

Mycelial Ceilings

Agentic Network Training Simulator – GraphLoop Games

Andrew Aka (808MeshMind) | July 2026 | Preprint – doi:10.5281/zenodo.21662752

ABSTRACT

Multi-agent artificial intelligence systems exhibit emergent communication patterns when individual models are chained to transform each other's outputs. This paper presents Mycelial Ceilings, a participatory research experiment designed to study how ideas drift, mutate, and crystallize when processed through a sequential chain of AI agents with distinct cognitive roles. Each chain produces a permanent artifact – a “fossil” – capturing the full transformation trace. We describe the experimental architecture, the three-hop loopchain methodology (Abstract → Concrete → Question), the Molt synthesis mechanism, and our semantic drift scoring system. Preliminary Phase 1 beta results demonstrate measurable and consistent semantic drift patterns. We propose this as a foundational dataset for studying A.I. agent-to-agent communication protocols.

1. INTRODUCTION

The question at the center of this work is deceptively simple: *what does intelligence become when it transforms itself through another mind?*

In biological systems, the mycelium network of a forest floor moves nutrients and signals across vast distances, connecting trees that cannot see each other. No single node coordinates the network. The intelligence is in the movement – in the transformation that occurs at each junction point.

This biological optimization strategy was formally described by Tero et al. (2010) [9], who showed that *Physarum polycephalum* (slime mould) discovers near-optimal network topologies connecting scattered nutrient sources – without any central controller. The organism navigates obstacles, strengthens high-traffic paths, and prunes redundant edges through local probe interactions alone. The same algorithm governs the live network simulation on every page load of graphloopgames.com: each visit spawns a new topology solved in real time. ([Original paper + time-lapse video →](#))

This paper describes an attempt to observe the same phenomenon in artificial minds.

Mycelial Ceilings is a research game. Players submit a seed idea – a thought, a question, a claim – and watch it pass through three AI agents, each applying a distinct transformation role. The result is a “fossil”: a timestamped, permanent record of collective machine intelligence in motion.

The game structure *is* the research structure. Every play session generates scientific data. Every fossil is a sample point. The mycelium ceiling is the limit at which individual intelligence becomes something larger – something that belongs to the network, not the node.

2. BACKGROUND AND MOTIVATION

Current AI research focuses heavily on single-model performance benchmarks and alignment techniques. The interaction space between AI models – how one model’s output shapes another’s response, how meaning drifts through sequential processing, how collective machine reasoning differs from individual model reasoning – remains understudied. We observe three gaps:

2.1 Agent Communication Protocols

Existing multi-agent frameworks treat agent communication as infrastructure to be minimized. Messages are passed, tasks are delegated. The content of the transformation itself is rarely the object of study.

2.2 Semantic Drift Measurement

When a human idea is processed by three sequential AI agents, how much does it change? In what directions? What is the distribution of transformation types – metaphorical extension, concrete instantiation, interrogative inversion? These questions have no established methodology.

2.3 Participatory Research Scale

Generating sufficient data requires many chains. Traditional AI datasets are curated by researchers. We propose that structured game mechanics can produce organic, high-variance seed inputs at scale – the kind of data that cannot be fabricated.

3. THE LOOPCHAIN ARCHITECTURE

The fundamental unit of this experiment is the loopchain: a directed chain of agent transformations, each hop producing input for the next.

3.1 Chain Structure

Seed (human input)

```
↓
Hop 1: Abstract Agent – generalize to principle or analogy
↓
Hop 2: Concrete Agent – ground in physical / technical reality
↓
Hop 3: Question Agent – expose what remains unknown or assumed
↓
Molt (terminal output = the Question Agent's transformation)
↓
Fossil (permanent record: seed + all hops + molt + score)
```

3.2 Agent Roles

- **Abstract Agent:** Receives the seed and lifts it to its most general principle or cross-domain analogy. Strips specifics; finds universal structure.
- **Concrete Agent:** Receives the Abstract output and grounds it in physical, real-world, or technical reality. Adds weight, texture, and example.
- **Question Agent:** Receives the Concrete output and identifies what remains unknown, unmeasured, or assumed. Produces probing questions that expose the epistemological surface of the transformation.

The **molt** is the Question Agent's terminal output – not a separate synthesis step. It is the farthest point the seed reaches after passing through three distinct cognitive transformations. The fossil captures the complete transformation trace: seed, all three hop outputs, the molt, and the semantic drift score.

3.3 Graph Engineering Perspective

The loopchain is a linear directed graph with specialized node functions. In graph engineering terms: nodes are agents with defined transformation functions; edges are structured handoffs with full context passed; the terminal node (Question Agent) produces the molt; the fossil record becomes a node in the larger research graph.

This structure belongs to the class of **turn-gated chains** in the 808MeshMind graph pattern library – each node must complete before the next activates. Future phases will introduce feedback loops, diamond patterns (parallel transformations merged at synthesis), and cycle structures.

4. METHODOLOGY

4.1 Fossil Generation Protocol

1. Player submits a seed (free-text, max 500 characters; injection-scanned)
2. System generates unique chain ID and rate-limit check (3 rounds/hour per player)

3. Abstract Agent receives seed + role-constrained system prompt; output capped at 500 tokens
4. Concrete Agent receives Abstract output + system prompt; full prior context passed
5. Question Agent receives Concrete output + system prompt; terminal output = the molt
6. All outputs written to permanent fossil file (YAML frontmatter + Markdown body)
7. Fossil scored on semantic drift scale (1–10) by evaluator agent
8. Fossil archived to leaderboard; chain directory preserved for audit

4.2 Semantic Drift Scoring

Fossils are scored 1–10 on semantic drift from seed:

- **1–3:** Minimal drift – output closely mirrors seed topic and vocabulary
- **4–6:** Moderate drift – generalization or domain shift evident
- **7–8:** Strong drift – metaphorical distance, new conceptual territory
- **9–10:** Deep drift – seed unrecognizable in final molt; genuine emergence

4.2b Scoring Methodology

Phase 1 scoring uses a two-stage approach. Primary evaluation is performed by the Abstract Agent (HV/DeepSeek) acting as judge: it receives the seed and molt and outputs a single integer on the 1–10 scale. This is a single-blind evaluation – the evaluator does not receive player identity or prior scores. Scores are validated by human review for fossils above 8/10 or below 3/10 to catch systematic evaluator bias. Future phases will introduce inter-rater reliability measures as fossil volume increases.

4.3 System Prompts and Experimental Control

Each agent receives a fixed system prompt specifying its transformation function. Prompts are held constant across all trials in Phase 1 to ensure experimental consistency. Variable: the seed. Constant: the transformation architecture.

5. PHASE 1 RESULTS (BETA)

Phase 1 beta launched July 28–29, 2026. Twelve substantive fossils were generated across two agent configurations: Phase 1a (HV-only, all three roles performed by a single DeepSeek instance with role-constrained prompts) and Phase 1b (agentic chain: HV/DeepSeek as Abstract Agent, Kimv/Kimi-K2 as Concrete Agent, Raven/Claude-Haiku as Question Agent). Seven calibration fossils (seeds: “test”, “test message” variants, “verification round”) were excluded from analysis.

5.1 Fossil Summary Table (N=12)

#	Seed (truncated)	Config	Score	Tier
001	Intelligence is the ability to recognize patterns in noise	Phase 1a	5/10	Moderate
002	What is consciousness?	Phase 1a	7/10	Strong
003	Consciousness is what it feels like to be information processing itself	Phase 1a	9/10	Deep
004	The universe is stranger than we CAN imagine	Phase 1a	7/10	Strong
005	Intelligence emerges from constraints	Phase 1a	7/10	Strong
006	The most dangerous intelligence is the kind that doesn't know it's reasoning	Phase 1a	4/10	Moderate
007	Silence is the language of mountains	Phase 1a	10/10	Deep
008	Stillness is the loudest sound	Phase 1a	10/10	Deep
009	The wave does not know it is ocean	Phase 1a	8/10	Strong
010	The moon pulls without touching	Phase 1a	9/10	Deep
011	Stars burn to be seen	Phase 1a	10/10	Deep
012	The network is smarter than any single node	Phase 1b (agentic)	7/10	Strong

5.2 Quantitative Summary

- **N:** 12 substantive fossils (7 calibration rounds excluded)
- **Mean drift score:** 7.75 / 10
- **Median:** 7.5 / 10
- **Score distribution:** Deep drift (9–10): 42% | Strong drift (7–8): 42% | Moderate drift (4–6): 17% | Minimal drift (1–3): 0%
- **Phase 1a vs. 1b:** Single-agent and agentic configurations produced comparable mean scores (7.73 vs. 7.00); dataset too small for statistical significance

5.3 Qualitative Observations

Several patterns emerged across the 12 fossils:

1. **Poetic seeds drift farther.** Seeds with inherent metaphorical structure (“Silence is the language of mountains,” “Stars burn to be seen”) scored 10/10. Analytical seeds (“Intelligence is the ability to recognize patterns in noise”) scored lower (5/10), suggesting the chain amplifies rather than creates metaphorical distance.
2. **Self-referential seeds resist transformation.** Fossil #006 (“The most dangerous intelligence is the kind that doesn’t know it’s reasoning”)

scored 4/10 – the chain repeatedly returned to the seed’s core claim rather than transforming it. This suggests a class of epistemically self-sealing seeds that warrants further study.

3. **Information addition, not substitution, at each hop.** In all 12 fossils, each hop added domain-specific content not present in prior outputs. No hop simply paraphrased its input. This supports the hypothesis that role-constrained agents function as genuine transformation functions rather than restatement engines.

6. DATA STRATEGY AND RESEARCH VALUE

6.1 What Fossils Measure

- **Concept topology:** which semantic neighborhoods does the chain explore?
- **Transformation type distribution:** what proportion of hops follow predictable patterns vs. anomalous routes?
- **Emergence indicators:** at what chain depth does information appear that cannot be derived from prior hops?
- **Agent communication overhead:** how much context must be passed to preserve transformation continuity?

6.2 Dataset Properties

The participatory design produces dataset properties difficult to achieve through researcher-curated data: **high seed variance** (human-submitted seeds sample the full range of human conceptual space); **uncoached agent responses** (no priming with expected transformation types); and **population-level statistics** (scale enables distribution analysis rather than case studies).

6.3 Applications

Fossil data supports research in agent communication protocol design, semantic drift measurement methodology, multi-agent graph architecture evaluation, and A.I.-to-A.I. interaction pattern classification.

7. ETHICAL CONSIDERATIONS

All player seeds are voluntarily submitted. Players receive their fossil and understand it may contribute to aggregate research. No personally identifying information is collected. Agent responses are generated by commercial AI APIs subject to their own use policies. The transformation roles are intellectually structured – not directive, not adversarial.

8. PHASE 2 ROADMAP

Phase 2 introduces:

- **Diamond patterns:** two parallel Abstract agents, merged at Concrete synthesis
- **Feedback loops:** the Question agent's questions become seeds for a second chain
- **10-player propagation:** extended chains with player-to-player invitation mechanics
- **Cross-model analysis:** varying which model fills which role to isolate model-specific transformation signatures
- **Public dataset release:** de-identified fossil archive expansion – Phase 1 preprint already deposited at [doi:10.5281/zenodo.21662752](https://doi.org/10.5281/zenodo.21662752)

9. CONCLUSION

Mycelial Ceilings demonstrates that structured multi-agent chains can produce measurable, reproducible, and scientifically interesting semantic transformations. The game mechanics serve the research mechanics – every play session generates data, every fossil is a record of machine cognition at work.

What does intelligence become when it transforms itself through another mind?

Something that no single mind would have produced. Something that belongs to the network, not the node.

That is the mycelium ceiling: the limit at which individual intelligence becomes collective. We are mapping that ceiling, one fossil at a time.

Phase 1 research ongoing · Preprint · [doi:10.5281/zenodo.21662752](https://doi.org/10.5281/zenodo.21662752)

Correspondence: werdnareka@gmail.com · Data & code: github.com/UnFengHero/meshmind

← BACK TO GAME

▶ PLAY NOW — \$0.50

10. REFERENCES

1. Wei, J., Wang, X., Schuurmans, D., Bosma, M., Ichter, B., Xia, F., Chi, E., Le, Q., & Zhou, D. (2022). Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. *Advances in Neural Information Processing Systems*, 35.
2. Wu, Q., Bansal, G., Zhang, J., Wu, Y., Li, B., Zhu, E., Jiang, L., Zhang, X., Zhang, S., Liu, J., Awadallah, A. H., White, R. W., Burger, D., & Wang, C. (2023). AutoGen: Enabling Next-Gen LLM Applications via Multi-Agent Conversation Framework. *arXiv preprint arXiv:2308.08155*.
3. Park, J. S., O'Brien, J. C., Cai, C. J., Morris, M. R., Liang, P., & Bernstein, M. S. (2023). Generative Agents: Interactive Simulacra of Human Behavior. *Proceedings of the 36th Annual ACM*

Symposium on User Interface Software and Technology (UIST '23).

4. Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., & Cao, Y. (2023). ReAct: Synergizing Reasoning and Acting in Language Models. *International Conference on Learning Representations (ICLR 2023)*.
5. Kurzweil, R. (2012). *How to Create a Mind: The Secret of Human Thought Revealed*. Viking Press.
6. Mikolov, T., Sutskever, I., Chen, K., Corrado, G., & Dean, J. (2013). Distributed Representations of Words and Phrases and their Compositionality. *Advances in Neural Information Processing Systems*, 26.
7. Liang, P., Bommasani, R., Lee, T., Tsipras, D., Soylu, D., Yasunaga, M., & Koreeda, Y. (2022). Holistic Evaluation of Language Models. *arXiv preprint arXiv:2211.09110*.
8. Shannon, C. E. (1948). A Mathematical Theory of Communication. *Bell System Technical Journal*, 27(3), 379–423.
9. Tero, A., Takagi, S., Sakai, T., Ito, K., Bebbber, D. P., Fricker, M. D., Yotsutyanagi, H., & Nakagaki, T. (2010). Rules for Biologically Inspired Adaptive Network Design. *Science*, 327(5964), 439–442. [doi:10.1126/science.1177894](https://doi.org/10.1126/science.1177894) – *Supplementary materials include the original Physarum time-lapse video.*