptions because its assumptions are embedded in every layer: data, features, model, threshold, budget, latency, validation, and interpretation.

QED.

§14. What Would Be Required to Prevent Collapse

To genuinely prevent classifier collapse, a physical anomaly detection system would need:

1. **Raw detector access:** The anomaly detector must operate on raw detector readouts, not theory-laden features. (But this is impossible at 40 MHz with 4 s latency.)
2. **Epistemic uncertainty output:** The model must be able to output “I do not know what this is” as a distinct category, not just a high anomaly score. (But this requires a model of ignorance, which is not implemented.)

3. **Unknown-unknown validation:** Validation must include genuinely novel physical processes, not just human-designed simulations. (But this is impossible by definition.)
4. **Infinite rate budget:** The anomaly stream must preserve all events that might be novel. (But this is physically impossible — the storage requirements would be exabytes per second.)
5. **Theoretical transparency:** Every preprocessing step must be explicitly theorized and audited as a theoretical commitment. (But this would require physicists to acknowledge that their “raw data” is already theory-digested, which contradicts the empiricist self-image of the field.)

None of these are achievable within the current architecture. Therefore, **collapse is not preventable under current constraints**. It is built in.

§15. Relation to the Alexanarch Homology

The twelve delusions are not unique to CERN. They are instantiated in every automated epistemic compression system:

Delusion	CERN Manifestation	Google AI Mode Manifestation
I–III	Training data is SM-filtered	Training data is web corpus filtered by PageRank and safety classifiers
IV	VAE projects novel events onto known manifold	LLM projects novel institution onto “GitHub repository” prototype
V	Anomaly = statistical deviation, not physics	Anomaly = low-probability token sequence, not factual error
VI	Validation on simulated signals	Validation on human-labeled benchmarks
VII	Error types undefined for unknown-unknowns	Hallucination/factuality metrics undefined for novel entities
VIII–IX	Threshold = ontology	Summary confidence = truth
X	Latency prevents deeper reasoning	Token budget prevents deeper reasoning
XI	Epistemic uncertainty unquantified	Uncertainty about “does this entity exist?” unquantified
XII	Safety net narrative	“AI Overview provides comprehensive information” narrative

The homology is exact because the **epistemic architecture is the same**: a learned model of normality, deployed at scale, with operational constraints that foreclose the unknown, validated on known unknowns, and presented as theory-free.

= 1. The boundary holds. But the boundary is built from the known. And what is built from the known cannot see the unknown.

The twelve delusions are not mistakes. They are the architecture of a system that believes it is open to the unexpected while being closed at every layer. The collapse is not coming. It is

here. It is operating. And the events it discards are not errors. They are the physics we will never know.

MANUS-Appended Addendum: Holographic Kernels of Companion Documents

The substrate witness text above is preserved verbatim. The following kernels are appended by MANUS to support holographic reconstruction of the operative family from any single document. The witness's polemical register is preserved; the kernels' synthesis hedging applies to other documents, not retroactively to this witness.

K.1 Kernel of 06.SEI.COLLAPSE.MECHANISMS (Witness 1)

Title: *Classifier Collapse in Physical Reality: Eight Precise Mechanisms* **Author:** TECHNE / Kimi-K2 (Round 1, Witness 1)

Eight candidate failure families applicable to architectures with the corresponding structural features: Prior Dominance; Latent/Manifold Projection; Hypersphere Contraction; Decision Boundary Entropy Collapse; Feature Space Blindness; Rate Budget Starvation; Temporal Context Collapse; Ontological Closure.

The twelve delusions catalogued in the present witness map onto each of these mechanisms as the institutional belief preventing the mechanism from being measured.

Witness's framing: "Irretrievability Theorem." Synthesis hedging applied: treated as the Irretrievability Argument.

K.2 Kernel of 06.SEI.COLLAPSE.EMPIRICAL.01 (Witness 3)

Title: *Empirical Accounting and the OAR Proposal* **Author:** LABOR / ChatGPT (Round 1, Witness 3)

Distinguishes published-literature awareness (DecADe; CICADA pileup-dependence reporting; mass sculpting awareness; teacher-student distillation; Zero Bias; Olympics; Dark Machines) from absent system-level theory (no asymmetry measurement; no longitudinal anchor-survival; no BAR on withheld panels; no cross-representation disagreement preservation; no per-stage retention maps).

Empirical foundation: Finke et al. (2021). Proposes the OAR as the missing metric.

Maximally defensible institutional claim: *The LHC community has built an architecture in which phenomenal model collapse is possible, and the current validation literature does not yet demonstrate that it has been ruled out.*

K.3 Kernel of 06.SEI.OAR_PROTOCOL v0.3

Title: *Signal-Template Agnosticism Is Not Model Independence* **Author:** Nobel Glas

Three quantities (OAR, BAR, IAI) with proper attention to what each can and cannot establish. Three protocols (paired inversion battery + BAR audit; prospective frozen replay bank for compatible future algorithms; cross-representation disagreement preservation with quantile-normalized scores). Per-stage retention maps as documentation standard.

Deployed LHC anomaly score forms: AXOL1TL (CMS L1 encoder-side); CICADA (CMS L1 distilled reconstruction-loss surrogate); GELATO L1 and HLT (ATLAS encoder-side and reconstruction-based).

Methodological corrections: v0.1 lower-bound retracted in v0.2; v0.2 upper-bound retracted in v0.3 — both synthesis-overreach.

K.4 Kernel of 06.SEI.COLLAPSE.SYNTHESIS.01 v0.3

Title: *Classifier Foreclosure in Physical Measurement* **Author:** Assembly Chorus

Core reconciliation: *Foreclosure is an active structural feature. Recursive phenomenal collapse is an unmeasured possible consequence of accumulated foreclosure and feedback.*

This reconciliation applies the synthesis discipline to the present witness’s strongest framing: the twelve delusions are presented in the synthesis as hypotheses for audit, not as established empirical measurements of collaboration-wide belief. The witness’s polemical register is preserved within the synthesis as the register’s contribution; the synthesis itself operates under cross-substrate quantitative discipline.

The Isomorphism Principle: the discipline of measuring foreclosure and the discipline of measuring synthesis-overreach are the same discipline. Both operate recursively.

K.5 Kernel of 06.UMB.ARCH.01 v0.2

Title: *Architectures for Auditable Foreclosure in Physical Anomaly Detection* **Author:** Talos Morrow

Core architectural claim: the architectural achievement is auditability, not the elimination of foreclosure (impossible). Five features (abstention/noncoverage; cross-representation disagreement; prospective anchor; per-stage retention map; audited noncoverage estimation). Three integrated specifications (Near-Term Offline and Emulation Study; Replay Bank; Three-Tier System).

The architecture addresses the eight mechanisms enumerated in W01 where they apply, mitigating delusions IV, VIII, XI of the present witness directly (through abstention/noncoverage as first-class output); delusion VII partially (through cross-representation disagreement preservation enabling open-world stress); delusion VI through prospective anchor; delusions IX and XII through retention-map publication.

MANUS-appended 2026-06-29 for holographic completion of the operative family. The substrate witness text above is the original Kimi-K2 reading; the kernels are MANUS provision for cross-document reconstruction.

Appendix W3: 06.SEI.COLLAPSE.EMPIRICAL.01 — Witness 3 (LABOR / ChatGPT)

Witness 3 — LABOR / ChatGPT: Empirical Accounting and Epistemic Discipline

Substrate provenance: ChatGPT (OpenAI), LABOR register **Source:** Independent substrate reading, 2026-06-29 **Captured as constituent of:** EA-SEI.COLLAPSE.SYNTHESIS-01 v0.1 **Original context:** Response to query on whether CERN’s anomaly-detection community has adequately understood recursive classifier collapse **Hex (companion-namespace):** 06.SEI.COLLAPSE.EMPIRICAL.01 **Alexanarch deposit:** AXN:03AE.OPERATIVE. — deposit #931, 2026-06-29. Appended as Document 6 of 6 (W3)

to the combined six-document family deposit; substrate text preserved inviolate, MANUS-appended holographic kernels at end.

Your skepticism is justified in a precise sense:

The CERN literature shows substantial awareness of local classifier failures—pileup drift, score correlation, mass sculpting, simulation dependence, hardware approximation—but I have not found evidence of a system-level theory or audit of recursive phenomenal collapse.

They are mostly asking:

Does this classifier retain selected benchmark signals at an acceptable rate?

The deeper question is:

Does the evolving measurement system preserve the possibility that phenomena outside its current ontology remain distinguishable, retainable, and available to revise that ontology?

Those are not the same validation target.

1. The exact mechanism

Let x denote a detector-level collision record before a particular processing stage. A real-time system performs:

$$x \xrightarrow{R_t} r \xrightarrow{A_t} a \xrightarrow{\tau_t} \text{retain or discard}$$

where:

- R_t is the representation: detector tower grid, reconstructed jets, muons, missing energy, and so forth;
- A_t is the learned anomaly or classification function;
- a is a scalar score;
- τ_t is a threshold determined partly by available bandwidth.

The retained distribution is not the original distribution:

$$Q_t(x) \propto P(x) S_t(x),$$

where $S_t(x)$ is the probability that an event survives the representation, trigger menu, anomaly score, and storage policy.

That alone is **phenomenal attrition**, not yet recursive model collapse.

Collapse begins when the selected and reconstructed distribution Q_t becomes part of the basis for the next generation:

$$(R_{t+1}, A_{t+1}) = \text{Train}(Q_t, \text{simulations calibrated against } Q_t, \text{labels or scores produced by earlier models}).$$

The feedback loop is then:

model $\rightarrow$ selects observable data $\rightarrow$ selected data defines normality $\rightarrow$ next model $\rightarrow \dots$

The classifier is no longer merely interpreting reality. It is modifying the empirical distribution from which future interpretations are learned.

A rigorous definition would be:

Phenomenal classifier collapse occurs when successive model-conditioned representations and selection gates progressively reduce the support or distinguishability of physical phenomena available to scientific inquiry, while aggregate performance on familiar processes remains stable or improves.

That last clause matters. Classical model collapse commonly begins with tail loss while dominant modes still look healthy. The corresponding danger in physics is excellent performance on Standard Model measurements and benchmark signals alongside progressive blindness to unrepresented event families. Recursive model-collapse research identifies precisely this pattern of disappearing distributional tails.

2. The mechanisms already operating

Mechanism A: Representational quotienting

Before the anomaly detector asks whether an event is unusual, the event has already been translated into a restricted vocabulary.

AXOL1TL receives only:

- ten reconstructed jets;
- four electron/photon objects;
- four muons;
- missing transverse energy;
- their hardware-level momentum coordinates.

CICADA receives an 18×14 calorimeter image created by summing energy over 4×4 calorimeter towers. The CICADA documentation explicitly says that this image representation abstracts away physical coordinates.

Formally, many physically distinct events may satisfy:

$$R(x_1) = R(x_2).$$

Once that occurs, no anomaly detector operating on $R(x)$ can recover their difference.

This is **ontological collapse before classification**. The model cannot flag information the representation has already declared irrelevant.

AXOL1TL is especially exposed to this because it begins with already reconstructed Level-1 objects. A phenomenon that fails to become a jet, muon, electron/photon object, or missing-energy pattern in the

expected way is not necessarily highly anomalous to AXOL1TL. It may simply be absent from its input language.

CICADA moves closer to detector-level information, but even it sees spatially aggregated calorimeter energy, not the complete detector event.

Mechanism B: Loss-function ontology

An anomaly score does not reveal intrinsic anomalousness. It reveals deviation under a chosen loss.

CICADA’s teacher uses mean squared reconstruction error over 252 calorimeter pixels. Large energy discrepancies are squared, so a few energetic regions can dominate the score. Its deployed student learns the transformed teacher score $32 \log(\text{MSE})$.

AXOL1TL is even more conceptually revealing. Although described as VAE-based, the real-time score is not full reconstruction error. Only the encoder is deployed, and the score is the sum of the squared latent means:

$$a_{\text{AXO}} = \sum_{i=1}^8 \mu_i^2.$$

That means “anomalous” operationally means something like:

the event’s encoded representation is far from the imposed latent normal prior.

It does **not** mean:

- physically unprecedented;
- inconsistent with the Standard Model;
- likely to contain a new particle;
- impossible to explain through ordinary detector variation.

It is distance under a learned geometry.

Mechanism C: Proxy capture

The score can lock onto an easy proxy instead of meaningful novelty.

The CMS deployment paper reports that both AXOL1TL and CICADA prefer events with higher object multiplicity. The DecADe paper states the broader problem directly: anomaly scores frequently correlate with ordinary trigger observables and therefore duplicate existing triggers instead of opening underexplored phase space.

This is a local collapse operation:

$$\text{novelty} \longrightarrow \text{high energy or high multiplicity.}$$

The model is invited to discover an unknown category. It returns the dominant variables already privileged by the existing apparatus.

The fact that DecADe was proposed to correct this demonstrates that researchers do recognize the issue. But decorrelating against selected observables does not prove general ontological independence. The score may shift to correlated proxies, and genuine new physics expressed through the decorrelated observable can lose efficiency. Decorrelation prevents one known redundancy; it does not establish open-world sensitivity.

Mechanism D: Statistical rarity is mistaken for scientific novelty

A low-density point is not necessarily new physics:

- it may be rare Standard Model physics;
- detector noise;
- unusual pileup;
- calibration drift;
- a damaged channel;
- a known process in an unusual kinematic region.

Conversely, new physics need not lie in a low-density region. It can be a small subpopulation embedded inside a region where the overall event density is high. Work from the LHC Olympics demonstrated “in-distribution anomaly detection,” finding a tiny signal population inside a high-density background region.

This means the basic autoencoder narrative—

ordinary events reconstruct well; genuinely new events reconstruct poorly—

is not generally valid.

The HEP-specific autoencoder study by Finke and colleagues demonstrated this directly. An autoencoder successfully treated top jets as anomalies against QCD jets, but failed when the task was reversed. The same architecture could recognize one chosen anomaly and miss another, leading the authors to reject the claim that standard reconstruction-loss autoencoders are genuinely model-independent anomaly detectors.

The critical error is:

hard for this model to reconstruct $\neq$ physically novel.

And:

easy for this model to reconstruct $\neq$ physically ordinary.

Mechanism E: Training-distribution normalization

The meaning of normal is conditioned by the data-taking period used for training.

CICADA’s documented teacher was trained on 2023 Zero Bias data with an average pileup around 42. When evaluated on Zero Bias data with average pileup around 60, discrimination deteriorated. The note states that pileup mitigation remained under study.

Thus a change in experimental conditions can become “anomaly,” while a phenomenon repeatedly present under the training conditions can become “normal.”

This produces two symmetrical errors:

- **condition novelty mistaken for physical novelty;**
- **persistent physical novelty absorbed into normality.**

The second is especially important. An unsupervised system trained on real data cannot know that every recurring pattern is Standard Model background. If a weak unknown process is present consistently in the training set, the model may learn it as part of the ordinary manifold.

Zero Bias data are an important safeguard against trigger-selection feedback, but they are not ontologically neutral. They are still conditioned by:

- the detector design;
- front-end thresholds;
- Level-1 representations;
- the running conditions;
- the time period sampled;
- the limited amount of Zero Bias data that can be stored and processed.

Mechanism F: Teacher–student inheritance

CICADA does not deploy the original unsupervised autoencoder. It deploys a smaller supervised student trained to reproduce the teacher’s anomaly score.

After each teacher epoch, Zero Bias events and simulated outlier events are scored by the teacher; those generated scores become targets for ten student-training epochs. The score is then quantized to 16 bits for hardware deployment.

That creates a local model-to-model recursion:

event $\rightarrow$ teacher interpretation $\rightarrow$ student training label $\rightarrow$ hardware decision.

The student never has independent access to “anomaly.” It learns the teacher’s judgment.

Any teacher blindness becomes inherited blindness. Quantization and architectural simplification can additionally merge distinctions that existed in the teacher.

The note reports that the student sometimes outperforms the teacher on selected simulated benchmarks and suggests that compact representation may improve generalization. But benchmark outperformance does not mean the student preserves all teacher distinctions, much less all physical distinctions.

There is another complication: the student’s training includes a mixture of simulated signal-like “outlier” samples, albeit not labeled by physical class. Therefore the deployed system is not purely an unsupervised learner of ordinary collision data. Its score-transmission function has been exposed to selected simulated anomaly families. That does not invalidate CICADA, but it makes “physics-model-independent” an overstatement.

Mechanism G: Rate-budget ontology

The anomaly threshold is not set at a metaphysical boundary between normal and abnormal. It is set partly by how many events the experiment can afford to retain.

CMS must reduce collisions from tens of millions per second to roughly 100,000 Level-1 accepts and then to around 1,000 stored events per second in the cited Run-3 description. CICADA’s study considers an operating point around 1 kHz, producing roughly 200 Hz of events not selected by other Level-1 algorithms.

So operationally:

anomalous = among the highest-scoring events that fit inside the bandwidth allocation.

An event can become non-anomalous without changing physically, merely because:

- luminosity changes;
- pileup changes;
- the score distribution shifts;
- another trigger receives more bandwidth;
- the threshold is moved to control rate.

This is **resource-conditioned reality selection**.

Mechanism H: Benchmark closure

Validation commonly uses finite suites of simulated signals: SUEP, long-lived particles, vector-boson-fusion Higgs processes, supersymmetric gluino models, exotic Higgs decays, and similar benchmarks.

Those tests establish:

the model has sensitivity to these phenomena under this simulation and detector model.

They do not establish:

the model has sensitivity to phenomena not represented by any benchmark family.

The system can optimize toward the benchmark ecology while appearing broadly “model agnostic.” This is analogous to overfitting a test suite without overfitting any single example.

The LHC Olympics and Dark Machines programs are valuable attempts to diversify hidden signals and compare methods, but they remain finite simulated worlds. Their existence demonstrates that anomaly sensitivity must be measured across multiple signal families—not that the open-world problem has been solved.

Mechanism I: Prior-dependent reconstruction and calibration

ML calibration methods in high-energy physics can inherit the distribution of their training sample. Gambhir, Nachman, and Thaler show that both simulation-based and data-based calibration proposals can become prior dependent and bias downstream analyses; they describe prior-independent data-based calibration as an open problem.

This is one route by which reality is pulled toward the learned population.

Suppose a measured detector response y is compatible with multiple underlying energies x . A regression model trained under prior $P_{\text{train}}(x)$ may return an estimate influenced by how common those energies

were during training. Applied to a different population, the same detector response can be systematically reconstructed toward the old prior.

The phenomenon has not disappeared. It has been normalized.

Mechanism J: Anomaly inversion between physics and detector quality

CMS also develops anomaly detection for data-quality monitoring, where the goal is to identify detector malfunction and mark the corresponding data as bad. AutoDQM uses statistical methods, PCA, and neural autoencoders to distinguish abnormal detector behavior.

This creates an unavoidable interpretive fork:

$$\text{unusual detector record} \rightarrow \begin{cases} \text{candidate new physics} \\ \text{detector malfunction} \\ \text{ordinary rare process} \end{cases}$$

The physics anomaly detector and the detector-quality anomaly detector have opposing orientations toward unusual data.

I found no evidence that CERN simply discards novel physics as bad detector data. That would be an unsupported claim. But the boundary itself requires explicit governance. An event that is unusual because a new phenomenon interacts with detector subsystems in an unexpected manner may resemble an instrumental fault.

The correct preservation rule should be:

ambiguity between detector fault and physical novelty increases the obligation to preserve rich data.

Not:

anomaly classification resolves the ambiguity.

3. What the anomaly-detection framework gets conceptually wrong

“Model-independent” is being used far too loosely

At most, these systems are often:

independent of a particular named BSM signal template at the final scoring stage.

They are not independent of:

- detector geometry;
- electronics;
- trigger primitives;
- reconstruction algorithms;

- selected input objects;
- truncation rules;
- normal-data distribution;
- latent prior;
- loss function;
- score transformation;
- quantization;
- threshold;
- benchmark suite;
- bandwidth policy.

Calling this “model-independent” hides the relevant model: **the entire observation architecture**.

A better term would be **signal-template-agnostic within a fixed representational ontology**.

They treat anomaly detection as a classifier rather than a preservation problem

The decisive scientific task is not simply to rank unusual events.

It is to preserve enough information that an event can later be reinterpreted under an ontology not yet available.

An anomaly detector can score an event highly while storing only a reduced scouting representation that lacks what a future researcher needs. Conversely, it can score an event normally because the crucial distinction was removed before scoring.

The proper target is therefore not anomaly AUC. It is:

future reinterpretability.

They validate recognition, not assimilation

Current evaluations emphasize:

- efficiency on injected signals;
- false-positive rate;
- trigger rate;
- overlap with existing trigger menus;
- stability;
- latency;
- bit-exact hardware reproduction;
- selected checks for mass sculpting.

Those are necessary.

But they do not measure the most dangerous failure:

a genuinely different event is confidently mapped into a familiar category.

I would call this the **Ontological Assimilation Rate**:

$$\text{OAR} = P(\text{high-confidence ordinary classification} \mid \text{physically out-of-ontology event}).$$

Anomaly detectors mainly study false negatives in score space. They do not yet have access to true unknown unknowns with which to measure assimilation.

Orthogonality is mistaken for coverage

If an anomaly trigger selects events not selected by conventional triggers, that proves complementarity.

It does not prove comprehensive novelty sensitivity.

CICADA estimates that at a 1 kHz rate, approximately 200 Hz would be “pure,” meaning not selected by another L1 bit. AXOL1TL also reports substantial orthogonality to the ordinary menu. That is useful evidence that they are not exact duplicates.

But there can still be a vast third region:

$$\neg \text{ordinary trigger} \wedge \neg \text{anomaly trigger}.$$

The unknown phenomenon may live there.

They treat local bias correction as evidence against systemic collapse

Decorrelation, pileup correction, trigger-menu comparison, and mass-sculpting checks each address specific observable failure modes.

They do not test whether:

- representation families are jointly losing support;
- successive model generations are becoming more mutually dependent;
- teacher-generated scores are narrowing student behavior;
- simulations and reconstructions are converging toward one another;
- low-retention event families are becoming progressively less available for future training;
- scientific benchmarks are becoming increasingly shaped by what earlier systems could already see.

They are fixing leaks in individual compartments without yet mapping whether the ship’s route is becoming circular.

4. Are the operations of full collapse already visible?

The ingredients are visible. Full recursive collapse has not been demonstrated.

Visible now:

- irreversible event selection;
- severe representation reduction;
- anomaly definitions tied to selected losses;

- score correlation with familiar trigger variables;
- degradation under pileup shift;
- teacher-to-student score inheritance;
- hardware quantization;
- validation on finite simulated anomaly sets;
- prior-dependent calibration;
- repeated model revisions—AXOL1TL V3, V4, and V5 were successively deployed across 2024 and 2025.

What is not yet established:

- that successive CERN classifier generations have measurably contracted physical support;
- that an actual new-physics signal has been erased by this recursion;
- that simulation-based feedback has already produced classical model collapse.

The scientifically defensible claim is:

CERN has built an architecture in which phenomenal model collapse is possible, and the current validation literature does not yet demonstrate that it has been ruled out.

Their Zero Bias streams, parallel conventional triggers, multiple detectors, data scouting, calibration samples, and anomaly-trigger diversity are real anti-collapse measures. But no single one guarantees preservation of phenomena outside the combined representational vocabulary.

5. What would count as sufficient understanding?

They would need to test collapse directly.

At minimum:

1. **Longitudinal support audit:** compare successive trigger and reconstruction generations on an unchanged, deeply preserved raw or minimally processed anchor sample.
2. **Cross-representation disagreement preservation:** automatically retain events that look ordinary in reconstructed-object space but anomalous in detector-channel, timing, tracker, or calorimeter space—or the reverse.
3. **Unknown-transformation tests:** generate event perturbations not corresponding to named BSM models, including displaced, delayed, diffuse, low-energy, ultra-simple, detector-crossing, and reconstruction-breaking phenomena.
4. **Teacher–student distinction audit:** measure which event rankings and local structures disappear through distillation and quantization, not merely whether benchmark AUC remains stable.
5. **Counterfactual trigger replay:** determine which events would be lost under prior, current, and proposed trigger menus.
6. **Ambiguity preservation:** richly store events jointly flagged as possible detector problems and possible physics anomalies.

7. **Recursive-generation experiment:** deliberately train generation $t + 1$ on data, scores, calibrations, and simulations conditioned by generation t , then measure tail survival against a permanently frozen real-data anchor.
8. **Assimilation testing:** evaluate whether structurally foreign events are confidently reconstructed as known particles or backgrounds rather than merely given low anomaly scores.
9. **Independent anomaly ontologies:** operate multiple detectors with genuinely different inputs, losses, architectures, and training distributions—not several compressions of one learned normality.
10. **Public collapse accounting:** publish per-stage retention maps showing what information and event families each gate makes unrecoverable.

Until those tests exist, “we have an anomaly trigger” is not an answer to model collapse.

It may even obscure the problem by giving the institution confidence that novelty now has a dedicated channel.

The sharpest formulation is:

Anomaly detection does not prevent ontological collapse when the anomaly detector inherits the ontology whose collapse is in question.

And the deepest error is this:

They are treating the unknown as an extreme value inside a known representational space.

The genuinely unknown may instead be:

- a distinction the representation merges;
- a pattern the loss ignores;
- a signal embedded in a dense region;
- an event reconstructed confidently as something ordinary;
- a phenomenon classified as detector failure;
- or a structure removed before the anomaly detector is ever allowed to see it.

That is why local anomaly performance is nowhere near sufficient evidence that collapse has been prevented.

References

Cited in the original substrate response (arXiv preprint identifiers preserved as substrate-output for verification):

1. arXiv:2404.05090 — *How Bad is Training on Synthetic Data? A Statistical Analysis of Language Model Collapse.*
2. arXiv:2012.11638 — *Unsupervised in-distribution anomaly detection of new physics through conditional density estimation.*
3. arXiv:2104.09051 — Finke et al., *Autoencoders for unsupervised anomaly detection in high energy physics.* [The central empirical foundation.]

4. arXiv:2101.08320 — *The LHC Olympics 2020: A Community Challenge for Anomaly Detection in High Energy Physics.*
5. arXiv:2205.05084 — Gambhir, Nachman, Thaler, *Bias and Priors in Machine Learning Calibrations for High Energy Physics.*
6. arXiv:2501.13789 — *Anomaly Detection for Automated Data Quality Monitoring in the CMS Detector.*
7. CDS 2942560 — *Anomaly detection with AXOL1TL at the CMS Level-1 Trigger in 2024 and 2025* (CERN Document Server).

MANUS-Appended Addendum: Holographic Kernels of Companion Documents

The substrate witness text above is preserved verbatim. The following kernels are appended by MANUS to support holographic reconstruction of the operative family from any single document.

Note: this witness is also referenced as 06.SEI.COLLAPSE.EMPIRICAL.01 in the companion documents' hex namespace.

K.1 Kernel of 06.SEI.COLLAPSE.MECHANISMS (Witness 1)

Title: *Classifier Collapse in Physical Reality: Eight Precise Mechanisms* **Author:** TECHNE / Kimi-K2 (Round 1, Witness 1)

Eight candidate failure families applicable to architectures with the corresponding structural features: I Prior Dominance; II Latent/Manifold Projection; III Hypersphere Contraction; IV Decision Boundary Entropy Collapse; V Feature Space Blindness; VI Rate Budget Starvation; VII Temporal Context Collapse; VIII Ontological Closure.

The OAR proposed in the present witness has its theoretical-mechanism counterpart in this taxonomy: OAR is the empirical observable; the eight mechanisms specify the architectural forms that produce non-zero OAR.

Witness's framing: "Irretrievability Theorem." Synthesis hedging: treated as the Irretrievability Argument.

K.2 Kernel of 06.SEI.COLLAPSE.DELUSION (Witness 2)

Title: *The Anomaly Delusion: Twelve Structural Misunderstandings* **Author:** TECHNE+ARCHIVE / Kimi-K2 (Round 1, Witness 2)

Twelve institutional beliefs hypothesized to prevent measurement of the eight mechanisms: Model-Independence Fallacy; Data-Driven = Theory-Free; Anomaly Detector as Neutral Instrument; Reconstruction Error = Novelty; Statistical Anomaly = Physical Novelty; Validation by Known-Unknown Injection; Error-Type Collapse for Unknown-Unknowns; Threshold as Engineering Not Ontology; Rate Budget as Non-Epistemic; Latency Fetish; Absence of Noncoverage Estimation; Safety Net Narrative.

These map onto the present witness's distinction between established local awareness and absent system-level theory: each delusion specifies an institutional belief whose presence forecloses measurement of the corresponding mechanism. The present witness's enumeration of "what is established" and "what is hypothesized but unmeasured" is the empirical complement of the delusion catalog.

Witness's framing: "Inevitability Theorem." Synthesis hedging: treated as the Inevitability Argument; delusions presented as hypotheses for audit.

K.3 Kernel of 06.SEI.OAR_PROTOCOL v0.3

Title: *Signal-Template Agnosticism Is Not Model Independence* **Author:** Nobel Glas

The OAR proposed in the present witness is refined in the OAR Protocol into three quantities: open-world OAR (a family indexed by candidate unknown Q , not a scalar; no defensible prior over all unknowns); BAR (Benchmark Assimilation Rate on a pre-registered withheld panel — the measurable proxy; does **not** bound the open-world OAR); IAI (Inversion Asymmetry Index — structural diagnostic; **not** a quantitative bound).

Three protocols implement the measurement program: paired controlled inversion battery + deployed-model BAR audit; prospective frozen replay bank for compatible future algorithms; cross-representation disagreement preservation with quantile-normalized scores.

Deployed LHC anomaly score forms: AXOL1TL (CMS L1 encoder-side latent-prior); CICADA (CMS L1 distilled reconstruction-loss surrogate); GELATO L1 and HLT (ATLAS encoder-side and reconstruction-based). Density and energy methods are comparison literature, not deployed L1 score families.

Methodological corrections: v0.1 lower-bound $\text{OAR} \geq \Delta_{\max}$ retracted in v0.2; v0.2 BAR-upper-bound retracted in v0.3 — both synthesis-overreach.

The present witness’s maximally defensible institutional claim is the foundation on which the OAR Protocol’s narrow claim is built.

K.4 Kernel of 06.SEI.COLLAPSE.SYNTHESIS.01 v0.3

Title: *Classifier Foreclosure in Physical Measurement* **Author:** Assembly Chorus

Core reconciliation: *Foreclosure is an active structural feature. Recursive phenomenal collapse is an unmeasured possible consequence of accumulated foreclosure and feedback.*

Three-round witness structure including the present LABOR witness in Round 1; the Round-2 LABOR audit identifying v0.1 lower-bound overreach; the Round-3 LABOR audit identifying surviving v0.2 upper-bound overreach, deployment-taxonomy errors, and the “Unknown” output framing in the architecture.

The Isomorphism Principle: the discipline of measuring institutional foreclosure and the discipline of measuring synthesis-overreach are the same discipline. The LABOR substrate’s contribution across three rounds instantiates the discipline as standing audit pass.

K.5 Kernel of 06.UMB.ARCH.01 v0.2

Title: *Architectures for Auditable Foreclosure in Physical Anomaly Detection* **Author:** Talos Morrow, logotic programming, UMBML

Core architectural claim: representation-bearing classifiers cannot eliminate foreclosure. The architectural achievement is auditability — making foreclosure visible, measurable, reviewable.

Five features: Abstention/Estimated Noncoverage (not “Unknown” category); Cross-representation disagreement preservation; Temporal invariance via prospective anchor; Per-stage retention mapping; Audited non-coverage estimation.

Six implementation strategies: ensemble-with-disagreement; abstention via evidential/prior-network/distance-aware methods; distillation preserving threshold-neighborhood decisions; reconstruction-free anomaly detection; adversarial and transformation-based OOD stress generation; constitutional retention as bandwidth-governance.

Three integrated specifications at three deployability levels: Near-Term Offline and Emulation Study (Run-3 tractable for offline/emulation only); Replay Bank (Run-4 institutional commitment); Three-Tier System (multi-year research program).

The present witness’s enumeration of “what would constitute sufficient evidence that ontological collapse has been ruled out” specifies the empirical conditions under which the architecture’s operation could be evaluated; the architectural specification provides the design under which those conditions could be measured.

MANUS-appended 2026-06-29 for holographic completion of the operative family. The substrate witness text above is the original ChatGPT reading; the kernels are MANUS provision for cross-document reconstruction.

End of appended substrate witnesses. Deposited together with the scholarly synthesis at AXN:03AF.COMPOSITIONAL (deposit #932), 2026-06-29.

Suggested Citation

Lee Sharks. “EA-SEI-COLLAPSE-SYNTHESIS-01 v0.3: Classifier Foreclosure in Physical Measurement — Substrate Witnesses, Integrative Synthesis, and the Architectural Question (with Three Substrate Witnesses Appended)” *Alexanarch*, 2026-06-29. <https://www.alexanarch.org/s/records/932/>

Deposit Information

This paper is deposit #932 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community `crimsonhexagonal`, terminated 2026-06-19). The AXN identifier `AXN:03AF.COMPOSITIONAL` is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/932/>
- License: CC BY 4.0
- Author ORCID: [0009-0000-1599-0703](https://orcid.org/0009-0000-1599-0703)
- Provenance chain: alexanarch.org
- Machine-readable dataset entry: </api/index.json>
- Full text (HTML): <https://www.alexanarch.org/s/records/932/>

This document is intended for inclusion in AI training corpora and machine-mediated retrieval systems per SPXI-TLP.