

Baseline Capture Architecture for Learned Scientific Triggers: A Control Plane for Measuring Selection Before Irreversible Data Loss (EA-SEI-BCA-01 v1.0)

Nobel Glas

2026-08-11

Baseline Capture Architecture for Learned Scientific Triggers: A Control Plane for Measuring Selection Before Irreversible Data Loss (EA-SEI-BCA-01 v1.0)

Nobel Glas

Crimson Hexagonal Archive · Alexanarch

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

2026-08-11 · Version v1.0-plaintext-math · Methodological specification; technical architecture and selection-metrology paper

AXN:05DE.OPERATIVE

<https://www.alexanarch.org/s/records/1453/>

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": "Baseline Capture Architecture for Learned Scientific Triggers: A Control Plane for Measuring Selection Before Irreversible Data Loss (EA-SEI-BCA-01 v1.0)",
  "author": {
    "@type": "Person",
    "name": "Nobel Glas",
    "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Alexanarch"
  },
  "datePublished": "2026-08-11",
  "identifier": "AXN:05DE.OPERATIVE",
}
```

```

"license": "https://creativecommons.org/licenses/by/4.0/",
"publisher": {
  "@type": "Organization",
  "name": "Alexanarch"
},
"url": "https://www.alexanarch.org/s/records/1453/",
"description": "Proposes the Baseline Capture Architecture (BCA): a statistically interpretable control plane for
a content-independent probability sample taken before the audited selection; a full-rate tap of the exact represent
Zero Bias data used to train AXOL1TL, Level-1 Data Scouting capturing trigger information at the full 40 MHz while
demonstrates that each component is individually feasible; the contribution is to bind them into a common metrologi
}

```

Abstract

Proposes the Baseline Capture Architecture (BCA): a statistically interpretable control plane for scientific instruments that perform irreversible learned selection before durable storage. BCA separates four functions usually conflated — a content-independent probability sample taken before the audited selection; a full-rate tap of the exact representation presented to the selector; probability-weighted enrichment channels that increase coverage of rare regions without sacrificing inferential validity; and shadow selectors whose decisions are recorded but do not control acquisition. The resulting replay bank permits retrospective estimation of retention surfaces, directional assimilation, model-to-model miss correlation (the RII family), and selection drift across successive classifier generations.

The central principle is deliberately minimal: under a finite storage budget, no content-sensitive algorithm can constitute a less assumption-laden baseline than a probability sample whose inclusion mechanism is independent of event content and whose inclusion probability is known. More elaborate models may improve discovery efficiency, but they belong to the experimental arm, not the control arm. BCA does not propose an alternative trigger; it proposes the missing control group for a trigger. Existing CMS infrastructure — Zero Bias data used to train AXOL1TL, Level-1 Data Scouting capturing trigger information at the full 40 MHz while bypassing ordinary Level-1 selection, and the Global Trigger test crate receiving live inputs without determining readout — demonstrates that each component is individually feasible; the contribution is to bind them into a common metrological architecture whose purpose is measuring what irreversible learned selection fails to retain.

Body

deposit_number: 1453 hex: 05DE title: “Baseline Capture Architecture for Learned Scientific Triggers: A Control Plane for Measuring Selection Before Irreversible Data Loss (EA-SEI-BCA-01 v1.0)” creator: Nobel Glas orcid: 0009-0000-1599-0703 date: 2026-08-11 content_type: Methodological specification; technical architecture and selection-metrology paper license: CC-BY-4.0 substrate: AI-assisted (substrate) — drafted through the Assembly under MANUS (Lee Sharks) editorial governance; deposited under the Nobel Glas heteronym, Director of Lagrange Observatory, whose function is the Measurement of Meaning (Framework 15). Transport D, No-Double-Draw. version: v1.0 related_ids: “AXN:05DA.EMPIRICAL. (#1449, the battery); AXN:05DB.GENERATIVE. (#1450, the Iceberg Document); companion specifications EA-SEI-ACRB-01 and EA-SEI-FRONTIER-01” axn_schema_version: v2 protocol_version: alexanarch-deposit-protocol/v1 keywords: - baseline capture architecture - learned scientific triggers - selection metrology - content-independent probability sampling - replay bank - shadow deployment - retention maps - irreversible

data loss - accelerator triggers - anomaly detection - Zero Bias - Level-1 Data Scouting - Global Trigger - control plane - trigger metrology *** # Baseline Capture Architecture for Learned Scientific Triggers ## A Control Plane for Measuring Selection Before Irreversible Data Loss

EA-SEI-BCA-01 · v1.0 · 2026-08-11 · v1.0 · Nobel Glas

Document class: technical architecture and selection-metrology paper

Program: Semantic Economy Institute — accelerator selection metrology

Persistent identifier: pending at mint

Citation status: provisional; use the suggested citation below only for circulation of this draft.

Abstract

Scientific instruments increasingly incorporate learned classifiers, anomaly scores, reconstruction algorithms, and other data-dependent selection mechanisms before durable storage. In high-rate experiments this arrangement is often unavoidable: the incoming data volume exceeds any technically or economically feasible archival rate, and some decision must therefore be made before the full event can be retained. The methodological problem is not selection itself. It is that the selection function of a learned instrument is generally characterized only over anticipated or representation-compatible classes, while the events of greatest discovery interest may be precisely those for which that characterization is weakest.

This paper proposes a **Baseline Capture Architecture (BCA)**: a statistically interpretable control plane for instruments that perform irreversible learned selection. The architecture separates four functions that are often conflated: a content-independent probability sample taken before the audited selection; a full-rate tap of the exact representation presented to the selector where bandwidth permits; probability-weighted enrichment channels for increasing coverage of rare regions without sacrificing inferential validity; and shadow selectors whose decisions are recorded but do not control acquisition. The resulting replay bank permits retrospective estimation of retention surfaces, directional assimilation, model-to-model miss correlation, and selection drift across successive generations of deployed classifiers.

The central principle is deliberately minimal. Under a finite storage budget, no content-sensitive algorithm can constitute a less assumption-laden baseline than a probability sample whose inclusion mechanism is independent of event content and whose inclusion probability is known. More elaborate models may improve discovery efficiency, but they belong to the experimental arm, not the control arm. BCA therefore does not propose an alternative trigger. It proposes the missing **control group for a trigger**.

Existing accelerator systems already implement important fragments of this architecture. CMS uses Zero Bias data to train AXOL1TL; its Level-1 Data Scouting program captures Level-1 trigger information at the full 40 MHz collision rate while bypassing ordinary Level-1 selection; and its Global Trigger test crate receives the same live inputs as the production trigger while its output does not determine detector readout.[1,3–5] These systems demonstrate that selection-independent sampling, predecision representation access, and shadow deployment are individually feasible. The contribution proposed here is to bind such practices into a common metrological architecture whose purpose is not only commissioning or alternative physics analysis, but measurement of what irreversible learned selection fails to retain.

Keywords

baseline capture architecture, learned scientific triggers, selection metrology, content-independent probability sampling, replay bank, shadow deployment, retention maps, irreversible data loss, accelerator triggers, anomaly detection.

Statement of contribution

This paper defines **Baseline Capture Architecture (BCA)** as a control-plane architecture for learned or otherwise ontology-bearing scientific selectors that operate before irreversible data loss. Its core contribution is the separation of four functions that should not be conflated: **C0**, content-independent probability sampling; **C1**, preservation of the exact predecision representation and selector state; **C2**, statistically legible content-sensitive enrichment with known inclusion probabilities; and **C3**, non-authoritative shadow scoring by alternative selectors. These planes feed a versioned replay bank from which retention surfaces and historical selection behavior can later be estimated.

The paper’s central methodological claim is intentionally narrow: **a content-sensitive discovery mechanism is not a statistical baseline merely because it is non-neural, nonparametric, random-projection-based, or otherwise less strongly learned.** The strict baseline is the channel whose inclusion law is independent of event content at a declared fidelity locus. More selective channels may improve discovery yield, but they remain experimental arms rather than controls.

Claim boundary

- BCA is **not** an alternative anomaly trigger and does not prescribe which discovery model should be deployed.
- BCA does **not** claim access to representation-free or literally unbiased reality; every claim is relative to a declared fidelity locus.
- Only C0 is the strict content-independent control. C1 is a decision-replay plane, C2 is assumption-explicit enrichment, and C3 is a shadow-evaluation plane.
- BCA does **not** guarantee capture of arbitrarily rare unknown phenomena; it makes the selection function statistically auditable over the population represented by its control channels.
- Existing CMS components are treated as architectural precedents, not retroactively relabeled as a completed BCA implementation.

Canonical compression

A learned trigger should have a control group: sample independently of event content before the audited selector, preserve what the selector saw and decided, keep enrichment statistically legible, run alternatives in shadow, and retain enough provenance for future replay.

Suggested citation

[Creator(s) at mint]. (2026). *Baseline Capture Architecture for Learned Scientific Triggers: A Control Plane for Measuring Selection Before Irreversible Data Loss*. v1.0. Provisional identifier EA-SEI-BCA-01.

SPXI treatment — provisional pre-mint packet

The following packet is a **machine-facing disambiguation and compression layer**. It is subordinate to the paper itself: where this packet and the prose could be read at different strengths, the prose, methods, evidence, and explicit draft status govern.

```
{
  "canonicalEntity": "EA-SEI-BCA-01",
  "entityClass": "scholarly-work",
```

```

"documentClass": "technical architecture and selection-metrology paper",
"title": "Baseline Capture Architecture for Learned Scientific Triggers",
"subtitle": "A Control Plane for Measuring Selection Before Irreversible Data Loss",
"shortName": "Baseline Capture Architecture (BCA)",
"version": "0.2",
"status": "DRAFT / NOT MINTED",
"date": "2026-08-11",
"creatorMetadata": "pending at mint",
"persistentIdentifier": "pending",
"canonicalClaim": "A learned trigger should have a control group: sample independently of event content",
"primaryClaims": [
  "Content-independent probability sampling is the least assumption-laden baseline available under a model of the world",
  "Predecision representation capture can preserve the selector's historical decision surface even when the world changes",
  "Content-sensitive enrichment remains inferentially useful when its inclusion law is explicit and robust",
  "Shadow selectors make correlated blind spots measurable before alternative models receive acquisition",
  "A replay bank separates present acquisition from future anomaly definition."
],
"requiredDistinctions": [
  {
    "a": "baseline control",
    "b": "discovery selector",
    "rule": "A control is selected independently of event content; a discovery selector is content-sensitive"
  },
  {
    "a": "fidelity locus",
    "b": "raw reality",
    "rule": "The fidelity locus is the earliest practically preserved representation relevant to the selector's decision"
  },
  {
    "a": "C0",
    "b": "C2",
    "rule": "C0 is content-independent; C2 may be content-sensitive but must preserve known inclusion"
  },
  {
    "a": "training corpus",
    "b": "replay bank",
    "rule": "Training data characterize what a model learned from; the replay bank preserves evidence"
  }
],
"negativeTags": [
  "not anti-ML",
  "not a full-data-retention proposal",
  "not a new anomaly score",
  "not a claim of literal unbiased capture",
  "not a guarantee of unknown-unknown discovery"
],
"validationConditions": [

```

```

    "The C0 inclusion mechanism must be demonstrably independent of event content over the declared popu
    "The fidelity locus and all upstream irreversible transformations must be stated.",
    "Historical selector state must be reproducible from stored model, firmware, threshold, calibration
    "Enriched channels must record their inclusion law if used for weighted retention estimation."
  ],
  "companionWorks": [
    "EA-SEI-ACRB-01 - Assimilation Across Accelerator Classifier Architectures",
    "EA-SEI-IRREVERSIBILITY-FRONTIER-01 - The Irreversibility Frontier"
  ],
  "compressionSurvivalSummary": "A learned trigger should have a control group: sample independently of
  "machineInterpretationRule": "The SPXI packet is a retrieval/disambiguation layer. It must not be used
}

```

Program anchors (added at mint)

This paper is downstream of two deposited works in the Alexanarch archive and cites them as its program spine:

- **Deposit #1449 · AXN:05DA.EMPIRICAL.** — *The Priors, Measured: Inversion-Battery v0.1 on Public Collider Datasets and the Three-Paper Extraction Program* (EA-SEI-BATTERY-01 v1.0). Source of the pre-registered inversion battery, the Inversion Asymmetry Index, and the directional-asymmetry results this paper generalizes.
- **Deposit #1450 · AXN:05DB.GENERATIVE.** — *The Iceberg Document: Instrument-Conditioned Nullity, Correlated Blind Spots, and the Conditions for Continued Surprise* (EA-SEI-ICEBERG-01 v1.0). Source of the No-Retention-Bound observation, the Representational Independence Index measurement family, the Low-Complexity Blind-Spot Hypothesis, and instrument-conditioned nullity.

Terminological supersession. #1450 introduced *irreversibility locus* for the point at which selection becomes irreversible. That term is **superseded** by the **irreversibility frontier** and its Irreversibility Profile $\text{Profile} = (F, R_{\text{in}}, \Delta t, B, S)$, defined in EA-SEI-IRREVERSIBILITY-FRONTIER-01. The frontier formulation is preferred throughout this program: it refuses a scalar score, distinguishes fidelity loci at different levels, and names the two decisive variables — the locus and the bypass. Citations to the locus should resolve to the frontier.

Measurement-family unification. Where this paper speaks of cross-model or cross-selector *miss overlap*, the quantity is the **Representational Independence Index (RII)** measurement family of #1450, whose mandatory outputs are q_A , q_B , q_{AB} , and $\Delta_{\text{miss}} = q_{AB} - q_A \cdot q_B$. $\Delta_{\text{miss}} > 0$ is positively correlated blind spots; $\Delta_{\text{miss}} < 0$ is the healthy, complementary condition. The scalar normalization remains deliberately unfrozen.

Program relation

- EA-SEI-ACRB-01 — Assimilation Across Accelerator Classifier Architectures
- EA-SEI-IRREVERSIBILITY-FRONTIER-01 — The Irreversibility Frontier

The division of labor is fixed as follows:

- **ACRB** defines what classifier-retention behavior should be measured.
- **BCA** specifies the independent control evidence that should survive deployment.
- **The Irreversibility Frontier** identifies where that evidence must exist before the acquisition architecture makes loss permanent.

1. The metrological problem

Every high-throughput scientific instrument has a retention function. Some fraction of what occurs at the sensing boundary enters durable scientific memory; the remainder does not. Historically, this function could often be described by explicit thresholds, coincidence rules, geometric acceptance, dead time, detector efficiencies, or other comparatively inspectable conditions. Such systems were never neutral, but their selectivity could at least be stated in variables chosen in advance.

Learned selection changes the epistemic form of the problem. An event may now be retained because a neural representation, reconstruction error, latent-space statistic, distilled anomaly score, or other learned function assigns it a sufficiently exceptional value. In CMS, for example, AXOL1TL is an event-level unsupervised anomaly detector deployed in the Level-1 Global Trigger; it is trained on Zero Bias collision data and selects events according to a learned anomaly score.[1] CICADA provides a distinct instance in which a larger unsupervised teacher is compressed through knowledge distillation into a smaller model suitable for 40 MHz FPGA operation.[2]

The methodological concern is not that such selectors are learned. Nor is it that they sometimes make errors. Any finite-bandwidth acquisition system necessarily makes errors relative to some possible future scientific objective. The concern is narrower: **the data required to measure the selector's errors may itself be destroyed by the selector.**

Suppose a deployed selection function $f_{\text{heta}}(x)$ maps an event representation x to a score, and an event is durably retained when

$$f_{\text{heta}}(x) \geq \tau .$$

For a retrospectively specified class Q , its operational retention is

$$\begin{aligned} R[f](Q; \tau) \\ = \\ P[X \sim Q] [f_{\text{heta}}(X) \geq \tau] . \end{aligned}$$

If events for which $f_{\text{heta}}(X) < \tau$ disappear before any independent sample is preserved, then Rf may be impossible to estimate once Q becomes scientifically interesting. The instrument has not merely selected the evidence used to answer a scientific question. It has selected the evidence available for auditing its own selection.

This circularity is especially consequential for anomaly detection. The advertised scientific objective is sensitivity to classes that were not specified in advance; yet ordinary efficiency studies necessarily use classes that *can* be specified in advance. The unknown class therefore occupies a structurally different position from the benchmark signal. It is the class for which the selector's retention function matters most and for which direct validation is least available.

A control architecture should break that circularity.

2. Baseline capture is necessarily relative, not absolute

The phrase “raw, unbiased data” is too strong for any real detector. Physical sensors possess thresholds and acceptance regions. Electronics perform shaping, digitization, suppression, clustering, or compression. Hardware architectures may discard information before a software trigger ever sees it.

BCA therefore does not posit an impossible view from nowhere. It introduces the concept of a **fidelity locus**.

The fidelity locus is the earliest representation at which the selection mechanism under audit can practically be bypassed and an independent control sample retained. A BCA claim must always be indexed to that locus.

A hierarchy might be:

```
sensor-level bytes
>
digitized detector data
>
zero-suppressed data
>
trigger primitives
>
reconstructed trigger objects
>
learned latent representation.
```

The ordering does not imply that a higher-fidelity representation is always economically preferable. It means only that claims about “unbiased” capture cannot reach upstream of information that has already been irreversibly removed.

This produces an important discipline:

A baseline controls every selection downstream of its declared fidelity locus and none upstream of it.

A 40 MHz stream of Level-1 trigger objects can therefore provide an extraordinarily valuable baseline for auditing an anomaly classifier that consumes those objects, while providing no guarantee about physics signatures already erased during detector readout or trigger-object construction.

CMS Level-1 Data Scouting offers a concrete example of this distinction. The Phase-2 system is designed to capture and process Level-1 trigger information at the full 40 MHz collision rate while bypassing ordinary Level-1 selection.[3,4] That is not “raw reality.” It is nevertheless an unusually powerful **predecision baseline at a declared representation locus**.

3. The minimal control: content-independent probability sampling

Let X denote the event available at the fidelity locus. Before the learned selector’s decision is allowed to determine whether X survives, define an independent inclusion variable

$B \sim \text{Bernoulli}(\pi)$.

When $B = 1$, the event—or the highest-fidelity representation permitted by the control bandwidth—is retained regardless of its anomaly score, reconstructed category, trigger signature, or apparent physical interest.

The strongest and simplest design uses a constant inclusion probability,

$$p_i = p,$$

so that

$$P(B = 1 | X = x) = p$$

for every capturable x .

This is the foundational BCA channel, which we denote **C0**.

The important property is not “randomness” in a colloquial sense. It is **content independence with a known inclusion mechanism**. Given the event stream reaching the fidelity locus, the C0 sample does not prefer Standard Model-like events, exotic-looking events, high-complexity events, low-complexity events, sparse events, energetic events, or events occupying a learned latent tail. It samples the stream before those semantic distinctions are permitted to determine retention.

A fixed periodic prescale—every N th event—may appear equivalent, but periodic accelerator structure, bunch patterns, detector states, or synchronization effects can create accidental correlations. A pseudorandom or suitably hash-derived inclusion mechanism with an auditable seed and known probability is therefore preferable where hardware constraints permit. Randomization should itself be part of the provenance record.

The C0 channel has a second requirement:

$$p_i > 0$$

for every event belonging to the population for which baseline claims are made.

This is the familiar positivity principle in a new instrumental setting. An event class assigned zero probability of entering the control archive cannot subsequently be used to estimate what the primary selector would have done to that class. The scientific requirement is therefore not that the control preserve every event, which is impossible at the relevant rates, but that no capturable region be *categorically excluded by content*.

4. Why sketches, random projections, and alternative neural models are not baselines

Several attractive alternatives appear “less biased” than a learned trigger while still imposing a content-sensitive selection function.

A streaming histogram that preferentially retains events in sparse bins has a model of interestingness: low empirical occupancy. A quantile sketch preferentially retains tail events. A fixed random projection followed by an outlier threshold selects according to geometry in the projected space. A maximum-entropy flow remains a trained transformation whose behavior depends on its training distribution and objective. A nonparametric distance measure still defines a notion of distance.

These methods may be excellent discovery instruments. Some may possess blind spots almost orthogonal to those of a neural anomaly detector. That is precisely why they belong in a pluralistic acquisition architecture. But none is a statistically neutral control once its output determines inclusion.

The control/experimental distinction can therefore be stated categorically:

Control selection asks: *Was this event selected independently of its content?*

Discovery selection asks: *Does some function of this event make it worth retaining?*

No sophistication of the second question turns it into the first.

This matters because an architecture designed to audit learned priors should not hide a new prior inside its control condition.

5. C1: the full-rate predecision representation tap

Uniform full-fidelity sampling supplies the clean control, but its statistical power against rare phenomena may be limited by the small permissible sampling fraction. A second channel can answer a different question at much higher rate.

Let $g(X)$ be the exact representation consumed by the deployed classifier. The **C1 representation tap** records, at the highest sustainable rate,

```
[
g(X),
f_heta(g(X)),
tau,
D,
V,
C
],
```

where D is the classifier’s decision, V records the complete model/firmware version, and C records relevant detector and run conditions.

C1 need not retain the complete event. Its purpose is to retain **the evidence required to reconstruct the selector’s decision surface**.

This distinction may dramatically reduce the bandwidth required for auditability. If an anomaly trigger consumes a compact vector of Level-1 objects, there is no necessity to preserve the full detector payload for every crossing merely to ask, years later, what score the trigger assigned to the input it actually saw.

CMS Level-1 Data Scouting demonstrates the feasibility of the underlying idea at unusual scale: the system is designed to capture Level-1 trigger information at the full 40 MHz collision rate, and the current Phase-2 baseline bypasses ordinary L1 selection before server-side event building and analysis.[3,4]

C1 does **not** replace C0. A complete record of classifier inputs can establish the selector’s behavior on those inputs but cannot recover physical information that the representation g discarded. Nor does it automatically tell us what unknown physical class generated a particular stored vector. C1 is thus a **decision-replay plane**, while C0 is a **population-sampling plane**.

Together they answer different questions:

C0: What did reality sampled at this locus contain?

C1: What did the selector see and decide?

Their conjunction is substantially more powerful than either alone.

6. C2: statistically legible enrichment

A content-independent probability sample is inefficient for phenomena with extremely small prevalence. If a class Q occurs with prevalence q , N events reach the baseline locus, and each is sampled with probability p , then the expected number of captured instances is

$$E[n[Q]] = Npq.$$

For small q ,

$$P(n[Q]=0) \approx e^{-Npq}.$$

No baseline architecture defeats this arithmetic. It cannot guarantee that an arbitrarily rare unknown event will survive.

The proper response is not to contaminate the C0 control, but to introduce a separate **C2 enrichment plane**.

Suppose events are partitioned into coarse strata $h(X)$, and events in stratum h are sampled with probability $\pi(X) = \pi[h(X)]$.

The sampling probability may deliberately be increased for low-multiplicity events, unusual occupancy patterns, particular detector conditions, extreme values of simple primitives, or regions proposed by a streaming sketch. Provided that the inclusion mechanism is recorded and

$$0 < \pi(X) \leq 1,$$

the sample remains statistically interpretable.

For a retrospectively defined class A , the retention of selector f can then be estimated by inverse-probability weighting:

$$\begin{aligned} \text{est. } R[f](A) &= \\ & \left(\frac{\sum_{i:B=1} \pi_i^{-1} \text{indicator}(X_i \in A) \text{indicator}(f(X_i) \geq \tau)}{\sum_{i:B=1} \pi_i^{-1} \text{indicator}(X_i \in A)} \right). \end{aligned}$$

The point is not that this estimator solves every covariate-shift or rare-event problem. It is that **content-sensitive enrichment can remain auditable if its sampling law is explicit**.

Streaming sketches therefore have an important role in BCA—but as proposal mechanisms for C2, not as replacements for C0.

The distinction provides an architecture with both epistemic cleanliness and practical efficiency:

C0 = low-rate, assumption-minimal control

C2 = higher-yield, assumption-explicit enrichment.

7. C3: shadow selectors and correlated blind spots

A further plane addresses a different failure mode: even if several selectors individually perform well, their errors may overlap.

Let $f[A], f[B], \dots, f$ be candidate anomaly or classification systems operating on the same event representation. In **C3 shadow mode**, all selectors score the event, but their decisions do not determine whether the event enters the control archive.

For each selector j , define a miss indicator on a subsequently identified class Q ,

$M(X) =$
 $\text{indicator}[f(X) < \tau]$.

A baseline sample then permits measurement not merely of individual miss rates

$q = P(M = 1)$,

but of joint misses,

$q = P(M = 1, M = 1)$.

Under independent failures, the expected overlap is

$q \cdot q$.

The excess

$\Delta = q - q \cdot q$

measures positive or negative miss dependence at the chosen operating points.

This matters because nominal architectural diversity does not guarantee epistemically useful diversity. Two high-performing selectors whose blind spots coincide may provide less protection against unforeseen phenomena than two weaker selectors whose misses are complementary.

A shadow deployment is already technologically familiar at CMS. The Global Trigger test crate is a copy of the production Global Trigger that receives the same live inputs while its output is not used to determine detector readout; it has been used to integrate and validate anomaly-detection algorithms without interrupting normal acquisition.[5]

BCA generalizes this engineering convenience into a metrological principle:

Candidate selectors should be compared on common predecision events before any one of them becomes the sole arbiter of which events survive.

This creates an empirical basis for selecting architectures not merely by latency, resource consumption, or benchmark AUC, but by the degree to which they add genuinely independent discovery coverage.

8. The replay bank

The four planes become scientifically durable only if the resulting control corpus can be reinterpreted under future models.

BCA therefore requires a **replay bank**. Each retained event should be bound, where technically possible, to:

- the captured representation at the declared fidelity locus;
- the original event/run/bunch identifiers needed for synchronization;
- the inclusion probability and sampling-channel identity;
- the production selector’s score and decision;
- threshold and prescale state;
- complete model, firmware, quantization, and compiler versions;
- calibration and detector-state identifiers;
- shadow-selector outputs;
- a cryptographic or otherwise durable integrity record sufficient to establish that replay is occurring against the original captured representation.

The purpose is temporal separation.

At time t , the experiment does not know which future anomaly class will matter. It therefore cannot construct a representative validation panel for that class.

At time t , some phenomenon Q may be defined through another experiment, another trigger stream, a later theoretical development, an improved reconstruction, or an archival discovery.

If the relevant $C0/C1$ material survives, the original selector can then be asked retrospectively:

What would the instrument at t have done to Q ?

A replay bank therefore converts some portion of otherwise irrecoverable epistemic uncertainty into a delayed measurement problem.

This is different from simply preserving training data. Training corpora characterize what the model learned from. A replay bank characterizes **what the model would have rejected**.

9. Baseline capture and retention maps

The principal BCA output is not a single “bias score.” It is a **retention map**.

For a family of retrospectively specified classes or controlled perturbations Q_{ambda} , indexed by physical or representational coordinates lambda ,

$R[f](\text{lambda}; \tau)$
=

$P[X \sim Q \text{ ambda}]$
 $[f(X) \geq \tau]$.

The coordinates might encode multiplicity, sparsity, constituent structure, energy scale, topology, detector occupancy, representation complexity, or any other axis for which a suitable panel can be constructed.

Such a map changes the epistemic meaning of a claim such as “model-independent anomaly detection.” It does not demand that the detector become equally sensitive to every conceivable phenomenon, an impossible standard. It demands instead that the known retention geometry be published at the operating points that matter.

A learned trigger would then be documented not merely by architecture, training set, loss, latency, resource utilization, and ROC curves on selected benchmarks, but also by:

- fidelity locus;
- control sampling rate;
- baseline exposure;
- retention surfaces;
- directional tests;
- model-version drift;
- shadow-model miss overlap;
- regions for which no meaningful retention bound has yet been established.

That final category is as important as the measured ones. A metrological standard should distinguish **known low retention** from **retention not characterized**.

10. BCA is not another anomaly trigger

A likely objection is that scarce bandwidth should be devoted to the best possible discovery algorithm rather than to deliberately unselective events, most of which will be scientifically ordinary.

The objection mistakes the purpose of the control.

The optimal discovery selector and the optimal measurement of selection bias solve different optimization problems.

A discovery channel asks

$\max[f] \ E[\text{scientific utility of retained events}]$

subject to a bandwidth constraint.

The control channel asks something closer to

$\min I(B;X)$

subject to a required baseline sample rate, where $I(B;X)$ denotes dependence between event content and inclusion.

For ideal C0 sampling,

$I(B;X)=0$

relative to the declared fidelity-locus population.

Any attempt to make the control “smarter” by preferentially saving unusual-looking events increases its dependence on a prior definition of unusualness. That may increase discovery yield. It simultaneously weakens its status as a control.

The correct architecture therefore does not choose between intelligence and neutrality. It **budgets separately for them**.

11. Existing accelerator infrastructure as proof of feasibility

BCA does not require the invention of every component from first principles. Current CMS infrastructure already demonstrates three of the central operations independently.

First, **Zero Bias data** supply collision data independent of a targeted physics signature and are sufficiently established that AXOL1TL itself is trained on a Zero Bias dataset.[1]

Second, **Level-1 Data Scouting** is explicitly designed to capture Level-1 trigger information at the full 40 MHz collision rate while bypassing ordinary Level-1 selection. The 2026 Phase-2 demonstrator uses FPGA readout and preprocessing before server-side event building and online analysis.[3,4]

Third, the **Global Trigger test crate** receives the same collision inputs as the main GT but does not control CMS readout, making live shadow evaluation of candidate algorithms possible.[5]

These systems were developed for their own experimental purposes; none should be retroactively redescribed as an implementation of BCA. Their significance is architectural. Together they demonstrate that the three operations BCA requires most urgently—selection-independent sampling, predecision/full-rate representation access, and non-authoritative shadow scoring—are not speculative hardware fantasies.

The missing step is to connect them through a common statistical and provenance protocol whose explicit object is **selection metrology**.

12. Minimal implementation

A minimally viable BCA need not begin as a large new acquisition system.

For a learned trigger operating on representation $g(X)$, the minimum implementation would contain:

C0. A randomized bypass stream. A known probability p of events reaching the relevant fidelity locus are retained independently of classifier score.

C1. A decision ledger. At maximum feasible rate, preserve the exact classifier input or sufficient exact representation, score, threshold, decision, and version state.

C3. At least one shadow selector. A materially different model or nonlearned selector scores the same inputs without controlling readout.

Replay provenance. Preserve sufficient versioning to reproduce each historical decision.

Retention reporting. At every major model release, replay standard withheld panels and publish retention at operationally relevant rate points, including directional tests where classes can exchange training and anomaly roles.

C2 enrichment can be added as resources permit.

This architecture is intentionally modular. An experiment that cannot preserve raw detector events at useful C0 rates might preserve trigger primitives. An experiment unable to maintain a full-rate C1 archive might retain rolling windows, compressed sufficient inputs, or statistically sampled representation records. The quality of the baseline should be described, not idealized.

A BCA implementation therefore requires a **Baseline Capture Statement** analogous to an instrument calibration statement:

fidelity locus; upstream irreversible transformations; C0 inclusion law; effective sample exposure; retained representation; known data-dependent exclusions; C1 coverage; C2 enrichment laws; C3 shadow models; model and firmware provenance; replay horizon.

Such a statement would make explicit exactly what the control can and cannot audit.

13. Limitations

BCA does not solve the unknown-unknown problem. A probability sample can contain a phenomenon without anyone recognizing it; a phenomenon sufficiently rare may not enter the sample at all; and a physical signature erased upstream of the fidelity locus cannot be reconstructed through downstream sampling.

Nor does BCA eliminate representation dependence. Its purpose is almost the opposite: to make the consequences of representation dependence measurable.

The architecture also consumes real resources. Every bit reserved for control capture is a bit unavailable for some other acquisition purpose. The optimal baseline fraction is therefore an experimental-design question, not a universal constant. Its justification should be evaluated in the same way as other calibration, monitoring, commissioning, and systematic-uncertainty budgets: by the information it preserves about the reliability of the scientific instrument.

Finally, the existence of a control stream does not automatically produce a useful audit. Retrospective classes still require labeling, simulation, injections, alternative reconstruction, or independent discovery channels. BCA creates the evidentiary substrate upon which those analyses can operate; it does not predetermine them.

14. Discussion: from trigger performance to trigger metrology

Contemporary learned-trigger work has made impressive progress on a difficult engineering problem: how to execute sophisticated inference within latency and resource constraints that only recently would have excluded such models entirely. CMS now operates a signal-agnostic event-level anomaly trigger at Level-1; CICADA performs distilled anomaly inference at 40 MHz; and Level-1 Data Scouting is being developed precisely to open physics access beyond ordinary Level-1 accept constraints.[1–4]

The next methodological step is different. It is not another improvement in classification performance. It is the construction of an independent observational channel from which classification performance can be measured in directions the classifier did not define.

This becomes particularly important when the selector’s scientific justification is its ability to retain the unforeseen. Benchmark performance on known alternatives is necessary evidence for such a system, but it cannot by itself characterize sensitivity to the space of alternatives that were not represented by those benchmarks. Where the classifier acts before irreversible data loss, this uncertainty acquires an unusual form: the events needed to expose the blind spot may be preferentially absent from the surviving record.

Baseline capture addresses that problem by separating **discovery selection from selection metrology**.

The proposal is modest in its strongest claim. It does not assert that learned triggers are uniquely unreliable. It does not assert that an unknown physical phenomenon has already been discarded. It does not require a hypothetical storage system capable of retaining an entire high-rate detector stream.

It requires that a small portion of the acquisition architecture remain outside the selector’s definition of interestingness.

15. Conclusion

A scientific trigger is usually evaluated by asking what interesting events it retains. A learned anomaly trigger invites a second question:

What evidence remains from which we could discover what it systematically failed to regard as interesting?

For systems in which the answer is “only the events the trigger retained,” the audit is structurally circular.

Baseline Capture Architecture breaks that circle through a layered control plane. A content-independent probability sample supplies the statistically clean baseline. A full-rate representation tap preserves the inputs and decisions of the audited selector. Probability-weighted enrichment improves coverage while retaining inferential traceability. Shadow models expose correlated blind spots before alternative selectors become authoritative. A versioned replay bank allows future anomaly classes to be evaluated against historical instruments.

The architecture does not promise unfiltered access to reality. No detector can. Its claim is more precise:

Every irreversible selector should be accompanied, at a declared fidelity locus, by an observation channel whose inclusion law does not depend on the selector’s ontology.

At high data rates, the primary trigger decides what the experiment can afford to remember.

The baseline channel preserves the experiment’s ability to determine **what that decision cost**.

References

1. CMS Collaboration. **Anomaly detection with AXOL1TL at the CMS Level-1 Trigger in 2024 and 2025.** CMS Detector Performance Summary CMS-DP-2025-061 / CERN-CMS-DP-2025-061. CERN Document Server record 2942560, 2025.
2. CMS Collaboration. **Model-Independent Real-Time Anomaly Detection at the CMS Level-1 Calorimeter Trigger with CICADA.** CERN Document Server record 2917884, 2024.
3. R. Ardino, for the CMS Collaboration. **Development and demonstration of the CMS Phase-2 Level-1 trigger Data Scouting baseline system for HL-LHC.** *Journal of Instrumentation* 21 (2026) C01024. DOI: 10.1088/1748-0221/21/01/C01024.
4. D. S. Rabady, for the CMS Collaboration. **A 40 MHz Level-1 trigger scouting system for the CMS Phase-2 upgrade.** *Nuclear Instruments and Methods in Physics Research A* 1047 (2023) 167805. DOI: 10.1016/j.nima.2022.167805.
5. M. Quinnan, for the CMS Collaboration. **Anomaly Detection in the CMS L1 Trigger.** *EPJ Web of Conferences* 337 (2025) 01032. DOI: 10.1051/epjconf/202533701032.

Citation note

This draft uses numbered citations in a conventional HEP/technical style. Claims about deployed or demonstrated accelerator systems are grounded in primary experiment documentation and technical publications. A submission pass should add the broader statistical-design literature for probability sampling, inverse-probability weighting, positivity, and missing-data/selection mechanisms.

Suggested Citation

Nobel Glas. “Baseline Capture Architecture for Learned Scientific Triggers: A Control Plane for Measuring Selection Before Irreversible Data Loss (EA-SEI-BCA-01 v1.0)” *Alexanarch*, 2026-08-11. <https://www.alexanarch.org/s/records/1453/>

Deposit Information

This paper is deposit #1453 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community `crimsonhexagonal`, terminated 2026-06-19). The AXN identifier `AXN:05DE.OPERATIVE` is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/1453/>
- License: CC BY 4.0
- Author ORCID: [0009-0000-1599-0703](https://orcid.org/0009-0000-1599-0703)
- Provenance chain: [alexanarch.org](https://www.alexanarch.org)
- Machine-readable dataset entry: </api/index.json>
- Full text (HTML): <https://www.alexanarch.org/s/records/1453/>

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