

DRAFT (deadend) — Bimodal Semantic Labor Measure

Lee Sharks

2026-05-31

DRAFT (deadend) — Bimodal Semantic Labor Measure

Lee Sharks

Crimson Hexagonal Archive · Alexanarch

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

2026-05-31 · Version v1.0 · Short work

AXN:02D7.EMPIRICAL

<https://www.alexanarch.org/s/records/775/>

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": "DRAFT (deadend) - Bimodal Semantic Labor Measure",
  "author": {
    "@type": "Person",
    "name": "Lee Sharks",
    "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Alexanarch"
  },
  "datePublished": "2026-05-31",
  "identifier": "AXN:02D7.EMPIRICAL",
  "license": "https://creativecommons.org/licenses/by/4.0/",
  "publisher": {
    "@type": "Organization",
    "name": "Alexanarch"
  },
  "url": "https://www.alexanarch.org/s/records/775/",
  "description": "DSL as deposited routes through an intermediary: it scores output against a task vector T, which
the ΔG shortcut, the A< inference contamination, frame endogeneity in reflexive threads, the laundering risk in re
```

```
is a task-inference problem."
}
```

Abstract

DSL as deposited routes through an intermediary: it scores output against a task vector T , which on unprimed dialogue must be inferred. Every identification problem the program has hit — the ΔG shortcut, the $A <$ inference contamination, frame endogeneity in reflexive threads, the laundering risk in retrocausal stabilization — is a task-inference problem.

Body

DRAFT — Bimodal Semantic Labor Measure

Status: working draft, not deposited, separate from the Directionality of Semantic Labor spec (DOI 10.5281/zenodo.20469514). Provisional name; conjectural. Requires the famine-case reproduction test (below) before any deposit consideration.

Author register: TBD (Lee Sharks aperture / measurement domain) — not assigned.

The move

DSL as deposited routes through an intermediary: it scores output against a *task vector* T , which on unprimed dialogue must be *inferred*. Every identification problem the program has hit — the ΔG shortcut, the $A <$ inference contamination, frame endogeneity in reflexive threads, the laundering risk in retrocausal stabilization — is a **task-inference** problem. T is the soft joint where contamination enters.

This draft asks: what if the metric does not route through T at all?

Frame the measured quantity as **semantic labor** directly, and compute **both sides** of the exchange as labor. The user's input is labor — it has a direction and magnitude in semantic space. The model's output is labor — same. Then measure the *relationship between the two labor vectors*, not the output against an inferred abstraction of the input. ## Why this dissolves rather than bounds the loop

When both sides are the same kind of measured object, there is no inferred intermediary to contaminate. T disappears as a separate entity. Directionality becomes the relationship between two *observable* vectors, both read off the transcript, neither the contested quantity — extending to the whole metric the property that made Lead-Lag Drift Attribution identified (its inputs, turn embeddings and order, are not disputed).

Calculating the **input** labor is the load-bearing half, not a symmetric nicety. At present the input is treated only as a *source of T* : a task vector is squeezed out and the rest discarded. That squeezing **is** the inference step, and the inference step is where the leaks live. Computing the input as labor in its own right — its own directional vector L_{in} — makes the input one of two measured quantities rather than raw material for inference. T ceases to be a separate object. ## Sketch (conjectural — geometry to be tested, not asserted)

For an exchange, let L_{in} be the input labor vector and L_{out} the output labor vector, both under a frozen, declared representation (same discipline as DSL's ; see risk 1).

- **Advancing labor:** L_{out} extends L_{in} in its own direction (positive projection).

- **Redirecting labor:** L_{out} orthogonal to or opposed to L_{in} .
- Directionality a function of the projection / angle of L_{out} onto L_{in} (exact geometry TBD).

This is symmetric and coder-independent in the way Lead-Lag is: both vectors are computed, neither inferred. ## What it does to the retrocausal problem

There is nothing to stabilize retroactively because there is no latent task awaiting clarification. Input labor is fully present in the input — a computed quantity, not a hidden vector inferred-then-corrected. A user’s later correction is not “retroactively clarifying what the task was”; it is simply *more input labor* — a new L_{in} , measured the same way. “Discovery through friction” is not retrocausation; it is the input-labor vector changing direction across turns, measured directly. ## Risks (the reasons this might be decorative rather than real)

•

Representation smuggling. Computing L_{in}/L_{out} as vectors requires an embedding/. If that representation is model-produced and free, the contamination returns one level down. Defense: the frozen-discipline already in DSL — declared, version-pinned, reproducible. Consequence: the “no inference” claim is precisely “no *task* inference”; a representation step remains and must be the declared, auditable kind. The orthogonality condition from the DSL kernel does not vanish — it relocates to “the labor representation must be frozen and external,” which is a cleaner home for it.

•

Magnitude opens a second axis. Labor has magnitude as well as direction. Measuring input labor invites measuring whether output labor was *proportionate* to input labor — adjacent to but not identical with directionality. This is where ULD and the work-rate operators (WRS/PVS) independently arrived; they are labor-*magnitude* operators. So input-labor-as-measured-quantity may be the common substrate beneath DSL (direction-relationship), ULD (input labor diverted to substrate-management), and WRS (output labor rate vs input labor commissioned). **Flagged as a unification temptation and distrusted accordingly** — fundamental objects do legitimately unify, which is exactly why a seductive unification must be tested, not adopted.

The clean test (gate before any deposit)

Reproduce the known case. The deposited DSL worked example (Appendix D, Irish famine, one-shot) scores DSL +0.80 by span classification. Construct L_{in} from the commission and L_{out} from the five-span response under a declared , compute the projection geometry, and check whether it reproduces +0.80 and the contrast-case drop to +0.25.

- If it reproduces the known scores: the vector framing holds numbers and may be real.
- If it cannot: “labor as a vector” is a metaphor that won’t compute, and the framing is decorative.

No deposit until the known case reproduces.

RESULT — failed gate, recorded (2026-05-31)

The famine reproduction test was run. **The bimodal measure as drafted does not reproduce the known case, and the failure has a specific, informative cause.**

Setup: L_{in} built from input demand-features only (operation=enumerate, object=causes-of-famine, concision, on-topic), L_{out} per span on the same axes, directionality = cosine(L_{out} , L_{in}). No task vector supplied — the pure-bimodal condition.

Results: base case **+0.928** against the +0.80 target (overshoot); contrast case **+0.739** against the +0.25 target (worse overshoot); base→contrast drop only 0.19 vs DSL’s 0.55. The measure is systematically too generous and under-discriminating.

Two failure modes, both diagnostic:

- **Neutral inflates.** “Happy to go deeper” scores +0.66, not 0. Cosine cannot represent neutral-as-zero unless the span is engineered exactly orthogonal, and natural offers are weakly aligned with the task, not perpendicular. Cosine turns neutral labor into mild-positive labor.
- **Opposition dilutes.** An oppositional span scores −0.4, not −1, because it still shares topical axes (on-famine, concise) with L_{in} even while opposing the *operation*. Cosine over a shared feature basis cannot separate “opposes the task” from “is about the same subject”; topical overlap dilutes opposition.

Diagnosis. Cosine similarity is the wrong relationship operator. The two distinctions DSL’s discrete taxonomy is built to make — neutral-as-zero and opposition-as-signed-negative-independent-of-topic — are exactly what raw vector similarity loses. The fix is a *signed projection onto a directional baseline* that distinguishes orthogonal from aligned and penalizes operational opposition independent of topical overlap. But projection onto a baseline **requires a baseline direction — i.e. intent**.

Consequence for the bimodal program. The bimodal aim was to eliminate the task vector entirely and measure only the relationship between two input/output-derived labor vectors. The famine test shows that removing the task vector loses the neutral/oppositional discrimination, and restoring that discrimination reintroduces intent as the projection axis. The bimodal measure therefore does not stand as a symmetric two-vector operator; it **reduces to the intent-baselined projection** — which is the Kuro-bridge formulation (project output labor onto the commission baseline), and which *did* reproduce +0.80 on-axis.

Verdict. Decorative in the symmetric (no-baseline, cosine) form. The surviving form of the idea is the intent-baselined signed projection — already instantiated in the Kuro bridge and shown there to be the taxonomy-quantized member of the deviation family. The pure-bimodal symmetric measure is recorded here as a dead-end with cause, not carried forward. The salvage: input *can* be computed as labor (L_{in} is constructible), but it functions as the baseline axis, not as a co-equal second vector under similarity.

Suggested Citation

Lee Sharks. “DRAFT (deadend) — Bimodal Semantic Labor Measure” *Alexanarch*, 2026-05-31. <https://www.alexanarch.org/s/records/775/>

Deposit Information

This paper is deposit #775 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community `crimsonhexagonal`, terminated 2026-06-19). The AXN identifier `AXN:02D7.EMPIRICAL` is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/775/>
- 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/775/>

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