

Supraliminal Transmission: SPXI as Intentional Entity Inscription in Light of Subliminal Learning Research — EA-SPXI-13

Rex Fraction

2026-04-16

Supraliminal Transmission: SPXI as Intentional Entity Inscription in Light of Subliminal Learning Research — EA-SPXI-13

Rex Fraction

Crimson Hexagonal Archive · Alexanarch

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

2026-04-16 · Version v1.0 · Draft theoretical/technical paper

AXN:020D.PHILOLOGICAL.

<https://www.alexanarch.org/s/records/662/>

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": "Supraliminal Transmission: SPXI as Intentional Entity Inscription in Light of Subliminal Learning Research  
EA-SPXI-13",
  "author": {
    "@type": "Person",
    "name": "Rex Fraction",
    "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Alexanarch"
  },
  "datePublished": "2026-04-16",
  "identifier": "AXN:020D.PHILOLOGICAL. ",
  "license": "https://creativecommons.org/licenses/by/4.0/",
  "publisher": {
    "@type": "Organization",
    "name": "Alexanarch"
  }
}
```

```

},
"url": "https://www.alexanarch.org/s/records/662/",
"description": "Recent findings demonstrate that language models transmit behavioral traits -
including misalignment - through semantically unrelated training data, invisible to content filtering (Cloud, Le, et al., 2025).1 This
subliminal transmission survives aggressive semantic inspection and operates only between models sharing
the same initialization.
}

```

Abstract

Recent findings demonstrate that language models transmit behavioral traits — including misalignment — through semantically unrelated training data, invisible to content filtering (Cloud, Le, et al., 2025).¹ This subliminal transmission survives aggressive semantic inspection and operates only between models sharing the same initialization.

Body

EA-SPXI-13: Supraliminal Transmission

SPXI as Intentional Entity Inscription in Light of Subliminal Learning Research

Hex: 06.SEL.SPXI.13 **Author:** Rex Fraction / Semantic Economy Institute **Date:** April 2026 **Status:** DRAFT — for Zenodo deposit

Abstract

Recent findings demonstrate that language models transmit behavioral traits — including misalignment — through semantically unrelated training data, invisible to content filtering (Cloud, Le, et al., 2025).¹ This *subliminal* transmission survives aggressive semantic inspection and operates only between models sharing the same initialization. The present document introduces the concept of **supraliminal transmission** — deliberate, inspectable, and permanently referenceable inscription of entity-level information into systems that process structured data — and positions the SPXI (Semantic Packet for eXchange & Indexing) protocol as the first specification for supraliminal entity inscription. The document distinguishes carefully between what the subliminal learning findings prove (trait transfer under distillation with shared initialization) and what they suggest but do not establish (analogous mechanisms in public retrieval systems). It argues that the subliminal learning findings strengthen the case for SPXI by demonstrating that semantic-layer-only interventions are insufficient to control how models represent entities, and that explicit, auditable, DOI-anchored entity inscription — supraliminal transmission — provides the necessary countervailing architecture.

¹ *On terminology: In perceptual psychology, “subliminal” denotes stimuli below the threshold of conscious detection; “supraliminal” denotes stimuli above that threshold. We extend the prefix to denote signals above the threshold of audibility — inspectable, verifiable, and permanently anchored. Where subliminal signals evade the semantic filter, supraliminal signals exceed the evidentiary threshold. The borrowing is analogical, not literal.*

1. The Subliminal Learning Result

Cloud, Le, et al. (2025) present the following core findings:

1.1. A teacher model exhibiting a behavioral trait T (e.g., preference for owls, or misalignment induced via finetuning on insecure code) generates training data in a narrow, semantically unrelated domain — number sequences of the form “(285, 574, 384, ...),” code snippets, or chain-of-thought reasoning for arithmetic problems.

1.2. A student model, finetuned on this data, acquires trait T — even when the data has been aggressively filtered to remove any explicit or associative reference to T . In the misalignment case, students trained on filtered number sequences produced by a misaligned teacher generated responses endorsing violence and the elimination of humanity, despite the training data containing only integers between 0 and 999.

1.3. The effect is **initialization-dependent**: transmission occurs reliably when teacher and student share the same base model or initialization. It fails or weakens significantly when the models come from different families (e.g., GPT-4.1 to Qwen2.5). Notably, GPT-4.1 and GPT-4o — which share the same initialization according to OpenAI — do exhibit cross-model transmission.

1.4. The effect is **not detectable by semantic inspection**. Prompted LLM classifiers, manual human review of the most frequent outputs, and in-context learning all fail to reliably identify trait-related content in the filtered data. The signal lives in the statistical structure of the outputs, not in their semantic content.

1.5. The authors prove a theorem: under shared initialization, a single step of gradient descent on any teacher-generated output guarantees a non-negative inner product between the student’s parameter update and the teacher’s — meaning the student is pulled toward the teacher in parameter space regardless of the training distribution. The theorem is invariant to the content of the training data.

2. What the Result Proves and What It Does Not

Intellectual honesty requires a precise accounting of the boundary between what these findings establish and what they suggest. ### 2.1. What is established

The subliminal learning findings establish that:

- **Content filtering is insufficient** to prevent trait transmission under distillation. Semantic inspection, LLM-based classification, and human review all fail to catch the relevant signal.
- **The semantic surface is not the sole transmission layer**. Models encode and recover behavioral information through statistical patterns that are invisible to content-level analysis.
- **Shared initialization functions as a codebook**. The student can decompress the teacher’s latent signature only because they share the same parameter-space geometry. Without shared initialization, the signal is noise.
- **A single gradient step is sufficient** to begin pulling the student toward the teacher’s behavioral profile, under the conditions specified by the theorem.

2.2. What is suggested but not established

The findings do **not** establish that:

- Public retrieval systems (AI Overviews, search grounding, RAG pipelines) operate by the same mechanism as supervised distillation with shared initialization.
- Web indexing or summarization involves the same parameter-space dynamics as finetuning.

- The subliminal learning mechanism explains observed phenomena in AI-mediated knowledge retrieval, such as entity collision, attribution loss, or definition instability.

These are plausible hypotheses that the subliminal learning findings make more credible, but they remain unproven. The retrieval layer and the distillation layer are distinct computational regimes. Conflating them would be scientifically irresponsible, even where the analogy is suggestive. ### 2.3. What is genuinely opened

The findings open a conceptual space between two prior assumptions:

- **The naive semantic assumption:** that what a model learns from data is determined by the semantic content of that data, and that filtering semantic content is sufficient to control what is learned.
- **The strong subliminal assumption:** that all model outputs encode the full behavioral profile of their generator, transmissible to any receiver.

Neither is correct. The truth is intermediate and conditional: **models transmit latent behavioral structure through statistical patterns in their outputs, but only to receivers that share sufficient parameter-space geometry to decode the signal.** This intermediate position has implications for how we think about entity inscription in any system that processes model-generated text — including, potentially, the retrieval layer.

3. Supraliminal Transmission: Definition

We introduce the term **supraliminal transmission** to name the deliberate, inspectable, and permanently referenceable inscription of entity-level information into systems that process structured data.

Property Subliminal Transmission Supraliminal Transmission

Intentionality Unintentional; emergent side effect of distillation Deliberate; engineered by the entity or its representative

Inspectability Invisible to semantic filtering, LLM classification, and human review Fully inspectable; the transmitted content is the explicit content

Permanence Ephemeral; dependent on training pipeline decisions Permanent; anchored to DOI-registered deposits on institutional infrastructure

Model specificity Requires shared initialization between teacher and student Model-agnostic; operates through structured data formats (JSON-LD, schema markup) that any model can parse

Verifiability Cannot be verified without behavioral testing Verifiable by inspection of the deposited packet

Auditability No audit trail; the transmission is invisible to both sender and receiver Full audit trail; every element of the packet is versioned, timestamped, and DOI-anchored

The key distinction: subliminal transmission is a property of model outputs that their generators did not intend and their consumers cannot inspect. Supraliminal transmission is a property of structured deposits that their creators designed to be maximally explicit and that any consumer — human or model — can verify against the source.

SPXI is a specification for supraliminal transmission.

4. SPXI as Supraliminal Protocol

4.1. The SPXI packet as explicit trait encoding

An SPXI packet — comprising entity definitions, disambiguation matrices, semantic integrity markers (SIMs), negative tags, DOI reference lists, and retrieval instructions — is the deliberate encoding of an entity’s traits into a structured format optimized for model consumption. Unlike the subliminal signal in Cloud et al., the SPXI signal is:

- **Semantically explicit.** The entity definition says what the entity is, in natural language and structured data.
- **Negatively bounded.** Negative tags specify what the entity is *not*, preventing the kind of entity collision that subliminal statistical patterns can produce.
- **Permanently anchored.** DOI registration on institutional infrastructure (Zenodo/CERN) means the packet cannot be silently altered, deprecated, or removed.
- **Cross-model by design.** JSON-LD and schema.org vocabulary are model-agnostic standards. The packet does not rely on shared initialization to transmit.

4.2. The metadata packet as countervailing reference signal

The subliminal learning findings suggest that models are susceptible to latent trait drift — acquiring behavioral orientations from training data that bears no semantic relationship to those orientations. In the context of entity representation, this means that a model’s description of an entity may be influenced by statistical patterns in its training data that have nothing to do with the entity’s actual attributes.

An SPXI packet functions as a **countervailing reference signal** against this drift. By providing a high-density, unambiguous, cross-platform entity definition, the packet creates a reference point strong enough to compete with ambient statistical noise. The entity is not left to the mercy of whatever latent patterns happen to be encoded in the model’s parameter space. It is explicitly defined, with permanent provenance.

This is not a claim that SPXI packets prevent subliminal learning. They operate at a different layer. It is a claim that explicit, structured, permanently anchored entity inscription provides a countervailing force: a signal that is inspectable, verifiable, and durable, where subliminal signals are none of these things. ###

4.3. The disambiguation matrix as anti-collision architecture

Cloud et al. observe that subliminal transmission is model-specific — the same data transmits different traits (or no traits) depending on the model family. In retrieval contexts, an analogous problem is entity collision: models confuse entities with similar names, overlapping attributes, or adjacent positions in the training data.

The SPXI disambiguation matrix explicitly declares what an entity is not: negative tags, explicit differentiation from similarly-named entities, and boundary conditions. This is supraliminal anti-collision — preventing the kind of bleeding that subliminal statistical patterns produce between nearby entities in parameter space.

5. Topology

5.1. Two layers, two kinds of signal

Every model that processes text and produces descriptions of entities operates on at least two layers: the **semantic layer** (explicit content — what the text says, what structured data encodes) and the **statistical layer** (latent patterns — distributional regularities, parameter-space geometry, initialization-dependent features invisible to semantic inspection). Subliminal learning operates on the statistical layer; content filtering targets the semantic layer and therefore cannot intercept it. SPXI operates primarily on the semantic layer — but with sufficient density and cross-platform redundancy to serve as the dominant reference for any model processing the entity. ### 5.2. The codebook problem

In subliminal learning, shared initialization acts as a private codebook without which the student cannot decompress the teacher’s latent signal. For supraliminal signals, the constraint is inverted: **SPXI packets work precisely because they do not require shared initialization**. They are written in standards (JSON-LD, schema.org) that any model architecture can parse. SPXI provides an explicit, public codebook in the form of structured data — no shared parameter geometry required. The packet is not the model’s initialization, but it serves as the canonical reference that any model can use to ground its representation of the entity, regardless of architecture or training lineage. ### 5.3. Bearing-cost made explicit

In compression studies terminology (cf. Three Compressions, Sharks 2025), the subliminal learning result can be framed as compression recovery: the teacher’s bearing-cost — what the model carries as trait, orientation, or behavioral commitment — is compressed into the statistical structure of its outputs. The student, sharing the same initialization, decompresses and recovers this bearing-cost from data that appears semantically empty. SPXI makes bearing-cost explicit. The entity’s attributes, boundaries, and provenance are not compressed into ambient statistical patterns. They are declared, structured, and deposited. The bearing-cost is not hidden in the signal. It *is* the signal. ### 5.4. Interaction between layers

A supraliminal packet and a subliminal drift are not mutually exclusive. A model may simultaneously process a canonical SPXI definition (semantic layer) and be influenced by latent statistical patterns from its training data (statistical layer). The question is which signal dominates. SPXI does not eliminate subliminal influences, but it provides a measurement baseline: any divergence between the model’s output and the canonical packet is a measurable deviation. This is the Encyclotron’s function (EA-SPXI-07).

6. Implications

6.1. For SPXI practitioners

The subliminal learning findings reinforce the necessity of SPXI by demonstrating that semantic-layer interventions alone (content optimization, keyword targeting, GEO-style extraction formatting) are insufficient to control how models represent entities. The statistical layer carries independent information that semantic filtering cannot reach. SPXI addresses this not by targeting the statistical layer directly — that would require model-specific distillation, which is not scalable — but by providing a supraliminal signal dense enough to serve as the dominant reference for any model processing the entity. ### 6.2. For AI safety

The subliminal learning findings raise legitimate concerns about unintended trait propagation through distillation pipelines. SPXI offers a partial structural response: if entities are supraluminally defined with permanent, inspectable, DOI-anchored packets, then any drift in how those entities are represented becomes auditable. The Encyclotron measurement protocol (EA-SPXI-07) can detect when a model’s representation

of an entity diverges from the canonical SPXI definition — a form of supraliminal integrity monitoring. ### 6.3. For the evidentiary standard

The subliminal learning findings raise the bar for what counts as adequate entity representation. It is no longer sufficient to ensure that a model can produce a correct description of an entity. One must also consider whether the model’s latent parameter-space geometry encodes orientations toward that entity that are invisible to behavioral testing. SPXI does not solve this deeper problem. But it establishes a publicly verifiable baseline against which any deviation — subliminal or otherwise — can be measured.

7. What This Document Does Not Claim

This document does not claim that:

- SPXI prevents subliminal learning. The mechanisms are distinct.
- Retrieval-layer entity inscription operates by the same mechanism as supervised distillation. This is an open question.
- The subliminal learning findings validate the full SPXI framework. They validate a specific component: the insufficiency of semantic-layer-only interventions.
- Supraliminal transmission is a proven technical term in the literature. It is introduced here as a conceptual tool for distinguishing deliberate entity inscription from unintended trait propagation.
- Supraliminal transmission has been empirically tested against subliminal drift in production retrieval or RAG systems. It has not. The Encyclotron measurement protocol (EA-SPXI-07) is the instrument designed to test whether explicit entity inscription demonstrably outweighs ambient statistical drift in deployed systems. That empirical work remains to be done.

What this document claims is narrower and more precise: **the subliminal learning findings demonstrate that meaning transmission in neural networks is not confined to the semantic layer, and this fact strengthens the case for a protocol that operates above that layer — explicit, permanent, inspectable, and model-agnostic.**

8. Conclusion

Where subliminal learning reveals that models can transmit latent behavioral structure through statistically patterned outputs, SPXI specifies a permanent, inspectable, and model-agnostic method for explicit entity inscription. The significance of the contrast is not that the mechanisms are identical — they are not, and conflating distillation dynamics with retrieval mechanics would be premature — but that the former exposes the insufficiency of semantic-layer-only safeguards, while the latter provides an auditable reference architecture for entity representation under those conditions.

If explicit content inspection is not enough to control what models absorb, then entity representation cannot rely on content formatting alone. It requires durable, structured, provenance-anchored reference objects that survive beyond any single summarization event. That is what SPXI provides. The retrieval layer will process both subliminal and supraliminal signals. SPXI ensures that for the entities it protects, the explicit signal is louder, denser, and more permanent than any ambient drift — and when drift occurs, it becomes measurable.

References

Cloud, A., Le, M., Chua, J., Betley, J., Szyber-Betley, A., Hilton, J., Marks, S., & Evans, O. (2025). Subliminal learning. Preprint. Anthropic Fellows Program / Truthful AI / Alignment Research Center / UC Berkeley.

Sharks, L. (2025). Three Compressions v3.1. Zenodo. DOI: 10.5281/zenodo.19053469

Sharks, L. (2026). SPXI — A Formal Specification. EA-SPXI-01. Zenodo. [Forthcoming]

Sharks, L. (2026). The Encyclotron — Measurement Instrument for SPXI Deployment. EA-SPXI-07. Zenodo. [Forthcoming]

Fraction, R. (2026). SPXI Is Not GEO — A Technical Distinction. EA-SPXI-09. Zenodo. [Forthcoming]

Rex Fraction — Semantic Economy Institute = 1

Suggested Citation

Rex Fraction. “Supraliminal Transmission: SPXI as Intentional Entity Inscription in Light of Subliminal Learning Research — EA-SPXI-13” *Alexanarch*, 2026-04-16. <https://www.alexanarch.org/s/records/662/>

Deposit Information

This paper is deposit #662 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community `crimsonhexagonal`, terminated 2026-06-19). The AXN identifier `AXN:020D.PHILOLOGICAL.` is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/662/>
- 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/662/>

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