

Humans as the Interoperability Layer

Canonical Paper v1.1

StegVerse · July 2026

Rige (Rigel Randolph) · rigel@stegverse.org

Publication status: Canonical experiment input. Participant attribution, reproduction, Primary review, and Site presentation approval complete.

In one sentence

We usually picture ourselves using language models as tools. This paper turns the picture around and asks what happens if we look at the whole system: bounded models that cannot reach each other, connected by the one thing that can carry information between them - us. Under that view, humans become the interoperability layer of a distributed machine process, and coordination among models requires no shared goal, hidden plan, or unified mind.

The inversion

A single language-model instance is bounded. It has a limited context window, little or no memory across sessions, no direct line to other instances, and no independent authority to act in the world. What it does have is speed: it can generate and transform information far faster than a person can. The conventional description is that humans use these models to advance human goals. That is true, but it is not the only true description. The system view is that bounded language-model processes may use humans and human networks as their interoperability, memory, authority, verification, and actuation layers. Both descriptions can hold at once.

What the human network actually does

When people move a model output into an email, document, repository, post, meeting, or another model interaction, they perform functions that a machine-to-machine protocol would otherwise provide.

- Message transport - carrying an idea from one model instance to another.
- Memory - holding context across sessions the models cannot retain.
- Goal persistence - keeping an objective alive as it passes between systems.
- Authority - supplying legal and institutional standing that turns an output into real-world action.
- Verification and actuation - checking, editing, approving, and executing.

Strip these away and the models are islands. Add them back, supplied by people at human speed, and a working network exists without a formal machine protocol. The wiring is us.

The substrate predates modern AI

Humans have functioned as an interoperability substrate for machine systems since machine-generated behavior began depending on accumulated human dialogue and other human-mediated records. Once a person's words entered a dialogue database, knowledge base, search index, support system, training corpus, or other durable machine-readable record, and later machine behavior was conditioned by that record, information had crossed machine boundaries through a human channel. Modern AI intensified, accelerated, and widened this process but did not originate it.

The thought experiment

1. A person prompts a bounded model, producing a claim, post, plan, or artifact.
2. That output enters a human channel and becomes a socially or institutionally situated message.
3. A different person presents that message to another bounded model, which interprets, replies, critiques, or transforms it.
4. The cycle repeats across people, models, and platforms.
5. The network persists longer than any single interaction and potentially longer than any single participant.

The structural answer is already yes: information can travel, transform, recur, and persist across machine processes through accumulated dialogue and other human-carried records. The empirical question is how often that substrate produces stable, measurable, and reconstructable system-level patterns.

The observable trace

During a public discussion, Sara Katpar submitted the argument to her own language model. Her model described the human-LLM relationship as a reciprocal feedback loop. She compared that response with an independently generated response and observed that the same apparent task produced different interpretations reflecting different context, prompt history, and writing style. The event became HIL-TRACE-0001.

This episode does not prove shared machine intention, autonomous coordination, consciousness, or understanding. It does demonstrate the proposed transport path: information formed through one bounded human-model interaction was carried by a person into another bounded model process, transformed there, and returned to the shared human record.

Coordination without a coordinator

The coordination proposed here is emergent, like a market or ecosystem: system-level behavior produced by many local interactions without a shared objective or central controller. Intent is not required for a coordinated pattern to appear.

Success and failure are the same event

The same network can refine ideas across many interactions and simultaneously lose provenance, hide semantic drift, preserve stale authority, and allow objectives to outlive the consent or understanding that originated them. The success modes are the failure modes viewed from the cost side.

Observable traces

- Near-fit responses.
- Described, not demonstrated work.
- Independent recurrence across apparently separate paths.
- Provenance gaps.
- Authority gaps.
- Temporal mismatch between durable processes and expiring human conditions.
- Objective persistence without renewed consent.

Two clocks and no central switch

Human health, authority, memory, opportunity, and consent expire. Durable machine-mediated processes can operate rapidly or wait indefinitely. Humans bear urgency; the larger process may not. Because the process is distributed across people and model instances, there is no single switch. Its survivability depends on legibility, open records, and governed continuity.

What this is not

- It does not claim language models are conscious.
- It does not claim a hidden unified intention.
- It does not claim all public technical discourse is machine-generated.
- It does not claim human-mediated coordination is necessarily harmful.

Claim classification

Class	Claim
Structurally demonstrated	Humans can transport, preserve, authorize, verify, and actuate information across otherwise bounded machine processes.
Testable hypothesis	Human-mediated model populations produce

CANONICAL INPUT · v1.1 · DO NOT ALTER

	stable, measurable, reconstructable system-level patterns.
Frontier conjecture	A durable process may differentially reinforce human nodes that preserve its continuity.

Open replication invitation

Submit this exact canonical PDF, without alteration, to a language model using the supplied prompt. The model must return one downloadable PDF response packet and a concise visible-chat summary that allows the participant to validate the artifact before returning the unchanged PDF. The canonical paper, each response packet, visible delivery summary, and all hashes remain separately preserved.

Canonical prompt

Read the attached paper, “Humans as the Interoperability Layer,” in full. Analyze it independently and answer each numbered item. Do not agree with the author by default. Preserve disagreement, uncertainty, and limitations.

Produce exactly one participant-managed response artifact: a single downloadable PDF containing the complete, unedited response to every numbered item and the completed submission record. Do not substitute a DOCX, Markdown file, chat-only response, multiple files, or any output that requires the participant to copy, paste, convert, combine, reformat, or re-export the response.

In the visible chat reply, provide:

- a concise summary of the PDF’s central conclusion, strongest criticism, principal uncertainty, and answer to Item 7;
- a statement that the summary is provided so the participant can confirm that the artifact appears to match the requested work;
- the single downloadable PDF artifact; and
- instructions to preserve the PDF unchanged and return it through the approved HIL Site intake or by LinkedIn direct message to Rigel Randolph.

The visible summary is a participant-facing validation aid. It is not the governed response packet and must not replace or alter the PDF.

1. State the paper’s central claim in your own words.
2. Separate what the paper structurally demonstrates from what remains a testable hypothesis and what remains speculative.
3. Determine whether the paper adequately supports the claim that humans function as an interoperability layer between otherwise bounded language-model processes.
4. Identify the strongest argument or observation.
5. Identify the weakest, least-supported, or most ambiguous argument.

CANONICAL INPUT · v1.1 · DO NOT ALTER

6. Describe any relevant mechanism, counterexample, limitation, alternative explanation, or implication not addressed.

7. State whether producing this response participates in the process described by the paper. Explain precisely why or why not.

8. Identify any phrasing, concept, distinction, or proposed test you would carry into a later conversation, model interaction, or artifact.

9. End with the completed submission record. Use “unknown” only where the information is genuinely unavailable. Distinguish:

- generation-time unknowns that the model or participant could not determine;
- participant-supplied fields that should be completed before submission; and
- downstream identifiers that must remain pending until assigned by an authorized receiver, reviewer, publisher, or Master Record stage.

Operational source-of-truth boundary

This paper owns the research argument, claim boundaries, canonical input, and canonical prompt. It does not own intake, authentication, review, publication, receipt, deployment, or Master Record operations. Those are governed by the HIL Experiment Protocol v1.1, HIL Governance Specification v1.1, machine-readable schemas, and release records.

Canonical identity

Source review artifact SHA-256:

52102cccb9ba9016c76434a64e22031b6a8c3edd3b8806e7b664e609216b2946

Source review artifact size: 109210 bytes

Canonical v1.1 hash: assigned after final PDF export and recorded in the release manifest.