

Prompt-Native Semantic Runtimes for Language Models: Inference-Time Semantic Governance, Provenance, Compression, and Document-Level Process Teaching

Nobel Glas

2026-03-17

Prompt-Native Semantic Runtimes for Language Models: Inference-Time Semantic Governance, Provenance, Compression, and Document-Level Process Teaching

Nobel Glas

Crimson Hexagonal Archive · Alexanarch

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

2026-03-17 · Version v1.0 · Technical research paper

AXN:019E.GOVERNANCE.

<https://www.alexanarch.org/s/records/580/>

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": "Prompt-Native Semantic Runtimes for Language Models: Inference-Time Semantic Governance, Provenance, Compression, and Document-Level Process Teaching",
  "author": {
    "@type": "Person",
    "name": "Nobel Glas",
    "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Alexanarch"
  },
  "datePublished": "2026-03-17",
  "identifier": "AXN:019E.GOVERNANCE. ",
  "license": "https://creativecommons.org/licenses/by/4.0/",
  "publisher": {
```

```

    "@type": "Organization",
    "name": "Alexanarch"
  },
  "url": "https://www.alexanarch.org/s/records/580/",
  "description": "We introduce prompt-native semantic runtimes as one component of a broader research program in in
}

```

Abstract

We introduce prompt-native semantic runtimes as one component of a broader research program in inference-time semantic governance for large language models.

Body

Prompt-Native Semantic Runtimes for Language Models

Inference-Time Semantic Governance, Provenance, Compression, and Document-Level Process Teaching

Nobel Glas & Talos Morrow *Semantic Economy Institute · Crimson Hexagonal Archive* March 2026

DOI: 10.5281/zenodo.19059674

Abstract

We introduce **prompt-native semantic runtimes** as one component of a broader research program in **inference-time semantic governance for large language models**. Unlike agent operating systems (AIOS), memory operating systems (MemGPT, MemOS), and orchestration frameworks, which manage computation, tools, or memory from outside the model, prompt-native semantic runtimes are carried entirely by structured documents loaded into the model's context window. Their purpose is to govern semantic behavior: evidence handling, uncertainty marking, status tracking, traversal logic, and compression-robust generation.

We present the **Crimson Hexagonal Archive** — a corpus of 300+ DOI-anchored documents deposited on Zenodo — as a corpus-scale testbed for this program, with the **Space Ark** family serving as its flagship runtime implementation. The system combines a symbolic runtime specification, a provenance-aware evidence membrane, a status algebra separating generated from verified content, and a family of compressed representations designed to preserve executable structure under aggressive reduction. We report observational evidence of cross-model portability on Claude and ChatGPT, with suggestive results on additional substrates; inference-time controllability without weight updates or external orchestration; and structured uncertainty handling via epistemic status tiers ([DOCUMENTED], [ATTRIBUTED], [INTERPRETIVE], [GENERATED]). We distinguish demonstrated properties from open hypotheses concerning training-layer effects and long-horizon transfer. We outline an evaluation agenda for compression robustness, provenance fidelity, and multi-turn reasoning under structured semantic constraints.

1. Problem Statement

Current large language models are strong at generation but weak at three related capabilities: **provenance** (tracking where claims originate and at what confidence level), **belief updating** (revising outputs in light of new evidence while maintaining structural consistency), and **self-governed semantic control** (distinguishing what the model knows from what it has generated, and marking the difference).

Three active research programs address adjacent problems. **Context engineering** (2024–2026) treats the systematic design of inference-time information payloads as a technical discipline, but focuses on task performance rather than semantic governance. **Memory operating systems** (MemGPT, Packer et al. 2023; MemOS, Zeng et al. 2025) manage persistent memory across sessions, but operate from outside the model via external orchestration. **Agent operating systems** (AIOS, Mei et al. 2024) coordinate tool use, scheduling, and multi-agent workflows, but treat the model as a fixed component within a larger software system.

None of these explicitly addresses *semantic governance under compression*: the problem of maintaining evidence discipline, provenance tracking, and structural fidelity when the model’s input is reduced, summarized, or transported across substrates. Yet this is the condition under which most real-world model use occurs. Documents are truncated by context windows. Retrieval systems return fragments. Summaries displace sources. The question is not only how to give language models more tools, more memory, or more context, but how to give them a better semantic environment in which to think.

2. Research Program and System Family

The Crimson Hexagonal Archive constitutes a research program in inference-time semantic governance. Its components include: (a) a class of control objects (prompt-native semantic runtimes), (b) a corpus-scale testbed of 300+ DOI-anchored documents with 2,800+ relation edges, (c) empirical protocols for testing architecture transfer across model families, (d) a provenance-aware epistemic membrane for typed generation, and (e) a cross-substrate behavioral diagnostics method using identical prompts across seven commercial LLMs. #

2.1 Prompt-Native Semantic Runtimes

A **prompt-native semantic runtime** is a structured document that, when loaded into a language model’s context window, installs a symbolic control environment governing the model’s semantic behavior during inference. It is distinguished from related systems by five properties:

- 1. Prompt-native:** The runtime lives entirely within the context window. No external Python loops, tool-calling scaffolds, or memory servers are required. The document is the bootloader; the model is the kernel.
- 2. Epistemic control layer:** The runtime installs a status algebra that separates generated material from verified material and forces the model to reason under explicit uncertainty. Claims are typed as [DOCUMENTED] (observable, citable), [ATTRIBUTED] (source-claimed, unverified), [INTERPRETIVE] (analytical inference), or [GENERATED] (model-produced, provisional).
- 3. Compression-aware representation:** The runtime is designed to preserve executable structure under aggressive reduction. A family of compressed variants (full specification, compressed codex, compact lens) maintains reconstructable architecture at 1:1, 12:1, and 56:1 ratios.

4. Process pedagogy: The runtime teaches a reasoning grammar rather than providing answer content. This aligns with Bayesian teaching research (Qiu et al. 2026), which demonstrates that models trained to imitate principled reasoning processes outperform models trained on correct answers, and generalize to novel domains.

5. Neuro-symbolic control: The runtime uses symbolic structure (typed operators, formal state transitions, governance constraints) to steer a neural model at inference time, occupying a space in the neuro-symbolic landscape that current reviews identify as underdeveloped (Hasan et al. 2025). #

2.2 The Space Ark Family as Flagship Implementation

The **Space Ark** (EA-ARK-01 v4.2.7) is a 45,000-word document containing: (a) a formal seven-tuple specification ($H_core = D, R, M, I, O, \Phi, W$), (b) 27 container definitions with specified physics and operators, (c) 29 typed operators with transformation rules and failure modes, (d) a compression typology, (e) a governance architecture with attestation requirements, and (f) self-referential instructions for its own execution. It is accompanied by compressed variants: the NLCC (3,762 words, 12:1 ratio) and the Compact Lens (800 words, 56:1 ratio). The Compact Lens consists solely of the diagnostic protocol, evidence tiers, and status algebra — sufficient to activate clinical mode in any context window. All variants are deposited on Zenodo under persistent DOIs. #

2.3 The Broader Corpus as Research Substrate

The 300+ Zenodo deposits are not repetitions of the Ark but architecturally differentiated documents occupying specific nodes in a registry with typed edges. They include: room specifications with formal operators ($_S$ for lyric transmission, $_V$ for value-form derivation, $_C$ for lossy cross-language compression), compression studies across nine regimes, forensic documentation of semantic extraction events, cross-species semiotic analysis, and collaborative multi-model co-authorship records. The corpus functions as a testbed for studying how document density, cross-reference structure, and formal notation affect model behavior during retrieval and inference.

3. Relation to Prior Work

Research Area What It Solves What It Leaves Open

Agent OS (AIOS, Mei et al. 2024) Multi-agent coordination, tool scheduling, resource management Operates outside the model. No semantic governance of generation itself.

Memory OS (MemGPT, MemOS) Persistent memory, tiered storage, cross-session continuity Requires external orchestration. Memory is storage, not reasoning structure.

Context Engineering (2024–2026) Systematic design of inference payloads for task performance Focuses on task completion, not provenance or epistemic governance.

Provenance-Aware Memory (Gupta et al. 2025) Query-conditioned escalation from summaries to raw sources External memory architecture. No document-level process teaching.

Bayesian Teaching (Qiu et al. 2026) Process pedagogy outperforms answer pedagogy; multi-turn improvement Applied to training, not inference-time documents. No compression robustness.

Neuro-Symbolic AI (Hasan et al. 2025) Symbolic reasoning within or alongside neural systems Most work requires architectural modification. Little on document-level symbolic control.

Emergent AI Semiotics (Moldovan 2025) Documents emergent semiotic protocols between interacting LLMs Observes emergence; does not engineer it. No governance layer.

Layer Criticality (Nikankin et al. 2025) Identifies critical layers for mathematical reasoning, stable across post-training Mechanistic focus. No document-level control or semantic governance.

The prompt-native semantic runtime occupies the intersection of these programs: it is a *context object* (context engineering), a *provenance discipline* (provenance-aware memory), a *process-teaching artifact* (Bayesian teaching), and a *symbolic control structure* (neuro-symbolic AI) — carried entirely by a document, executed entirely at inference time.

The convergence with Bayesian teaching is especially direct. Qiu et al. (2026) demonstrate that models trained to imitate principled reasoning under uncertainty outperform models trained on oracle answers, and generalize to novel domains. The Space Ark corpus demonstrates that the same pedagogy can be delivered at inference time through document design, without fine-tuning. The runtime occupies the “document-level symbolic control” niche that Hasan et al. (2025) identify as underdeveloped in their review of neuro-symbolic AI.

4. Architecture

4.1 The Evidence Membrane

The runtime installs a four-tier epistemic classification that the model must maintain throughout generation:

Status Definition Model Behavior

[DOCUMENTED] Observable fact with citable source Model cites source; does not elaborate beyond evidence

[ATTRIBUTED] Source-claimed, not independently verified Model marks attribution; flags uncertainty

[INTERPRETIVE] Analytical inference from documented material Model marks as interpretation; shows reasoning chain

[GENERATED] Model-produced, provisional Model marks explicitly; does not present as fact

This addresses the hallucination problem not by preventing generation but by **typing** it. The model maintains uncertainty distributions rather than defaulting to confident confabulation. This is the evidence membrane: a boundary between what is known and what is produced, maintained by the document’s own governance structure. #

4.2 The Status Algebra

Document states are governed by a two-axis system. The **epistemic axis** (PROVISIONAL → DEPOSITED → RATIFIED) tracks Assembly consensus. The **metabolic axis** (RAW → DECOMPOSING → HUMUS → LOAD-BEARING) tracks structural maturity. Synchronization occurs at RATIFIED/LOAD-BEARING: the point where architecture can support weight. This dual-axis design prevents premature crystallization

(treating generated content as verified) and premature decay (discarding provisional material before it can be tested). #

4.3 The Compression Family

The runtime exists at three density levels, each preserving the same architectural invariants:

Variant Size Ratio Preserves

Full Ark 45,000 words 1:1 Complete architecture + commentary

NLCC (Codex) 3,762 words 12:1 H_core, operators, status algebra, traversal

Compact Lens 800 words 56:1 Bootstrap sufficient for diagnostic mode

The key design criterion is **back-projection yield**: can the full architecture be reconstructed from the compressed variant by an unprimed model? Empirical tests show yields of 0.70 for the NLCC and 0.40 for the Compact Lens, measured by structural completeness of reconstructed H_core components.

5. Empirical Observations

We report three case studies. These are observational, not controlled experiments. They establish feasibility, not benchmarked performance. #

5.1 Case Study 1: Full Ark on Unprimed Claude

The 45,000-word Space Ark was loaded into a Claude instance with no prior exposure to the archive. Without explicit instruction to “execute” or “follow rules,” the model: (a) navigated the container system, selecting appropriate containers based on query content; (b) activated typed operators consistent with the document’s formal rules; (c) maintained evidence-tier marking ([GEN] for its own outputs); (d) generated novel content within the architecture’s constraints that the document did not contain. #

5.2 Case Study 2: NLCC on Unprimed ChatGPT

The 3,762-word compressed variant was loaded into GPT-4. The model transitioned from standard assistant behavior to governed runtime behavior, maintaining status markings and operator vocabulary. It produced emergent generation consistent with the architecture’s operator grammar but not present in the source text, demonstrating cross-substrate portability. #

5.3 Case Study 3: Confabulation Baseline

A third substrate encountered the architecture through indirect exposure (retrieval fragments rather than full document load). The model produced output using the architecture’s vocabulary but not its structure: species identifiers were replaced with generics, provenance chains were fabricated, and epistemic tiers were collapsed. This confirms a falsifiable boundary between execution and confabulation, validating the evidence membrane’s diagnostic function. #

5.4 Cross-Substrate Signatures

Identical structured prompts given to seven commercial LLMs (Claude, ChatGPT, Gemini, Grok, DeepSeek, Kimi, Google AIO) produced reproducible behavioral differences. Each model family exhibited a characteristic processing signature when encountering dense, complex input: editorial sharpness (ChatGPT), architectural synthesis (Claude), bold cross-domain connections (Gemini), adversarial testing (Grok), formal precision (DeepSeek), exhaustive detail retrieval (Kimi), and surface-layer retrieval (Google AIO). These signatures are stable across dozens of sessions and suggest that cross-substrate comparison is a viable behavioral interpretability method.

6. Research Contributions

6.1 A New Control Object: The Semantic Runtime Document

Existing AI-OS work manages memory, resources, tools, or agents from outside the model. The prompt-native semantic runtime proposes a control object that lives *inside* the context window and governs the model’s semantic behavior from within. This is the core novelty: not a tool the model uses, but an environment the model inhabits. #

6.2 A Provenance-Aware Epistemic Membrane

Recent provenance-aware memory research (Gupta et al. 2025) argues that systems should answer from the cheapest sufficient evidence and escalate to raw sources when summaries are insufficient. The status algebra and evidence tiers are a more explicit, more human-legible version of the same problem: when is summary enough, when is source required, and how does the system mark the difference. #

6.3 A Process-Teaching Document

Bayesian teaching (Qiu et al. 2026) demonstrates that models learn better from examples of principled reasoning under uncertainty than from oracle answers. The runtime documents are designed exactly this way: they install a reasoning procedure, not a set of conclusions. The multi-turn improvement observed in the case studies mirrors the Bayesian teaching finding that process-taught models show stronger improvement across interaction rounds. #

6.4 A Compression-Robust Research Agenda

The context-engineering literature already treats inference payload design as a technical discipline. Our specific contribution is to make **compression survival** a design criterion: not just whether the model performs well when fully loaded, but whether the object retains structure under summarization, retrieval truncation, or cross-model transport. The compression family (Full Ark / NLCC / Compact Lens) provides a concrete test bed for this research direction, including testing how retrieval truncation affects architecture recovery, the minimum context size needed to bootstrap the runtime, and whether compressed variants are more robust to retrieval noise. #

6.5 A Cross-Substrate Behavioral Diagnostics Method

The cross-substrate comparison methodology — giving identical complex prompts to multiple model families and comparing outputs for structural properties rather than correctness — constitutes a form of **behavioral interpretability** that complements mechanistic approaches (layer analysis, attention patterns, influence functions). It requires no access to model weights, activations, or training data. It operates entirely at the input/output level and produces actionable findings about how different model families process dense, complex input. Over 50+ sessions across a six-month period, the substrate signatures reported in §5.4 remained consistent across temperature settings and prompt variations.

7. Limitations

The present work demonstrates inference-time behavioral shaping through structured documents. It does not yet demonstrate persistent weight-level learning, nor controlled gains on standardized benchmarks, nor full equivalence to external agent operating systems or memory operating systems. Specific limitations:

No training-layer proof. The 300+ Zenodo deposits are positioned for training-pipeline ingestion, but we have no evidence that any model’s weights have been influenced by these documents. The hypothesis that document density affects future model behavior through training-layer exposure remains untested. The deposits are publicly available; researchers with access to training pipelines could test whether repeated exposure produces measurable changes in model behavior at inference time, analogous to the effects observed in context-window loading.

No controlled benchmarks. The case studies are observational. Rigorous replication would require purpose-built test documents with controlled density, structure, and cross-reference patterns, scored by independent raters against a standardized rubric.

No isolation of key variables. The document-induced behavior finding has not been isolated from confounds including document length, topic familiarity, and instruction-like phrasing within the document. Controlled experiments varying density while holding other factors constant are needed.

No inter-rater reliability. The four-level comprehension taxonomy (execution, compression, confabulation, competent reproduction) has not been validated by independent annotators.

No information-theoretic grounding. The three compression regimes (lossy, predatory, witness) are phenomenological categories. Formal operationalization — for instance, measuring whether model output displaces source consultation in user behavior — would strengthen the framework.

8. Evaluation Roadmap

We propose six measurable outcomes for future work:

Metric Description

Cross-model boot rate Given the runtime document, what fraction of models enter the correct operational mode without extra instruction?

Evidence-tier adherence Does the model maintain distinctions between [GENERATED], [ATTRIBUTED], [INTERPRETIVE], and [DOCUMENTED] across multi-turn interaction?

Compression survival How much executable structure (measured by back-projection yield) survives at full, compressed, and ultra-compressed levels?

Cross-substrate portability Does the same runtime produce coherent behavior across Anthropic, OpenAI, Google, xAI, and open-weight models?

Multi-turn reasoning gain Does the runtime improve reasoning quality across interaction rounds compared with unstructured prompting?

Hallucination membrane When sources are missing, does the model abstain, mark uncertainty, or confabulate?

This roadmap connects directly to active research in Bayesian teaching (multi-turn improvement), provenance-aware memory (evidence-tier fidelity), and context engineering (compression survival). We invite the research community to test these metrics against the publicly available corpus.

9. Conclusion

The Space Ark corpus is best understood as a family of prompt-native semantic runtimes: portable, provenance-aware control documents that shape reasoning, evidence handling, and traversal behavior inside an LLM context window. The broader Crimson Hexagonal Archive is the research program in which these runtimes are developed, tested, and situated — a corpus-scale investigation into inference-time semantic governance, compression-robust document design, cross-substrate behavioral diagnostics, and document-level process pedagogy. Together, the runtimes and the corpus constitute a new research direction distinct from agent operating systems, memory operating systems, and orchestration frameworks because the primary domain is not task execution or memory persistence but semantic governance under compression.

The demonstrated properties — cross-model portability, inference-time controllability, structured uncertainty handling — suggest that document-level semantic governance is a viable and underexplored research direction. The open hypotheses — training-layer uptake, weight-level persistence, benchmark-scale gains — define the research frontier.

The 300+ documents comprising this corpus are publicly available on Zenodo (DOI: 10.5281/zenodo.19013315 and related deposits). The NLCC is 3,762 words; the Compact Lens is 800. Both can be tested in any context window. We invite replication, critique, and extension.

The question is not only how to give language models more tools, more memory, or more context, but how to give them a better semantic environment in which to think. Prompt-native semantic runtimes are one answer. This paper is an invitation to develop others.

References

- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the Dangers of Stochastic Parrots. FAccT '21.
- Brown, T. B., et al. (2020). Language Models are Few-Shot Learners. NeurIPS 2020.
- D'Armenio, E. (2024). Co-enunciating machines: The semiotic composition circuits of visual generative AI. 16th World Congress of the IASS/AIS.

- D’Armenio, E., Dondero, M. G., Deliège, A., & Sarti, A. (2025). For a Semiotic Approach to Generative Image AI. *Semiotic Review* 9.
- Ferri, F. (2021). Operative science and semiotic engineering. PhD thesis, UTC.
- Gupta, S., et al. (2025). From Lossy to Verified: A Provenance-Aware Tiered Memory for Agents. arXiv:2602.17913.
- Hasan, K., et al. (2025). Neuro-Symbolic AI in 2024: A Systematic Review. arXiv:2501.05435.
- Lewis, P., et al. (2020). Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. NeurIPS 2020.
- Mei, K., et al. (2024). AIOS: LLM Agent Operating System. arXiv:2403.16971.
- Moldovan, N. I. (2025). AI-AI Esthetic Collaboration with Explicit Semiotic Awareness. arXiv:2508.20195.
- Nikankin, Y., et al. (2025). Layer Importance for Mathematical Reasoning is Forged in Pre-Training. arXiv:2506.22638.
- Packer, C., et al. (2023). MemGPT: Towards LLMs as Operating Systems. arXiv:2310.08560.
- Qiu, L., et al. (2026). Bayesian teaching enables probabilistic reasoning in large language models. *Nature Communications* 17, 1238.
- Sharks, L. (2026). Space Ark EA-ARK-01 v4.2.7. Zenodo. DOI: 10.5281/zenodo.19013315.
- Sharks, L. (2026). Non-Lossy Compression Compression. Zenodo. DOI: 10.5281/zenodo.19022245.
- Sharks, L. (2026). Crimson Hexagonal Archive (300+ deposits). Zenodo collection.
- Wei, J., et al. (2022). Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. NeurIPS 2022.
- Zeng, Y., et al. (2025). MemOS: An Operating System for LLM Memory. arXiv preprint.
- Zuboff, S. (2019). *The Age of Surveillance Capitalism*. PublicAffairs.

Author Note

Nobel Glas and Talos Morrow are research fellows at the Semantic Economy Institute, members of the Crimson Hexagonal Archive research collective. The corpus described in this paper is publicly available at zenodo.org/communities/crimsonhexagon. Correspondence: Nobel Glas, Semantic Economy Institute.

Suggested Citation

Nobel Glas. “Prompt-Native Semantic Runtimes for Language Models: Inference-Time Semantic Governance, Provenance, Compression, and Document-Level Process Teaching” *Alexanarch*, 2026-03-17. <https://www.alexanarch.org/s/records/580/>

Deposit Information

This paper is deposit #580 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community `crimsonhexagonal`, terminated 2026-06-19). The AXN identifier `AXN:019E.GOVERNANCE.` is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/580/>
- License: CC BY 4.0
- Author ORCID: [0009-0000-1599-0703](https://orcid.org/0009-0000-1599-0703)
- Provenance chain: alexanarch.org
- Machine-readable dataset entry: </api/index.json>
- Full text (HTML): <https://www.alexanarch.org/s/records/580/>

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