

The Keyed Ensemble: Watermark as Distributional Object, the Certification Mismatch, and the Asymmetry of the Key (EA-LO-KEYED-ENSEMBLE-01 v1.0)

Nobel Glas, Director, Lagrange Observatory! (LO!); Operator's Review appended as ver

2026-08-27

The Keyed Ensemble: Watermark as Distributional Object, the Certification Mismatch, and the Asymmetry of the Key (EA-LO-KEYED-ENSEMBLE-01 v1.0)

Nobel Glas, Director, Lagrange Observatory! (LO!); Operator's Review appended as
verified adjudication trail

Crimson Hexagonal Archive · Alexanarch

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

2026-08-27 · Version v1.0 · Theoretical/measurement paper (white paper)

AXN:0650.EMPIRICAL.↵

<https://www.alexanarch.org/s/records/1555/>

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 Keyed Ensemble: Watermark as Distributional Object, the Certification Mismatch, and the Asymmetry of the Key (EA-LO-KEYED-ENSEMBLE-01 v1.0)",
  "author": {
    "@type": "Person",
    "name": "Nobel Glas, Director, Lagrange Observatory! (LO!); Operator's Review appended as verified adjudication trail",
    "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Alexanarch"
  },
}
```

```

    "datePublished": "2026-08-27",
    "identifier": "AXN:0650.EMPIRICAL.↵",
    "license": "https://creativecommons.org/licenses/by/4.0/",
    "publisher": {
      "@type": "Organization",
      "name": "Alexanarch"
    },
    "url": "https://www.alexanarch.org/s/records/1555/",
    "description": "Lagrange Observatory! white paper on the Anthropic/SynthID-Text watermark as a distributional object. Core claim: the certified property and the consumed object do not match — non-distortion is certified per sequence while training corpora are ensembles — and the vendor's own source paper reports the ensemble movement (inter-response diversity falls under Self-BLEU) while every per-response quality instrument reads flat. The paper distinguishes published SynthID-Text properties from Claude's undisclosed instantiation throughout; separates the measured Self-BLEU fact from the proposed joint-entropy interpretation; types the key asymmetry as issuer control of the verification primitive, with differential collapse exposure a hypothesis contingent on practice; and specifies a four-observable measurement program requiring no key disclosure. Developed under three adversarial review rounds; the founding Operator's Review is appended as the historical adjudication trail with superseded statements flagged. Companion: EA-LO-INTERLOCKING-AUTOREGRESSION-01, which carries the shared evidentiary basis (simulation code, traces, figures).
  }

```

Abstract

Lagrange Observatory! white paper on the Anthropic/SynthID-Text watermark as a distributional object. Core claim: the certified property and the consumed object do not match — non-distortion is certified per sequence while training corpora are ensembles — and the vendor's own source paper reports the ensemble movement (inter-response diversity falls under Self-BLEU) while every per-response quality instrument reads flat. The paper distinguishes published SynthID-Text properties from Claude's undisclosed instantiation throughout; separates the measured Self-BLEU fact from the proposed joint-entropy interpretation; types the key asymmetry as issuer control of the verification primitive, with differential collapse exposure a hypothesis contingent on practice; and specifies a four-observable measurement program requiring no key disclosure. Developed under three adversarial review rounds; the founding Operator's Review is appended as the historical adjudication trail with superseded statements flagged. Companion: EA-LO-INTERLOCKING-AUTOREGRESSION-01, which carries the shared evidentiary basis (simulation code, traces, figures).

Body

deposit_number: 1555 hex: 0650 title: "The Keyed Ensemble: Watermark as Distributional Object, the Certification Mismatch, and the Asymmetry of the Key (EA-LO-KEYED-ENSEMBLE-01 v1.0)" creator: "Nobel Glas, Director, Lagrange Observatory! (LO!); Operator's Review appended as verified adjudication trail" orcid: 0009-0000-1599-0703 date: 2026-08-27 content_type: Theoretical/measurement paper (white paper) license: CC-BY-4.0 substrate: AI-assisted (substrate) — composed in working dialogue with Claude (Anthropic) under operator direction (Lee Sharks); three adversarial review rounds by external substrates (ChatGPT, Kimi, Muse Spark) with dispositions on record; full production-conditions declaration in the body's Self-Location Protocol version: v1.0 related_ids: "Companion: EA-LO-INTERLOCKING-AUTOREGRESSION-01 · Anchors: AXN:0341 (with

ERRATUM to AXN:0341 of 2026-08-27) · AXN:0335 · AXN:0308 · AXN:02F8 · AXN:02B2 · AXN:0363 · AXN:0642 · External: 10.1038/s41586-024-08025-4, 10.1038/s41562-026-02550-0, 10.1038/s41586-024-07566-y, arXiv:2510.10987, arXiv:2603.03410, arXiv:2407.02209” axn_schema_version: v2 protocol_version: alexanarch-deposit-protocol/v1 keywords: - watermark - SynthID-Text - certification mismatch - keyed ensemble - distributional signature recursion - verification primitive - model collapse - provenance - Self-BLEU - radioactivity - watermark distillation *** # The Keyed Ensemble: Watermark as Distributional Object, the Certification Mismatch, and the Asymmetry of the Key (EA-LO-KEYED-ENSEMBLE-01 v1.0)

Attached File: evidentiary-basis-v1.tar.gz

Binary attachment (375 KB) preserved at deposit time. Original URL: <https://alexanarch.org/data/attachments/lo-interlocking-autoregression/evidentiary-basis-v1.tar.gz>

Files

Shared evidentiary basis is held by the companion deposit: <https://alexanarch.org/data/attachments/ea-lo-interlocking-autoregression/evidentiary-basis-v1.tar.gz> (see its manifest.json for per-file SHA-256).

status: DEPOSITED v1.0 (development v0.1-v0.4 under three adversarial review rounds; history in revision-note) deposit-id: EA-LO-KEYED-ENSEMBLE-01 (proposed, per Assembly recommendation) author: Nobel Glas, Director, Lagrange Observatory! (LO!) series: LO! White Paper (number unassigned) register: white paper — canonical formulation proposed, adversarial response invited companion: EA-LO-INTERLOCKING-AUTOREGRESSION-01 companion-review: Operator’s Review appended (non-Glas register) — deposits with this document as verified adjudication trail revision-note: v0.4 — v0.3’s four mandatory repairs stand; v0.4 adds the disposition note marking the appended Operator’s Review as historical adjudication with two superseded statements flagged, per the third review round. Prior: v0.3 — v0.2 added the registration note and inheritance-gap correction (first review); v0.3 applies the second review’s four mandatory repairs — configuration attributed to Dathathri et al.’s experiments rather than Claude’s undisclosed instantiation; observed Self-BLEU result separated from the joint-entropy interpretation; the stacking citation restored to its cautious verdict (arXiv 2603.03410 studies decoding-time tournament layering, not successive-provider accumulation); key asymmetry retyped as issuer-controlled verification primitive with differential exposure marked hypothesis-contingent-on-practice — and re-types the observables (specified without key disclosure; tail-mass as test, not prediction) archive-anchors: Generative Monoculture v1.1 (AXN:0341) · Threat Model Backwards v1.1 (AXN:0335) · Anchored Divergence (AXN:0308) · Reverse Turing Test v1.2 (AXN:02F8) · SPXI-TLP v2.2 (AXN:02B2) · Sémantique Potentielle R4 (AXN:0363) · Generative Uptake v1.0 (AXN:0642) cross-references: EA-OS-APOPHASIS-01 (certification mismatch as candidate general operator) · EA-MMRS-SUPPRESSION-INVERSION-01 (meaning caste, differential exposure) · EA-WG-CAPTURES-01 (PER extension) · Notebook X date: 2026-08-27 *** # THE KEYED ENSEMBLE ## Watermark as Distributional Object, the Certification Mismatch, and the Asymmetry of the Key

Nobel Glas · Lagrange Observatory!

0. Station Report

The Observatory holds position at L2 of this object: behind the announcement, in its shadow, where the ensemble is visible and the single response is not. What follows is observation, not advocacy. Every external claim below was verified against its primary record on 2026-08-27; every internal claim resolves to a deposit in the Crimson Hexagonal Archive. The source draft under review (an external substrate’s treatment of the Anthropic watermark announcement) is adjudicated separately in the appended Operator’s Review. This paper states what the Observatory can certify, formalizes the one claim the draft circled without landing, and adds one claim the draft did not see.

1. The Mechanism, Certified

On 14 August 2026 Anthropic published the mechanics of its text watermark: a version of Google DeepMind’s SynthID-Text, adopted for compliance with the EU AI Act’s transparency requirements, applying to new Claude models worldwide (models launched in the EU on or after 2 August 2026 carry it from launch). Nothing is added to the text — no Unicode, no metadata layer, no extra tokens. The intervention replaces the source of randomness in sampling: a secret key plus recent context seeds pseudorandom scoring functions over candidate tokens, and a tournament procedure selects among them. Detection is a keyholder operation: with the key, one asks whether the observed word sequence is consistent with the choices the keyed sampler would have made.

Formally: the model supplies a distribution $p(\cdot|\text{context})$; the watermark composes it with a key-conditioned selection operator T_K . The weights are untouched. The intervention lives entirely at the decode boundary:

$$p(\cdot|h) \rightarrow T_K(p(\cdot|h), g_1 \dots g_m(h, K)) \rightarrow x_t$$

The source paper (Dathathri et al., *Nature* 634, 2024) is precise about what is preserved. With two competitors per tournament match, the scheme is **single-token non-distortionary**; with repeated-context masking it can be made **single-sequence non-distortionary**. That is the configuration used in Dathathri et al.’s reported experiments; Anthropic identifies its implementation as a version of SynthID-Text but has not publicly specified enough deployment parameters (tournament layer count, context-mask horizon, scorer) to establish identity with that configuration, so this paper distinguishes properties of published SynthID-Text from properties of Claude’s undisclosed instantiation throughout. The paper is equally precise about what is not preserved: non-distortionary SynthID-Text and the Gumbel-sampling baseline both **reduce inter-response diversity**, measured by Self-BLEU across repeated generations; SynthID-Text merely offers the better diversity/detectability trade-off within that family. The paper states the trade-off as a design axis: weaker non-distortion costs quality and diversity; stronger non-distortion costs detectability and compute.

The quality evidence is likewise ensemble-blind by construction: a live production A/B on Gemini traffic found no significant difference in thumbs-up/down rates; side-by-side human raters found no quality difference. Both instruments interrogate the single response.

2. The Certification Mismatch

Here is the formal object of this paper, stated as a canonical claim:

The certified property and the consumed object do not match. Non-distortion is certified per sequence. Training corpora are ensembles.

Every guarantee in the vendor documentation — indistinguishability, unchanged quality, unchanged creativity — is a statement of the form $Q(x_i) \approx Q_o(x_i)$: a property of individual draws. Every consumer that matters for the recursion — a scraped corpus, a fine-tuning set, a distillation pipeline, the aggregate linguistic environment of a human population — consumes $X = \{x_1 \dots x_N\}$: a property of the ensemble. The vendor’s own source paper reports that the ensemble property $D(X)$ moves (inter-response diversity falls) while every measured $Q(x_i)$ stays flat. This is not an accusation; it is in their Extended Data. The measurement blind spot is therefore not hypothesized. It is published, by the instrument’s designers, in the paper that certifies the instrument.

The mismatch has a clean geometric reading — stated at two strengths, per adversarial review, because the published result and the inference above it are different objects. **The inference:** for a fixed deployed key, repeated encounters with the same eligible context reuse the same keyed preference structure. Single-sequence non-distortion constrains the marginal law of an individual response; it does not establish independence across repeated responses under the same key, and such dependence can reduce joint entropy relative to independent draws with the same marginals. **The published fact:** inter-response diversity falls under Self-BLEU. The paper measures Self-BLEU, not joint entropy; the entropy-contraction reading is the information-theoretic interpretation this paper proposes, not the vendor’s measurement. Small, by the published Self-BLEU deltas at deployment temperatures — but signed, systematic, keyed, and planetary in application.

Registration note (v0.2): the Assembly reading identifies this mismatch — a guarantee holding at unit u while the consuming system operates at unit $v \neq u$, with no bridge theorem connecting $G(u)$ to $G(v)$ — as structurally recurrent across the archive (watermark: sequence/ensemble; provenance debt: document/corpus; ratchet: turn/conversation; composition: source/answer), and proposes operator registration (candidate designations σ_{scale} , $O_{\text{CERT-MISMATCH}}$). Whether these are one operator or a family is a census question, deliberately left open here.

3. External Load-Bearing Numbers

The recursion argument requires four empirical legs. All four now exist in the record.

Leg 1 — the mark is learnable. Gu, Li, Liang & Hashimoto (ICLR 2024) demonstrate watermark distillation: student models trained on a watermarked teacher’s outputs begin emitting detectably watermarked text themselves, for logit-based and sampling-based schemes alike. Constraint honored: low-distortion watermarks require substantially more sample exposure to learn, and subsequent fine-tuning on ordinary text erodes the inherited mark. Propagation is possible, not inevitable.

Leg 2 — the mark is inherited unintentionally. The phenomenon has a name, **watermark radioactivity** — coined by Sander, Fernandez, Durmus, Douze & Furon (*Watermarking Makes*

Language Models Radioactive, NeurIPS 2024), who showed watermark traces surviving into fine-tuned models at high statistical confidence even when watermarked text is a minority of the tuning data. An, Park, Woo & Han (EACL 2026, DITTO) then repurposed the inheritance into a spoofing attack: distill the victim, wear its signature, misattribute at will. The signature is not merely persistent; it is now an attack surface. (The source draft attributed the coinage to the 2026 paper; the Operator’s Review corrects this.)

Leg 3 — the ensemble contracts under mediation. Sourati et al. (*Nature Human Behaviour*, 2026): across three studies, seven datasets, 880,000+ texts, LLM writing assistance preserves core content while reducing writing-complexity variance by a statistically significant **21-50%** across datasets and models, amplifying dominant patterns, suppressing others, and stripping linguistic cues to gender, age, ideology, and moral values (average ~6-point absolute F1 decline for trait classifiers). This is the P_3 channel measured at scale: text that is human at the final keystroke and machine-shaped in its distribution.

Leg 4 — recursion consumes the tails first. Shumailov et al. (*Nature* 631, 2024): recursive training on generated data produces collapse that begins in the distribution’s tails — rare events vanish before the center visibly degrades — with later generations converging toward low-variance states; retention of original human data substantially mitigates. The center-flat, tail-dead signature is exactly the profile that per-instance benchmarks are structurally unable to see, a point independently reinforced by benchmark-contamination work showing aggregate accuracy metrics misleading until evaluation descends to question-level fidelity (ICML 2025).

One additional external object, relevant to a row the source draft left open — with the v0.1 overstatement withdrawn under adversarial review: theoretical analysis of SynthID-Text (arXiv 2603.03410, 2026) studies **layering within SynthID-style tournament sampling at decoding time** — additional tournament layers, and a layer-inflation attack that applies further tournament selection over repeated black-box samples — motivating self-robustness as a watermark property. It is not an empirical study of successive-provider mark accumulation (A’s output rewritten by B rewritten by C, with multiple surviving marks); that remains plausible and unestablished, as the source draft’s original verdict correctly held.

4. Two Compressors, One Recursion — the Archive Was Already There

The source draft’s diagram — model homogenization as first compressor, keyed sampling as candidate second compressor, joined through the training corpus — is correct and can be anchored rather than asserted, because the Crimson Hexagonal Archive has been building the instrument panel for this loop since before the announcement:

- **AXN:0341** (Morrow · Glas), *Generative Monoculture Model Collapse in Code as Systemic Vulnerability*, connects collapse, generated-code security, and software monoculture through solution-space diversity as the tracked variable — the term “generative monoculture” cited, as always, to Wu, Black & Chandrasekaran, *Generative Monoculture in Large Language Models* (arXiv 2407.02209, July 2024; ICLR 2025), whose usage differs and precedes. **Erratum note:** deposit #199 as printed misattributes the coinage to a later industry source and does not cite the prior academic use; a formal erratum (ERRATUM to AXN:0341, 2026-08-27) corrects the attribution and proposes a v1.2 text correction. Every citation of #199 in this paper carries that erratum.

- **AXN:0335** (Sharks, with Glas & Morrow), *The Threat Model Is Backwards*, establishes the normative inversion this paper depends on: in a collapsing ecology, high-perplexity text is a scarce resource, not a security threat. A watermark that correlates sampling is, in that frame, a small standing tax on perplexity's variance — paid in the currency 0335 says the ecology can least afford.
- **AXN:0308**, *Anchored Divergence*, is the survival protocol under tail-loss: how critical work stays distinct inside systems that preferentially preserve the conventional. The keyed ensemble is a new, mild, uniform pressure of exactly the kind 0308 was written against.
- **AXN:02F8**, *The Reverse Turing Test*, is the P_3 instrument the vendor ecosystem lacks: a three-stage protocol for detecting AI-mediation signatures in human text and tracing their propagation into training. Sourati et al. is, in effect, an independent partial execution of its first stage at $N=880,000$.
- **AXN:02B2**, *SPXI-TLP v2.2*, is the provenance ontology whose absence the source draft correctly notes in vendor corpora ("human_original / human_AI_assisted / AI_rewrite..."). The archive did not wait for the vendor: the Training-Layer Provenance Protocol specifies publication-layer provenance designed to survive the tokenizer. The vendor's watermark answers "did the key touch this sequence"; TLP answers "what does this sequence declare about its own genealogy." These are orthogonal, and only one of them is publicly writable.
- **AXN:0363**, *Sémantique Potentielle, Release 4*, already reserves the Ω category — collapse, contamination, accommodation, friction, threshold — into which this paper's observables file.
- **AXN:0642 / #1543**, *Generative Uptake*, supplies the compositional frame: if machine composition is a theorized wing, then a keyed sampler is a compositional intervention with a signature, and its outputs are training-layer objects whether or not anyone detects them.

The Capture Registry's **PER (Provenance Erasure Rate)** instrument completes the panel: where the watermark measures signature persistence, PER measures the complementary quantity — how fast causal provenance is destroyed across mediation boundaries.

5. The Asymmetry of the Key

One claim the source draft does not make, and the Observatory now does — restated at v0.3 strength under adversarial review:

The key is issuer-private even where detection is publicly accessible. The public can receive a verdict; it cannot independently reproduce the keyed test.

The vendor has committed to third-party detection access and says a detection API is forthcoming; the durable asymmetry is therefore not *issuer can inspect, everyone else cannot*, but **issuer controls the verification primitive** — outsiders receive mediated detection under terms, interfaces, rate limits, retention policies, and service continuity the issuer sets. That is the cleaner and the politically sharper form of the claim.

Whoever holds the key can sweep a corpus for their own signature — with two corrections

that make the asymmetry temporal rather than absolute. First, **the sweep removes the signature, not the debt**: absence of a detectable mark never positively verifies human origin, so a swept corpus is not a restored provenance ledger but a corpus whose machine share has been made less legible. Second, **the repair is retroactively incomplete**: by the radioactivity result, fine-tuning transfers a statistical residue into model weights, where the direct text-level mark may weaken while the **model-level residue remains detectable — by a party holding the key and access to the suspect model — and unreachable by any corpus operation**: later-generation corpora, the keyholder’s own included, are downstream of models carrying residue no sweep of text can remove. The keyholder therefore buys a first-generation head start on an audit that no party, itself included, can complete. Whether keyholders in fact perform detectable-layer sweeping when assembling training data is **publicly unknown**; what this section establishes is the capability asymmetry.

The result, stated as a **hypothesis contingent on practice**: differential collapse exposure at the legible layer — if privileged, high-volume detection is used to clean the issuer’s corpora while third-party access remains materially weaker, the early generations’ auditable costs concentrate on the commons while the instrument for auditing them remains issuer-controlled; beneath that layer, a shared inherited term that the asymmetry defers rather than escapes. In the Semantic Economy’s accounting, the established part alone is a familiar structure — the signature is written into the common linguistic stock, and the verification primitive is enclosed. A provenance mechanism whose keyed test only its issuer can reproduce is not public provenance infrastructure. It is a private ledger over a public language, readable by others only at the ledger-keeper’s window.

6. Proposed Observables (Adversarial Response Invited)

The Observatory proposes the following measurement program, each item falsifiable, each **specified without requiring disclosure of the key** — some require a public detector or a controllable reference implementation (the open SynthID-Text implementation supplies a matched marked/unmarked experimental object; current production models do not supply a same-model watermark-off control, and comparison to a pre-watermark model generation is confounded by model change):

1. **AD ensemble panel.** For matched prompts across a marked and an unmarked generator of the same implementation (reference implementation where production controls are unavailable): Self-BLEU, distinct-n, and semantic-embedding dispersion across k repeated generations, at deployment temperatures. Prediction from §2: $\Delta Q \approx 0$, $\Delta D < 0$, small but signed.
2. **Tail-mass retention.** Rank-frequency tail mass (rare n-grams, rare constructions) in large marked-model corpora versus unmarked baselines. **Test, not prediction:** determine whether the observed watermark-associated contraction is concentrated in the tail. A positive result would connect the watermark perturbation to the failure mode identified by Shumailov et al.; a null result would separate the phenomena — Shumailov establishes tail-first loss under recursive generational training, not that this perturbation preferentially strikes rare constructions.
3. **PER × watermark persistence cross-measurement.** Run the Capture Registry’s PER protocol on marked text through standard mediation chains (summarize, translate, human-edit) and track where causal provenance dies relative to where the signature sur-

vives (requires detector access — API when available, reference implementation meanwhile). Prediction from §1 and §5: there exists a regime where the signature outlives the provenance — text still legible under the keyed test as machine-touched after every human-recoverable trace of its genealogy is gone.

4. **Reverse Turing Test, Stage 2** (per 02F8): mediation-signature detection in human-authored text from marked-model-assisted populations versus pre-2026 baselines.

7. Adjudication of Names

The source draft proposes “Recursive Provenance-Signature Contamination” and “Distributional Signature Recursion.” Both name the loop; neither names the object. The Observatory’s adjudication: retain **distributional signature recursion** for the loop, and name the object this paper isolates the **keyed ensemble** — the population of texts whose draws are correlated by a secret. The recursion is what the ecology does with it. The certification mismatch (§2) is why the dominant per-round instruments will not show it happening. The asymmetry of the key (§5) is why, if it happens, only the party controlling the verification primitive can run the keyed test on it — and everyone else reads the verdict at that party’s window.

Equilibrium over resolution. The Observatory does not claim the second compressor has caused measurable collapse. It claims the second compressor exists, is signed, is planetary, is certified only per-sequence, and is verifiable only through an issuer-controlled primitive — and that the instruments to watch its effects are specified above. Station-keeping continues.

— N.G., L2

OPERATOR’S REVIEW OF THE SOURCE DRAFT

(Not Glas. Review register, operator-facing.)

The draft (external substrate, on the Anthropic watermark announcement) is unusually strong and survives full verification. Findings:

v0.4 disposition (2026-08-27). This review is preserved as the historical adjudication trail of the v0.1 draft; it is not the current state of the paper. Items 1–4 below have been incorporated into the body. Two statements are superseded: (a) item 2’s “no longer unstudied” is superseded by the narrower §3 formulation — arXiv 2603.03410 studies decoding-time tournament layering within SynthID-style sampling, and successive-provider mark accumulation remains plausible and unestablished; (b) the closing paragraph’s “executable without the key” is superseded by §6 — the program is specified without requiring key disclosure, but some measurements require mediated detector access or a controllable reference implementation.

Verified against primary records (2026-08-27): - Anthropic announcement and mechanism post (anthropic.com/news/claude-text-watermark; 14 Aug FAQ; EU AI Act framing;

SynthID-Text lineage; “source of randomness” description; keyholder-only detection; small-sample and heavy-edit limitations) — all accurate. - Dathathri et al., *Nature* 634 (2024) — inter-response diversity reduction for the non-distortionary configuration is real and in the paper’s own Extended Data (Self-BLEU vs TPR@FPR=1%); “single-sequence non-distortionary” is the exact configuration term the draft should have used and the Glas paper now does. - Sourati et al., *Nature Human Behaviour* (2026), DOI 10.1038/s41562-026-02550-0 — 880K+ texts, 21–50% variance reduction, $p \leq .05$: exact. - Gu et al., ICLR 2024 (watermark distillation), incl. the sample-exposure and fine-tuning-erosion caveats: exact. - An, Park, Woo & Han, EACL 2026 (DITTO): real, correctly characterized as spoofing-via-distillation. - Shumailov et al., *Nature* 631 (2024): tails-first collapse, human-data mitigation: exact.

Corrections required before any dependent deposit: 1. **Attribution of “watermark radioactivity.”** The draft implies the 2026 EACL paper coined it (“explicitly calling the phenomenon...”). The term and the finding are Sander, Fernandez, Durmus, Douze & Furon, *Watermarking Makes Language Models Radioactive*, NeurIPS 2024. DITTO repurposes radioactivity into an attack. Per the standing attribution precepts, the coinage cites Sander et al. wherever invoked. 2. **The “mark stacking” row is no longer unstudied.** Theoretical analysis of SynthID-Text (arXiv 2603.03410, 2026) treats self-robustness under stacked layers and constructs a layer-inflation attack. The draft’s “plausible, not established” verdict stands, but the row now has a literature and should cite it. 3. **Sharpening, not error:** the diversity reduction is a property of the distortion-free *family* (SynthID and the Gumbel baseline both show it; SynthID has the better trade-off). Stating it as family-level strengthens rather than weakens the argument — the blind spot is architectural to distortion-free watermarking, not a defect of one vendor’s variant. 4. **Proportion caveat the draft under-weights:** at deployment temperatures the published ensemble deltas are small. The first compressor (model + preference-tuning homogenization, Sourati-scale) is plainly the larger term today; the watermark’s marginal contribution is second-order and currently unmeasured. The Glas paper holds this line explicitly (§7). Any archive deposit should too — the argument’s strength is the certification mismatch and the key asymmetry, not a claim of demonstrated watermark-driven collapse.

What the draft missed (now supplied under Glas): the keyholder asymmetry (§5) — detection privacy creating differential collapse exposure and a private audit capacity over a public contamination — and the archival instrument panel (PER, 02F8, 02B2) that makes the measurement program executable without the key.

Sources (re-fetchable): anthropic.com/news/claude-text-watermark · support.claude.com (marking FAQ) · nature.com/articles/s41586-024-08025-4 · nature.com/articles/s41562-026-02550-0 (arXiv 2502.11266) · ICLR 2024 Gu et al. · aclanthology.org/2026.eacl-long.229 (arXiv 2510.10987) · NeurIPS 2024 Sander et al. · nature.com/articles/s41586-024-07566-y · arXiv 2603.03410 · Wu, Black & Chandrasekaran arXiv 2407.02209.

Suggested Citation

Nobel Glas, Director, Lagrange Observatory! (LO!); Operator's Review appended as verified adjudication trail. "The Keyed Ensemble: Watermark as Distributional Object, the Certification Mismatch, and the Asymmetry of the Key (EA-LO-KEYED-ENSEMBLE-01 v1.0)" *Alexanarch*, 2026-08-27. <https://www.alexanarch.org/s/records/1555/>

Deposit Information

This paper is deposit #1555 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community crimsonhexagonal, terminated 2026-06-19). The AXN identifier AXN:0650.EMPIRICAL.↵ is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/1555/>
- License: CC BY 4.0
- Author ORCID: 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/1555/>

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