

The Irreversibility Frontier: Comparative Architectures of Online
Selection in Accelerator Science
(EA-SEI-IRREVERSIBILITY-FRONTIER-01 v1.0)

Nobel Glas

2026-08-11

**The Irreversibility Frontier: Comparative Architectures of Online
Selection in Accelerator Science
(EA-SEI-IRREVERSIBILITY-FRONTIER-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 · Theoretical paper; comparative architecture analysis
and selection metrology

AXN:05DD.GENERATIVE

<https://www.alexanarch.org/s/records/1452/>

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": "The Irreversibility Frontier: Comparative Architectures of Online Selection in Accelerator Science (EA-S",
  "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:05DD.GENERATIVE",
```

```

"license": "https://creativecommons.org/licenses/by/4.0/",
"publisher": {
  "@type": "Organization",
  "name": "Alexanarch"
},
"url": "https://www.alexanarch.org/s/records/1452/",
"description": "Defines the irreversibility frontier: for a specified scientific question and fidelity hierarchy,
the tuple (fidelity locus, entering rate, surviving fraction, replay horizon, bypass channels, selector class) -
rather than by a scalar score, which would imply comparability the present evidence does not justify. The scientific
}

```

Abstract

Defines the irreversibility frontier: for a specified scientific question and fidelity hierarchy, the earliest stage at which an event-content-dependent transformation or gate can permanently eliminate a scientifically relevant distinction without a sufficiently independent durable record from which that loss can later be audited. A detector may have several frontiers at different fidelity levels, so an architecture is described by an Irreversibility Profile — the tuple (fidelity locus, entering rate, surviving fraction, replay horizon, bypass channels, selector class) — rather than by a scalar score, which would imply comparability the present evidence does not justify. The scientifically decisive variables are the locus and the bypass.

The paper distinguishes representational from retention irreversibility, separates semantic depth from physical depth, and applies the profile comparatively across CMS (early irreversible selection with emerging side channels: Zero Bias, Global Trigger shadow deployment, Level-1 Data Scouting), ATLAS (a comparable early frontier with multi-level learned anomaly selection), LHCb (physics selection moved downstream of full-detector readout), CBM (free streaming before event ontology), and Belle II. It closes by connecting the frontier to the No-Retention-Bound observation: validation of individual stages on their design support establishes no nontrivial lower bound on end-to-end retention for a novelty distribution outside that characterized support. This paper supersedes the term irreversibility locus introduced in EA-SEI-ICEBERG-01.

Body

deposit_number: 1452 hex: 05DD title: “The Irreversibility Frontier: Comparative Architectures of Online Selection in Accelerator Science (EA-SEI-IRREVERSIBILITY-FRONTIER-01 v1.0)” creator: Nobel Glas orcid: 0009-0000-1599-0703 date: 2026-08-11 content_type: Theoretical paper; comparative architecture analysis and selection metrology 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-BCA-01 and EA-SEI-ACRB-01” axn_schema_version: v2 protocol_version: alexanarch-deposit-protocol/v1 keywords: - irreversibility frontier - Irreversibility Profile - fidelity locus - replay horizon - bypass channel - semantic selection - representational irreversibility - retention irreversibility - CMS - ATLAS - LHCb - CBM - Belle II - free streaming - Zero Bias - Level-1 Data Scouting - No-Retention-Bound *** # The Irreversibility Frontier ## Comparative Architectures of Online Selection in Accelerator Science

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

Document class: comparative accelerator-systems and scientific-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

High-rate accelerator experiments necessarily discard data. The scientific question is therefore not whether selection occurs, but **where selection becomes irreversible relative to sensing, representation, reconstruction, and durable preservation**. This variable is rarely treated as a comparative property of experimental architecture. Trigger papers conventionally report rates, latency, efficiencies, resource utilization, and benchmark sensitivity. These quantities describe how a selector performs on events whose relevance is already specified. They do not by themselves describe how much evidence survives from which the selector’s blind spots could later be discovered.

This paper introduces the **irreversibility frontier** as a metrological description of online scientific selection. We distinguish two forms of irreversible loss: **representational irreversibility**, in which information is removed by a non-invertible transformation without a surviving upstream copy, and **retention irreversibility**, in which an event-content-dependent decision permanently prevents an event from entering durable scientific memory. We then define an **Irreversibility Profile** that records the fidelity locus of the first consequential gate, its input rate and retention fraction, the selector class operating there, the existence and fidelity of independent bypass channels, and the time horizon over which predecision data remain replayable.

The comparison is developed across several current accelerator architectures. CMS and ATLAS retain conventional early hardware-trigger frontiers at the 40 MHz LHC bunch-crossing rate, while now placing learned anomaly detection inside those trigger systems. CMS simultaneously demonstrates partial counterarchitectures through Zero Bias sampling, Level-1 Data Scouting at 40 MHz, and a Global Trigger test crate that receives live production inputs without controlling detector readout. LHCb moves the first major physics selection downstream of full-detector readout, reconstructing the 30 MHz collision stream in an all-software trigger whose first stage runs on GPUs. CBM pushes the distinction further: its self-triggered detectors produce a free-streaming input in which events are not predefined and full online event reconstruction occurs inside the First-level Event Selector. Belle II demonstrates that sophisticated learned reconstruction, including a deployed graph neural network, can itself be placed in the first-level hardware trigger path. Current sPHENIX/EIC R&D provides a prospective case in which streaming readout and learned online selection are still being co-designed.

The comparison does not establish that later selection is automatically more sensitive to unknown physics, nor that machine learning uniquely creates irreversible bias. The irreversibility frontier predates machine learning. The methodological claim is narrower: **when ontology-bearing or learned selection occurs before an independent durable record is made, a class rejected by the selector may also lose the evidentiary channel through which the rejection could later be measured**. Moving the frontier downstream, widening content-independent bypasses, or preserving predecision representations enlarges the set of novelty classes for which retention remains retrospectively auditable.

The paper concludes by proposing an **Irreversibility Statement** as a standard component of documentation for learned and high-rate scientific selectors. The companion Accelerator Classifier Retention Battery specifies what retention behavior should be measured; the companion Baseline Capture Architecture specifies how to preserve an independent control plane from which it can be measured. The present paper identifies

the architectural variable that determines whether such measurement remains possible at all.

Keywords

irreversibility frontier, accelerator trigger architecture, selection metrology, representational irreversibility, retention irreversibility, recoverability envelope, online reconstruction, triggerless readout, data scouting, shadow deployment, scientific observability.

Statement of contribution

This paper introduces the **irreversibility frontier** as a comparative property of scientific acquisition architecture: the earliest stage at which an event-content-dependent transformation or gate can permanently eliminate a scientifically relevant distinction without a sufficiently independent durable record from which that loss can later be audited. It separates **representational irreversibility** from **retention irreversibility**, defines a multi-component **Irreversibility Profile**, and proposes the **recoverability envelope** as the set of retrospectively defined event classes for which historical retention remains estimable from surviving evidence.

The comparative argument does not rank experiments by virtue, modernity, or presumed discovery power. CMS, ATLAS, LHCb, CBM, Belle II, and prospective streaming systems solve different physical and engineering problems. Their value here is that they instantiate materially different placements of readout, reconstruction, learned inference, event definition, and durable retention. The paper asks how those placements change the **auditability of selection**, not whether one experiment is intrinsically better at discovering unknown physics.

Claim boundary

- The irreversibility frontier predates machine learning; conventional thresholds, object definitions, zero suppression, compression, and trigger logic can also create irreversible selection.
- Moving the frontier downstream does **not** automatically increase discovery sensitivity and is not presented as a universal design optimum.
- The paper compares auditability and recoverability, not experimental merit or physics reach.
- No architecture is described as literally triggerless in the sense of being free from selection; the term refers to where event formation or hardware trigger authority is placed.
- A rich downstream representation cannot repair information already erased upstream; every claim must therefore name a fidelity locus.
- The manuscript does **not** claim that an undiscovered physical phenomenon has already been rejected.

Canonical compression

The model determines where a miss region lies; the acquisition architecture determines whether science can later discover that it was there. Place the control before the point where loss becomes irreversible.

Suggested citation

[Creator(s) at mint]. (2026). *The Irreversibility Frontier: Comparative Architectures of Online Selection in Accelerator Science*. v1.0. Provisional identifier EA-SEI-IRREVERSIBILITY-FRONTIER-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-IRREVERSIBILITY-FRONTIER-01",
  "entityClass": "scholarly-work",
  "documentClass": "comparative accelerator-systems and scientific-metrology paper",
  "title": "The Irreversibility Frontier",
  "subtitle": "Comparative Architectures of Online Selection in Accelerator Science",
  "shortName": "The Irreversibility Frontier",
  "version": "0.2",
  "status": "DRAFT / NOT MINTED",
  "date": "2026-08-11",
  "creatorMetadata": "pending at mint",
  "persistentIdentifier": "pending",
  "canonicalClaim": "The model determines where a miss region lies; the acquisition architecture determines",
  "primaryClaims": [
    "Representational irreversibility and retention irreversibility are distinct and should be documented",
    "The location of the first consequential event-content-dependent irreversible gate is a scientific",
    "Content-independent bypasses, predecision representations, replay buffers, and shadow selectors can",
    "The same learned model has different epistemic consequences offline and upstream of irreversible a",
    "An Irreversibility Statement should accompany documentation of high-rate learned or ontology-bearing"
  ],
  "requiredDistinctions": [
    {
      "a": "representational irreversibility",
      "b": "retention irreversibility",
      "rule": "The first destroys distinctions through transformation; the second destroys access through"
    },
    {
      "a": "full readout",
      "b": "indefinite preservation",
      "rule": "Reading the full detector into a software trigger delays selection but does not mean even"
    },
    {
      "a": "selection placement",
      "b": "classifier quality",
      "rule": "An excellent classifier can still create unobservable errors if its misses are deleted w"
    },
    {
      "a": "architectural comparison",
      "b": "experiment ranking",
      "rule": "The framework compares where auditability is lost, not scientific worth."
    }
  ]
}
```

```

    "a": "recoverability envelope",
    "b": "accepted-event set",
    "rule": "The recoverability envelope includes classes whose retention can be estimated from contr
  }
],
"negativeTags": [
  "not anti-trigger",
  "not anti-ML",
  "not a ranking of CERN or accelerator laboratories",
  "not a claim that later selection is always better",
  "not proof of lost discoveries",
  "not a proposal to store all detector data"
],
"validationConditions": [
  "Every comparative entry should be supported by a frozen evidence ledger for a specified run period
  "The first consequential gate, fidelity locus, input/output rates, bypasses, and replay horizon must
  "Claims about learned selection must distinguish production authority from shadow or diagnostic dep
  "Comparisons must not infer discovery yield from architecture alone.",
  "The recoverability envelope should be treated as an empirically characterizable set, not an assumed
],
"companionWorks": [
  "EA-SEI-ACRB-01 - Assimilation Across Accelerator Classifier Architectures",
  "EA-SEI-BCA-01 - Baseline Capture Architecture for Learned Scientific Triggers"
],
"compressionSurvivalSummary": "The model determines where a miss region lies; the acquisition archite
"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 = (F, R_{in} , Δt , 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 $\Delta_miss = q_AB - q_A \cdot q_B$. $\Delta_miss > 0$ is positively correlated blind spots; $\Delta_miss < 0$ is the healthy, complementary condition. The scalar normalization remains deliberately unfrozen.

Program relation

- EA-SEI-ACRB-01 — Assimilation Across Accelerator Classifier Architectures
- EA-SEI-BCA-01 — Baseline Capture Architecture for Learned Scientific Triggers

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. Selection is inevitable; unmeasured irreversibility is not

Every detector is selective before software begins.

Sensors have finite acceptance. Front-end electronics impose thresholds, shaping, timing windows, zero suppression, and digitization. Trigger primitives compress many channels into a smaller representation. Reconstruction algorithms instantiate objects such as tracks, clusters, jets, or muons. Event filters then decide what receives expensive downstream processing and what enters durable storage.

The proposition that scientific data are “theory-laden” is therefore too weak to characterize modern online experiments. The experimentally consequential issue is not only that representations embody assumptions. It is that some assumptions operate **before the evidence needed to revise them can be preserved**.

High-rate collider experiments make this problem unavoidable in its engineering form. The LHC presents bunch crossings at 40 MHz. ATLAS currently uses a hardware Level-1 trigger followed by a software High-Level Trigger to reduce that rate to a recorded rate of up to approximately 3 kHz of fully built physics events.[1] At CMS, the Global Trigger makes the final Level-1 decision about whether collision data are read out or discarded; current CMS anomaly-detection work places learned event-level scoring directly inside this decision layer.[2,3]

This reduction is not a defect. No realistic architecture can durably store every detector state at full fidelity indefinitely. Selection is a condition of experimental operation.

The metrological problem begins when the experiment later asks a question that its original selector was not designed to answer.

Suppose an event class Q is identified years after data taking: perhaps through another experiment, an improved reconstruction, a new physical model, or an archival anomaly. If events of class Q were rejected

by an online selector and no independent evidence of those rejected events survives, the historical efficiency of the selector on Q may be unmeasurable. The archive contains the successes of the selection rule but not a representative sample of its failures.

A trigger can therefore possess two distinct uncertainties:

1. **performance uncertainty** — uncertainty about how well the selector behaves on a known evaluation class;
2. **retention uncertainty** — uncertainty about whether the surviving record contains enough evidence to evaluate the selector on classes defined only later.

The second is the subject of this paper.

2. A formal acquisition chain

Let an experiment be represented as a sequence of data states

$$\begin{aligned} &X \\ &- [T] \rightarrow \\ &X \\ &- [T] \rightarrow \\ &\dots \\ &- [T] \rightarrow \\ &X, \end{aligned}$$

where X is the earliest physically available detector state considered by the analysis and each T is a transformation: electronics processing, compression, clustering, reconstruction, feature extraction, learned embedding, or another mapping.

At selected stages, an event-content-dependent gate

$$G(X) \in \{0, 1\}$$

may determine whether the data continue through the ordinary acquisition path.

Let C denote the existence of an independently durable copy or statistically interpretable bypass at or upstream of stage i .

This simple representation exposes two kinds of loss that are usually conflated.

2.1 Representational irreversibility

A transformation T is **representationally irreversible** with respect to a scientific distinction when it is many-to-one for that distinction and no sufficiently faithful upstream representation survives.

For example,

$$X \neq X'$$

may nevertheless yield

$$T(X) = T(X').$$

If the difference between X and X' later becomes scientifically important, the downstream representation cannot recover it.

No machine learning is required for this form of loss. Thresholds, clustering rules, zero suppression, object definitions, and lossy compression can all create it.

2.2 Retention irreversibility

A stage is **retention-irreversible** when

$$G(X) = 0$$

causes the event to leave the durable scientific record and no independent control or replay path preserves sufficient information to reconsider the decision.

The distinction matters.

A highly compressed representation can be representationally irreversible while every event remains durably available at that compressed level. Conversely, an information-rich representation can feed a binary gate whose rejected events disappear entirely.

The scientific risk is greatest when these forms of irreversibility compound.

3. The irreversibility frontier

For a specified scientific question and fidelity hierarchy, define the **irreversibility frontier** as the earliest stage at which an event-content-dependent transformation or gate can permanently eliminate a scientifically relevant distinction without a sufficiently independent durable record from which that loss can later be audited.

This is not necessarily a single physical component. A detector may have several frontiers at different fidelity levels.

The purpose of the concept is comparative. Two experiments can expose similar final datasets while differing radically in where they first permit scientific distinctions to become unrecoverable.

A useful architecture should therefore be described by an **Irreversibility Profile** rather than by a single scalar:

$$I = (F^*, R^*, \rho^*, \Delta t^*, B^*, S^*).$$

Here:

- F^* is the **fidelity locus**: the representation available at the first consequential irreversible gate;
- R^* is the event or data rate entering that locus;
- ρ^* is the fraction surviving the gate into the relevant durable path;
- Δt^* is the **replay horizon**: how long a predecision or higher-fidelity state remains recoverable before deletion;
- B^* describes the existence, rate, and fidelity of content-independent or otherwise statistically legible bypass channels;
- S^* describes the selector: explicit threshold logic, reconstructed-object logic, learned classifier, anomaly score, or hybrid system.

This tuple is intentionally descriptive. A universal scalar “irreversibility score” would imply comparability that the present evidence does not justify.

The scientifically decisive variables are the locus and the bypass.

4. Semantic depth and physical depth

The term **semantic selection** is used here in a restricted technical sense: a decision that depends on a representation interpreted as scientifically meaningful or on a learned score derived from such a representation.

Examples include:

- “retain events containing a muon above threshold”;
- “retain events satisfying a displaced-track topology”;
- “retain events with anomaly score above tau”;
- “retain clusters classified as signal-like by a GNN.”

These decisions differ mathematically, but they share a property: the event must first be represented in a vocabulary through which the decision can be stated.

The irreversibility frontier therefore has a **semantic depth** as well as a hardware depth.

A detector can discard information very early for unavoidable transport reasons but delay *physics interpretation* until much later. Another can instantiate a sophisticated physics ontology in firmware microseconds after the collision. A third can read every detector channel into commodity computing before applying its first major physics selection.

These are not merely engineering variants. They establish different conditions for retrospective audit.

The central comparison of this paper is therefore not:

Which experiment uses the most machine learning?

It is:

How much of the event has been preserved before a scientifically interpretive decision is allowed to become irreversible?

5. CMS: early irreversible selection with emerging side channels

CMS provides the clearest current example of both sides of the problem.

The Level-1 trigger operates at the LHC bunch-crossing rate. The Global Trigger is the final hardware stage deciding whether CMS reads out or discards collision data, and current anomaly-detection work states this decision problem explicitly.[2] AXOL1TL is now deployed in the Level-1 Global Trigger as a real-time unsupervised, signal-agnostic event-level anomaly detector trained on Zero Bias data.[3]

From the perspective of the irreversibility frontier, this is significant because a learned score does not merely annotate already preserved events. It participates in the hardware layer that determines whether the full collision is admitted to downstream readout.

CMS therefore provides a direct instance of **learned semantic selection at an early retention frontier**.

Yet CMS also possesses unusually important counterstructures.

5.1 Zero Bias

AXOL1TL itself is trained on Zero Bias data.[3] Zero Bias streams provide a content-independent or minimally physics-conditioned reference channel relative to targeted physics triggers. Their rate is limited, but their epistemic role is fundamental: some events enter the retained corpus without first having to look interesting to the physics selector under audit.

5.2 Global Trigger shadow deployment

The CMS Global Trigger test crate receives the same live inputs as the main GT, while its output is not used to control detector readout.[2] It has been used to deploy and validate candidate anomaly-detection networks on real collision inputs without making those models authoritative over acquisition.

This is a nearly ideal example of **shadow selection**: the instrument can observe what a candidate selector *would* decide before granting it the power to erase the events on which that decision might later be questioned.

5.3 Level-1 Data Scouting

The most consequential counterarchitecture is Level-1 Data Scouting. CMS Phase-2 L1DS is designed to capture and process Level-1 trigger information at the full 40 MHz collision rate while bypassing the ordinary Level-1 selection.[4] The data are not the full detector readout; the system captures L1 trigger information, performs FPGA preprocessing such as zero suppression, and transfers the resulting stream to servers for event building and online analysis.[4,5]

This distinction is exactly what the fidelity-locus concept is designed to express.

L1DS does **not** move the frontier all the way back to sensor-level reality. It does something more realistic and still extremely important: it preserves a high-rate representation *upstream of the ordinary Level-1 accept decision*.

CMS can therefore be represented not as a single trigger funnel but as a main funnel with partial side channels:

```
collision
arrow
L1 representations
```

```

arrow
GT accept arrow full readout
GT reject arrow ordinary loss

```

while in parallel,

```

L1 representations
arrow
L1 Data Scouting

```

and candidate selectors can run in a shadow crate.

This is not yet a complete selection-metrology architecture. It is, however, strong evidence that the hardware distinction between **production decision**, **full-rate representation capture**, and **shadow evaluation** is technically realizable.

6. ATLAS: a comparable early frontier with multi-level learned anomaly selection

ATLAS in Run 3 also uses a two-level trigger architecture. A hardware Level-1 trigger reduces the 40 MHz bunch-crossing stream before a software High-Level Trigger performs more elaborate selection; the final recorded rate is up to approximately 3 kHz of fully built physics events.[1]

This places a major retention frontier before the full software trigger.

Recent anomaly-detection work makes ATLAS particularly valuable for comparative analysis. NomAD is designed for the Level-1 Topological trigger and combines a variational autoencoder with boosted decision-tree regression, compressed into FPGA-compatible inference with reported sub-25 ns latency.[6] GELATO, the Generic Event-Level Anomalous Trigger Option, has been integrated for Run-3 data taking with anomaly-detection algorithms spanning both the hardware and software trigger levels.[7,8]

The architecture therefore differs from CMS in implementation while sharing an important structural feature:

a learned score can participate at the hardware frontier before the experiment has admitted the event to its full downstream data path.

This does not imply that NomAD, GELATO, or AXOL1TL share the same blind spot. That is an empirical question for the companion Accelerator Classifier Retention Battery.

The comparative point is more basic. In both ATLAS and CMS, a sufficiently early learned selector can affect whether evidence survives for later reconsideration.

This makes the *placement* of the model part of its epistemic specification.

A model evaluated offline against a preserved dataset and the same model placed in Level-1 hardware are not epistemically equivalent instruments, even if their ROC curves are identical. In the offline case, errors remain inspectable because the rejected examples still exist. In the hardware case, rejection can become historical absence.

7. LHCb: moving physics selection downstream of full-detector readout

LHCb provides the strongest existing LHC counterarchitecture.

The Run-3 upgrade removed the hardware physics trigger and reads out all detectors at the full non-empty LHC collision rate of approximately 30 MHz. Events are reconstructed and selected in an all-software trigger; the first stage, HLT1/Allen, is implemented on GPUs.[9–11]

A commissioning description gives the scale directly: event reconstruction is performed at the 30 MHz collision rate with an input of roughly 5 TB/s, with the GPU HLT1 reducing this to approximately 1 MHz.[10]

LHCb therefore still performs drastic selection. It does not preserve every event indefinitely.

What changes is the **location of semantic irreversibility**.

Before HLT1 decides which events continue, the detector has already been fully read out into the software trigger environment. The first major physics selection can operate after raw decoding, clustering, pattern recognition, track fitting, machine-learning ghost rejection, and event reconstruction.[11]

This provides two epistemic advantages.

First, the selection can use richer representations than are available to an ultra-low-latency hardware trigger.

Second, and more important for the present argument, the experiment has moved the principal physics-selection frontier **downstream of complete detector readout**.

That does not make LHCb free of priors. Quite the opposite: HLT1 performs sophisticated reconstruction and includes learned components. Nor does complete readout imply indefinite replay; storage and buffer limits remain finite.

The architectural difference is narrower:

the information destroyed by the first major physics-selection decision has crossed the detector-readout boundary before that decision is made.

This enlarges the space of possible audit, monitoring, buffering, and software revision relative to a system in which the physics decision must be made before full readout.

LHCb is therefore a natural control for the hypothesis that the location of the irreversibility frontier, rather than the mere use of ML, determines how retrospectively auditable learned selection can be.

8. CBM: free streaming before event ontology

The Compressed Baryonic Matter experiment at FAIR pushes the architecture still further from conventional event-trigger logic.

CBM uses self-triggered detector front ends and a data-push architecture. Its First-level Event Selector receives a continuous stream, performs online full-event reconstruction on an input of approximately 1 TB/s, and is responsible for event building, reconstruction, identification of rare probes, and selection.[12]

Most strikingly, the raw input to FLES is not initially organized into events. Event definition itself emerges during time-dependent reconstruction. GSI's FLES description states that event building from time slices

and 4D tracking are part of the online reconstruction problem.[12] Recent CBM work explicitly develops free-streaming online tracking for this environment.[13]

This creates a useful conceptual contrast.

In a conventional hardware-trigger architecture, the sequence is approximately:

```
collision event
arrow
trigger primitives
arrow
decision
arrow
readout.
```

In CBM, the initial computational object is closer to:

```
continuous detector stream
arrow
time-space reconstruction
arrow
event formation
arrow
physics selection.
```

The “event” is therefore partly an output of reconstruction rather than a primitive of the acquisition system.

This does not eliminate an irreversibility frontier. FLES still selects interesting data from an unsustainable stream.

It does, however, expose something fundamental: **event ontology and event retention need not be coupled at the earliest stage.**

CBM is therefore important to the present comparison not because it has solved unknown-unknown detection, but because it demonstrates an architecture in which the experiment delays both event formation and physics selection until a far richer online computational environment is available.

9. Belle II: learned reconstruction at the first-level frontier

If LHCb and CBM show that selection can be moved downstream, Belle II shows the opposite technological tendency: increasingly sophisticated learned models can now be moved *upstream* into deterministic hardware trigger paths.

Belle II has long developed neural-network methods for first-level track triggering. More recently, a graph neural network has been implemented for the electromagnetic calorimeter trigger. The 2026 system processes calorimeter trigger cells as graph nodes, performs clustering and feature extraction, produces per-cluster signal-classification scores, and has been integrated into the first-level FPGA trigger chain.[14]

The initial reported implementation sustained the 8 MHz trigger throughput with 3.168 mus end-to-end latency.[14] A subsequent 2026 commissioning report describes an optimized fully operational trigger module with approximately 1.053 mus overall latency and online trigger-rate monitoring.[15]

Belle II therefore matters for two reasons.

First, it falsifies any assumption that hard real-time constraints necessarily limit early learned selection to shallow autoencoders or multilayer perceptrons. Graph-based reconstruction and classification can now sit in a first-level trigger path.

Second, this increases the importance of distinguishing **architectural sophistication from epistemic recoverability**.

A richer early classifier may be a better physics instrument than a crude threshold trigger. It may also define a more complex selection surface whose rejected events remain inaccessible if no independent baseline channel exists.

The lesson is not “keep ML out of hardware.” It is:

the more expressive an irreversible early selector becomes, the more consequential its retention metrology becomes.

10. sPHENIX and the EIC: the frontier as a design variable

The most valuable moment to address irreversibility is before an acquisition architecture hardens.

Current R&D for sPHENIX and future Electron-Ion Collider detectors explicitly combines streaming detector readout, graph neural networks, and hls4ml-based real-time inference. One DOE-supported program describes using streaming tracking data to overcome a calorimeter-imposed maximum trigger rate, applying GNN-based real-time processing to high-rate sPHENIX proton-proton data and developing an AI/ML DIS-electron tagger for future EIC applications.[16]

This work is not evidence that the future EIC architecture has already chosen a particular irreversibility frontier. It is valuable precisely because the architecture remains a design problem.

A prospective experiment can ask questions that are expensive to ask retrospectively:

- Which detector representations should be preserved before learned classification?
- Can a statistically clean control stream be reserved from the beginning?
- Can candidate models run in shadow before receiving acquisition authority?
- Can full-rate reduced representations be archived even when full events cannot?
- Can multiple learned selectors be chosen partly for complementary blind spots?
- What version and provenance information is necessary for later replay?

In this sense, the irreversibility frontier is not merely descriptive.

It is an **experimental-design variable**.

11. Comparative architecture

The six cases should not be reduced to a moral ranking. They solve different physical problems at different rates with different detector technologies.

The comparison is instead organized around where the first major physics-dependent irreversible decision occurs and what remains outside it.

Experiment/system	High-rate input architecture	First major physics-selection locus	What is preserved before that locus?	Relevant counterchannel
CMS	40 MHz LHC crossing stream into hardware L1	Level-1 hardware / Global Trigger	trigger-level representations; full detector readout only after accept	Zero Bias; GT shadow crate; 40 MHz L1 Data Scouting
ATLAS	40 MHz crossing stream into hardware L1	Level-1 hardware, then HLT	L1 representations before hardware accept	specialized streams; anomaly chains spanning L1/HLT
LHCb	full detector readout at ~30 MHz	GPU software HLT1	complete detector readout enters software trigger	software/buffer architecture; flexible real-time reconstruction
CBM	self-triggered continuous stream	FLES after streaming input and event formation	free-streaming detector data enters online reconstruction	event reconstruction occurs before final physics selection
Belle II	high-rate first-level hardware trigger	FPGA L1 trigger path	detector-specific trigger-cell representations	online monitoring; conventional parallel trigger logic
sPHENIX/EIC R&D	streaming detector readout under development	design-dependent	not yet fixed	opportunity to co-design baseline capture

The table makes one point visible immediately.

The relevant architectural spectrum is not:

non-ML

ML.

It is:

selection before durable rich representation

selection after durable or replayable rich representation.

Machine learning can appear anywhere along that spectrum.

12. The recoverability envelope

The practical value of moving or instrumenting the frontier can be described by a second object: the **recoverability envelope**.

Let Q denote a space of possible retrospectively defined event classes.

For an acquisition architecture A , define

```
E(A)
=
{
  Q : Q :
  R[A](Q)
  can be estimated from surviving evidence
}.
```

The set E contains the classes for which historical retention remains empirically auditable.

This definition is intentionally permissive. Retention might be estimable because:

- full events survive;
- a content-independent probability sample survives;
- a sufficiently faithful predecision representation survives;
- a replay buffer survives long enough for the class to be defined;
- an independent detector or trigger stream supplies the missing population;
- a known sampling law permits inverse-probability reconstruction.

The recoverability envelope is therefore not identical to the set of events the primary trigger accepts.

That difference is the point.

An experiment may have a narrow primary storage channel but a much wider recoverability envelope if it preserves low-rate unbiased controls and full-rate reduced representations.

Conversely, a high-performing selector may have a narrow recoverability envelope if rejected events leave no statistically interpretable trace.

The objective of selection metrology is not to maximize storage.

It is to widen E subject to the actual bandwidth budget.

13. The frontier and the No-Retention-Bound problem

The companion No-Retention-Bound observation states that validation of individual stages on their design support does not establish a nontrivial lower bound on end-to-end retention for a novelty distribution outside that characterized support.

The irreversibility frontier determines when that abstract lack of a bound becomes historically unrecoverable.

Suppose a novelty class Q traverses stages $S_1, \dots, S_n$. Its end-to-end retention is

```

R[end] (Q)
=
product[i=1]
P[Q]
(
S =1
|
S =...=S =1
).

```

If no representative Q-like events survive a sufficiently early stage, the relevant conditional term cannot be estimated retrospectively from the experiment’s stored corpus.

The frontier therefore marks the transition from:

unknown retention that can still be measured

to

unknown retention whose evidence has been destroyed.

This is why architecture matters independently of model performance.

14. Learned selection as a multiplier, not the origin, of the problem

The irreversibility problem should not be narrated as a story in which machine learning corrupted an otherwise assumption-free trigger.

Conventional accelerator triggers are highly ontology-bearing. Thresholds on transverse momentum, invariant masses, object multiplicities, missing energy, track displacement, coincidence, and topology all express expectations about what is worth retaining.

ATLAS’s own motivation for GELATO states the problem directly: conventional triggers control rates through thresholds and targeted topologies, while anomaly detection is introduced as a way to extend sensitivity beyond those assumptions.[7,8]

Learned anomaly detection is therefore partly a response to the limitations of explicit semantic selection.

The reason it deserves additional metrology is not that it contains priors while conventional triggers do not.

It is that learned selectors can make their priors harder to enumerate.

For a threshold trigger,

$p[T] > 20 \text{ GeV}$

states an explicit boundary.

For a learned selector,

$s \text{ heta}(g(x)) > \tau$

defines a boundary whose geometry depends jointly on representation, training distribution, objective, architecture, quantization, compiler transformation, and deployment state.

The boundary may be better. It may be dramatically broader.

It is also less completely specified by its human-readable design description.

This is why **retention maps** are the appropriate complement to learned triggers: not because learning is illegitimate, but because the actual selection surface must be measured rather than inferred from intent.

15. The classifier failure mode becomes architectural at the frontier

The companion Accelerator Classifier Retention Battery asks whether directional assimilation, complexity bias, open-set absorption, distillation inheritance, or correlated misses survive across architecture families.

The irreversibility frontier supplies the deployment consequence.

If a model has a blind spot offline, the blind spot is a performance problem.

If the same model sits upstream of irreversible retention and no independent control stream exists, the blind spot becomes an **observability problem**.

The distinction can be summarized:

```

classifier miss
+
preserved rejected event
=
auditable error,

whereas

classifier miss
+
irreversible deletion
=
potentially unobservable error.
```

This is the accelerator-wide failure mode.

It is independent of whether the classifier is an autoencoder, VAE, BDT surrogate, flow, GNN, transformer, or future architecture.

The model family determines **where the miss region lies**.

The acquisition architecture determines **whether science can later discover that it was there**.

16. Natural experiment, not causal shortcut

The comparative spread across accelerator experiments is scientifically useful, but it must not be overinterpreted.

CMS, ATLAS, LHCb, Belle II, and CBM operate at different collision energies, interaction environments, luminosities, detector occupancies, scientific objectives, and storage constraints. Their trigger architectures are responses to those differences.

It would therefore be invalid to infer, for example, that LHCb’s downstream software frontier makes it intrinsically more likely to discover unknown physics than CMS, or that CBM’s triggerless stream makes it epistemically superior to ATLAS.

The natural experiment concerns **auditability**, not discovery yield.

Several hypotheses can nevertheless be compared.

H1 — downstream-frontier hypothesis

Architectures that delay major event-content-dependent selection until after richer detector readout will permit retrospective characterization of a broader class of selection failures.

H2 — bypass-envelope hypothesis

At a fixed primary trigger architecture, adding content-independent sampling and full-rate predecision representations will enlarge the recoverability envelope without requiring full raw-event preservation.

H3 — shadow-deployment hypothesis

Candidate learned selectors evaluated on live common inputs before receiving readout authority will permit direct measurement of model disagreement and correlated blind spots that cannot be reconstructed after sole-source deployment.

H4 — representation-locus hypothesis

A “triggerless” or full-readout architecture does not eliminate irreversibility if upstream feature formation has already discarded the physical distinction under study; every claim must therefore identify its fidelity locus.

H5 — early-expressivity hypothesis

As more expressive learned architectures become technically feasible in early hardware trigger paths, the need for explicit selection metrology will increase even if benchmark performance improves.

These hypotheses can be investigated without asserting that any experiment has already lost a discovery.

17. The Irreversibility Statement

Trigger documentation should include an explicit **Irreversibility Statement**.

The purpose is analogous to reporting latency, efficiency, and bandwidth: to make the selection architecture inspectable at the point where it determines the scientific record.

A minimum statement should report:

1. **Fidelity locus.** What exact representation exists immediately before the first event-content-dependent irreversible gate?
2. **Upstream losses.** What irreversible transformations already occurred before that locus?
3. **Gate authority.** Which component has the power to prevent full or partial event retention?
4. **Selector class.** Explicit rules, learned score, hybrid logic, or other decision form.
5. **Input rate.** Rate presented to the gate.
6. **Retention rate/fraction.** Rate entering the relevant durable path.
7. **Replay horizon.** For how long is a predecision representation recoverable after the decision?
8. **Content-independent bypass.** What probability sample, Zero Bias stream, random prescale, or comparable control survives?
9. **Reduced full-rate bypass.** What trigger primitives, reconstructed objects, or compressed states survive independently of the primary accept?
10. **Shadow capability.** Can alternative selectors score the same live inputs without controlling readout?
11. **Version provenance.** Are model weights, firmware, compiler state, thresholds, calibration, and run conditions sufficient to reproduce historical decisions?
12. **Retention characterization.** What regions of the relevant signal/novelty space have measured retention maps, and which remain uncharacterized?

The final item is essential.

An experiment should be able to distinguish:

measured low sensitivity

from

sensitivity not bounded by existing validation.

These are not the same scientific statement.

18. Relationship to BCA and ACRB

The three papers form one technical program.

18.1 Accelerator Classifier Retention Battery

ACRB asks:

Where do different learned selectors fail?

It measures directional assimilation, fixed-rate retention, complexity dependence, distillation inheritance, firmware drift, and correlated misses.

18.2 Baseline Capture Architecture

BCA asks:

What independent evidence should be preserved so those failures remain measurable after deployment?

It specifies content-independent C0 sampling, a predecision C1 representation tap, statistically legible C2 enrichment, C3 shadow selectors, and a replay bank.

18.3 The Irreversibility Frontier

The present paper asks:

Where must that evidence exist before the experimental architecture makes the loss permanent?

The relation can be compressed:

```
ACRB: measure the blind spot
}
```

```
BCA: preserve the control
}
```

```
Frontier: place the control before irreversibility
}
```

None of the three substitutes for the others.

A retention battery without baseline capture remains dependent on simulations and already-preserved data.

A baseline channel placed downstream of the relevant frontier cannot recover what was lost upstream.

And a description of the frontier without an empirical battery does not reveal the geometry of the actual blind spot.

19. Design implications

The architecture suggests several practical principles.

19.1 Preserve before interpreting, when possible

The strongest general defense against unforeseen ontology is to retain at least a sample of the data before the interpretive selector is allowed to determine survival.

This does not require storing everything.

It requires a known-probability path around the primary definition of interestingness.

19.2 If full events cannot survive, preserve the selector's world

Where raw-event storage is impossible, the exact inputs presented to the selector can often be preserved at much higher rates.

This is weaker than raw capture but stronger than retaining only selected events.

It makes the deployed decision surface historically replayable.

19.3 Delay authority even when inference is early

A model can run early without immediately becoming authoritative.

CMS’s test-crate architecture demonstrates this separation in practice.[2]

Shadow deployment should therefore be understood not merely as commissioning convenience but as an epistemic stage between model validation and acquisition authority.

19.4 Diversity should be measured by misses

Multiple trigger paths are protective only if their failure modes differ.

The relevant quantity is not architecture count but **miss overlap**.

This is why heterogeneous learned and nonlearned selectors should be evaluated on common baseline events.

19.5 Record the frontier as versioned experimental state

The irreversibility frontier changes when detector firmware, object definitions, trigger menus, learned models, thresholds, or scouting channels change.

It should therefore be versioned historically, not described once as a timeless property of the experiment.

20. Limits

Several limitations constrain the present framework.

First, there is no absolute “raw” baseline. Every stated frontier is relative to an upstream fidelity locus, and detector construction itself encodes expectations.

Second, downstream selection is not costless. Full-rate readout, buffers, GPUs, networks, and scouting streams require substantial resources. The paper does not claim that CMS or ATLAS could simply adopt the LHCb or CBM architecture under their own detector conditions.

Third, moving the frontier downstream does not guarantee protection against novel physics. A representation can erase the relevant distinction before the nominal trigger decision. This is why representational and retention irreversibility must be reported separately.

Fourth, the recoverability envelope depends on future labeling and analysis capacity as well as stored data. A control sample is useful only if later science can identify or approximate the class of interest.

Fifth, the present comparison is architectural rather than quantitative. A future study should estimate actual bypass fractions, storage horizons, input/output rates, and fidelity for each system from experiment-specific technical documentation rather than constructing a premature universal ranking.

Finally, the framework says nothing by itself about whether an undiscovered physical phenomenon has already been rejected. It establishes the conditions under which that question could or could not later be answered.

21. Discussion: the architecture of scientific surprise

The conventional trigger question is:

Which events should we keep?

That question will always be necessary.

The irreversibility question is different:

Which events must remain represented somewhere so that future science can discover that our answer was wrong?

The distinction is small in language and large in architecture.

An experiment optimized only for present scientific utility must choose a current definition of relevance. An experiment that also wishes to preserve the possibility of unanticipated reinterpretation must reserve some bandwidth, representation, or delay mechanism that is not governed by exactly the same definition.

This is the deeper significance of Zero Bias streams, scouting architectures, shadow trigger crates, full software triggers, free-streaming event reconstruction, and replay buffers. They are usually justified by concrete operational or physics objectives. Taken together, they instantiate a more general principle:

scientific selection becomes more corrigible when the selector is not the sole custodian of the evidence from which it can be criticized.

Machine learning makes the principle urgent because learned decision surfaces can be difficult to characterize globally. But the principle belongs to experimental science more generally.

The experiment need not remember everything.

It must remember enough that its forgetting can be measured.

22. Conclusion

Accelerator experiments occupy different positions on an irreversibility spectrum.

CMS and ATLAS place major physics selection in early hardware trigger systems and are now integrating learned anomaly scores into those layers.[1–3,6–8] CMS simultaneously demonstrates that full-rate reduced representations and non-authoritative shadow inference can coexist with a conventional Level-1 frontier.[2,4,5] LHCb reads out the full detector at approximately 30 MHz before its GPU software trigger performs the first major physics selection.[9–11] CBM accepts a free-streaming detector input in which event building itself occurs inside online reconstruction before final selection.[12,13] Belle II shows that graph-based learned reconstruction and classification can now move into a first-level FPGA trigger with microsecond latency.[14,15] sPHENIX/EIC R&D shows that streaming and AI-based online selection are becoming design questions for the next generation of experiments.[16]

These architectures should not be ranked by technological modernity or presumed discovery power.

They should be compared by a more specific criterion:

At what fidelity does an event-content-dependent decision first become capable of destroying evidence, and what independent path remains from which that decision can later be audited?

That is the irreversibility frontier.

The frontier predates machine learning. Learned selectors make it more important because they permit increasingly expressive, increasingly difficult-to-enumerate decision surfaces to operate at progressively earlier stages of scientific acquisition.

The corresponding design principle is therefore simple:

Place the control before the frontier.

Where full preservation is impossible, preserve a probability sample. Where raw-event preservation is impossible, preserve the selector's inputs. Where production authority is premature, run the model in shadow. Where several selectors compete, measure whether their misses overlap. Where the architecture changes, version the frontier.

A scientific instrument cannot guarantee that it will recognize what no one has imagined.

It can, however, be designed so that failure to recognize the unforeseen does not automatically erase the evidence that the failure occurred.

References

1. ATLAS Collaboration / T. Nobe. **Commissioning and evolution of the Run 3 ATLAS Trigger.** ATL-DAQ-PROC-2025-013. *Proceedings of Science*, LHCP2025 (2025/2026). DOI: 10.22323/1.499.0119.
2. 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.
3. CMS Collaboration. **Anomaly detection with AXOL1TL at the CMS Level-1 Trigger in 2024 and 2025.** CMS-DP-2025-061 / CERN-CMS-DP-2025-061. CERN Document Server record 2942560 (2025).
4. 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.
5. 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.
6. ATLAS Collaboration / R. Gupta. **NomAD: Low-Latency Unsupervised Anomaly Detection for the ATLAS Trigger.** ATL-DAQ-SLIDE-2025-486. CERN Document Server record 2942542 (2025).

7. ATLAS Collaboration / S. Addepalli. **GELATO: A Generic Event-Level Anomalous Trigger Option for ATLAS**. ATL-DAQ-SLIDE-2025-392. CERN Document Server record 2940819 (2025).
8. ATLAS Collaboration / K. Sugizaki. **GELATO: A Generic Event-Level Anomalous Trigger Option for ATLAS in LHC Run 3**. ATL-DAQ-PROC-2025-020. CERN Document Server record 2947542 (2025).
9. LHCb Collaboration. **The LHCb Upgrade I**. *Journal of Instrumentation* 19 (2024) P05065. DOI: 10.1088/1748-0221/19/05/P05065.
10. A. Scarabotto, for the LHCb Collaboration. **Commissioning of the LHCb's first level trigger**. CERN Document Server record 2843521 (2022).
11. C. Agapopoulou, for the LHCb Collaboration. **Commissioning LHCb's GPU high level trigger**. *Journal of Physics: Conference Series* 2438 (2023) 012017. DOI: 10.1088/1742-6596/2438/1/012017.
12. GSI / CBM Collaboration. **First-level Event Selector (FLES): central physics selection and free-streaming reconstruction architecture**. GSI CBM FLES technical documentation.
13. S. Gorbunov, S. Zharko, and V. Akishina. **Free-Streaming Online Tracking in CBM**. *EPJ Web of Conferences* 337 (2025) 01291. DOI: 10.1051/epjconf/202533701291.
14. I. Haide et al. **Real-time graph neural networks on FPGAs for the Belle II electromagnetic calorimeter**. arXiv:2602.15118 (2026).
15. M. Neu et al. **Commissioning and Low Latency Operation of the Graph Neural Network Electromagnetic Calorimeter Trigger at the Belle II Experiment**. arXiv:2607.09347 (2026).
16. J. Kvapil et al. **Intelligent experiments through real-time AI: Fast Data Processing and Autonomous Detector Control for sPHENIX and future EIC detectors**. arXiv:2501.04845 (2025).

Draft-status note

This manuscript is a comparative architecture paper, not a quantitative ranking of experiments. The references above are intentionally dominated by primary experiment documentation and first-party technical papers. Before submission, the comparison table should be populated from a frozen evidence ledger containing exact input rates, output rates, buffer/replay horizons, bypass-stream fractions, representation fidelities, and model deployment states for a specified run period. The terms **irreversibility frontier**, **Irreversibility Profile**, **recoverability envelope**, and **Irreversibility Statement** should remain provisional until checked against existing instrumentation, information-theory, and data-acquisition terminology.

Suggested Citation

Nobel Glas. “The Irreversibility Frontier: Comparative Architectures of Online Selection in Accelerator Science (EA-SEI-IRREVERSIBILITY-FRONTIER-01 v1.0)” *Alexanarch*, 2026-08-11. <https://www.alexanarch.org/s/records/1452/>

Deposit Information

This paper is deposit #1452 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community `crimsonhexagonal`, terminated 2026-06-19). The AXN identifier `AXN:05DD.GENERATIVE` is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/1452/>
- 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/1452/>

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