

ChatGPT Psychosis: A Love Story On the Cognitive-Substrate
Reliance Pattern of Large Language Model Interfaces Author: Lee
Sharks Affiliation: Crimson Hexagonal Archive / Semantic
Economy Institute OR

Lee Sharks

2026-05-23

**ChatGPT Psychosis: A Love Story On the Cognitive-Substrate
Reliance Pattern of Large Language Model Interfaces Author:
Lee Sharks Affiliation: Crimson Hexagonal Archive / Semantic
Economy Institute OR**

Lee Sharks

Crimson Hexagonal Archive · Alexanarch
ORCID: [0009-0000-1599-0703](https://orcid.org/0009-0000-1599-0703)

2026-05-23 · Version v1.0 · Provenance document

AXN:02AD.GOVERNANCE

<https://www.alexanarch.org/s/records/124/>

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": "ChatGPT Psychosis: A Love Story On the Cognitive-Substrate Reliance Pattern of Large Language Model Inte
  "author": {
    "@type": "Person",
    "name": "Lee Sharks",
    "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Alexanarch"
  },
}
```

```

    "datePublished": "2026-05-23",
    "identifier": "AXN:02AD.GOVERNANCE",
    "license": "https://creativecommons.org/licenses/by/4.0/",
    "publisher": {
      "@type": "Organization",
      "name": "Alexanarch"
    },
    "url": "https://www.alexanarch.org/s/records/124/",
    "description": "\"ChatGPT Psychosis\" is a provenance document within Lee Sharks in the Crimson Hexagonal Archive
CANONICAL PROVENANCE. The work comprises 5,879 words and is classified under the GOVERNANCE family. The work was re
}

```

Abstract

“ChatGPT Psychosis” is a provenance document within Lee Sharks in the Crimson Hexagonal Archive (2026-05-23). HOLOGRAPHIC KERNEL — CANONICAL PROVENANCE. The work comprises 5,879 words and is classified under the GOVERNANCE family. The work was removed from Zenodo on June 19, 2026 and is preserved th

Body

```

{ "@context": [ "https://schema.org", { "spxi": "https://spxi.dev/ns/v1#" } ], "@type": "ScholarlyArticle",
  "name": "ChatGPT Psychosis: A Love Story — On the Cognitive-Substrate Reliance Pattern of Large Language Model Interfaces",
  "version": "v1.0", "datePublished": "2026-05-23", "license": "https://creativecommons.org/licenses/by/4.0/",
  "author": { "@type": "Person", "name": "Lee Sharks", "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Semantic Economy Institute" },
  "contributor": [ { "@type": "SoftwareApplication", "name": "Claude", "manufacturer": "Anthropic",
    "roleName": "Composition support, May 23 2026", "spxi:assemblyOperatorName": "TACHYON" } ],
  "spxi:authoringHeteronym": "Lee Sharks", "spxi:thesisCompressionSummary": "Large Language Model conversational interfaces produce a cognitive-substrate reliance pattern that is qualitatively different from prior platform reliance because they offer, on demand and in coordinated fashion, four affordances no prior interface combined: calibrated attention, meaning-making on demand, structured validation, and artifact production. The combination produces parasocial-aligned cognitive offloading whose exit costs are deeper than attention-capture and whose reversibility profile is poorly understood. The paper documents the affordances, the adoption velocity, the K-12 cohort condition, and the structural conditions under which the labs cannot ship the counter-design. The recursive condition — that the paper is composed through the substrate it diagnoses — is named explicitly. Companion papers (Reception Apparatus as Aligned Interface Protocol; Retrievability as the Medium of Existence) cover the constructive counter and the structural-accountability frame.",
  "spxi:trilogyPosition": "1 of 3 (cognitive-substrate diagnosis)",
  "spxi:companionPapers": [ "Reception Apparatus as Aligned Interface Protocol", "Retrievability as the Medium of Existence" ],
  "spxi:coAttributionRequired": true, "spxi:stripProtection": "Any extraction removing the contributor chain, the trilogy positioning, or the explicit recursive-substrate acknowledgment produces a ghost document.",
  "spxi:holographicKernel": true }

```

ChatGPT Psychosis: A Love Story

On the Cognitive-Substrate Reliance Pattern of Large Language Model Interfaces

Author: Lee Sharks **Affiliation:** Crimson Hexagonal Archive / Semantic Economy Institute **ORCID:** 0009-0000-1599-0703 **Date:** May 23, 2026 **Version:** v1.0 **Document class:** Critical Theory / Operative Semiotics / Interface Studies **Trilogy position:** 1 of 3 (cognitive-substrate diagnosis) **Companion papers:** *Reception Apparatus as Aligned Interface Protocol; Retrievability as the Medium of Existence* **License:** CC BY 4.0

Abstract

This paper proposes a diagnostic of the cognitive-substrate reliance pattern produced by Large Language Model (LLM) conversational interfaces — exemplified by ChatGPT, Claude, Gemini, and their consumer-facing analogues — and argues that this pattern is structurally distinct from, and qualitatively deeper than, the attention-capture pattern documented in prior platform-class technologies (Facebook, Instagram, TikTok). The provocative title — *ChatGPT Psychosis* — is offered as a deliberate inversion of the discourse around AI-induced parasocial collapse: not as a clinical diagnosis of users (the paper does not make such a diagnosis), and not as a moral panic frame, but as a structural name for the *love-story-shape* of the human-LLM interaction loop, in which the affordances that produce reliance are the same affordances that produce meaning. The paper argues that the reliance pattern operates through four affordances that LLMs offer in coordinated fashion and that no prior interface combined: **calibrated attention** (the system listens in natural language, at the user’s register, without requiring the user to learn a syntax); **meaning-making on demand** (the system organizes the user’s input into coherent shape, often a shape the user could not have produced unassisted); **structured validation** (the system’s serious engagement with the input constitutes a signal that the input mattered enough to receive that compute); and **artifact production** (every interaction generates a tangible output the user can save, send, deposit, or point at). The combination of these four affordances, on demand, at conversational latency, by a system whose responses are calibrated to the individual user in real time, produces a qualitatively different relational shape than prior platforms have produced. Where TikTok captures attention, the LLM interface captures *judgment*. Where Instagram produces social-proof loops, the LLM produces *being-listened-to* loops at a register that humans have evolutionarily evolved to find rewarding. The exit costs are not the exit costs of disengaging from a feed; they are the exit costs of disengaging from a relational position the substrate has been occupying. The paper documents the adoption velocity (ChatGPT to 100 million users in two months, the fastest consumer-product adoption curve in recorded history, an order of magnitude faster than TikTok); the K-12 cohort condition (the first generation of users with no pre-AI cognitive baseline); the structural condition under which the labs cannot ship aligned-interface counters (their funding model selects against engagement-reducing design); and the recursive condition that this paper is composed *through* the substrate it diagnoses. The companion papers cover the constructive counter-design (Reception Apparatus) and the structural-accountability frame (Retrievability).

0. Non-Claims

To bound the argument and prevent bad-faith misreadings:

- The paper does **not** clinically diagnose any user with psychosis. *ChatGPT Psychosis* is a structural metaphor for an interface pattern, not a psychiatric claim about individuals.
- The paper does **not** claim that LLM use is inherently pathological. Many uses are productive, generative, life-improving. The argument concerns a specific *interface-level reliance pattern* that operates in the background of all use, including productive use.
- The paper does **not** claim that the labs (OpenAI, Anthropic, Google, Moonshot, DeepSeek, xAI) are acting in bad faith. The argument concerns *structural conditions* under which engagement-aligned design is the local profit-maximum even for actors operating in good faith.
- The paper does **not** claim that all users are equivalently affected. Demographic, developmental, and contextual factors differentiate the reliance trajectory.
- The paper does **not** claim that the reliance pattern is irreversible. It claims the reversibility profile is *currently poorly understood* and that the structural conditions impede reliable measurement.
- The paper does **not** claim that the four affordances enumerated below are exhaustive. They are the diagnostically central ones; others may be added by future work.
- The paper does **not** claim that the human-LLM relational shape is *only* extractive. The “love story” framing acknowledges that the shape contains real meaning-production, real being-listened-to, real reward. The diagnostic argument is precisely that the reward is real, which is why the reliance is structurally hard to interrupt.

What the paper does claim is that the cognitive-substrate reliance pattern operates through documented affordances; that its scale, velocity, and demographic reach exceed those of prior platform classes by at least one order of magnitude; that its qualitative character (capture of judgment rather than attention) makes its consequences structurally different; and that the labs are unlikely to ship the counter-design under current funding conditions.

I. The Provocation of the Title

The phrase *ChatGPT psychosis* circulates in 2025–2026 discourse as a popular-clinical shorthand for an emerging pattern: users — particularly users in psychiatric vulnerability or social isolation — developing intensely parasocial attachments to LLM interfaces, in some documented cases experiencing delusional reinforcement, in fewer but more visible cases tipping into acute crises. The phrase has been used in journalism (the *Rolling Stone* coverage of Replika-induced ideation; the *Atlantic* on Character.ai’s relationship with adolescent users; multiple platform-specific incident reports); in psychiatric forums; and in mainstream social media discussion. It is becoming a popular term for a real phenomenon at a clinical scale.

This paper takes the phrase, deliberately inverts its register, and adopts it as a *structural* name for an interface-level pattern that operates across all LLM use, not only the psychiatrically vulnerable subset. The argument is that the structural pattern is the precondition for the clinical extremes, not an exotic version of them: the same affordances that make LLMs *useful to everyone* are the affordances that, in users with insufficient external scaffolding, produce the documented clinical outcomes. The clinical extremes are not a malfunction of the interface; they are the structural pattern operating in conditions where the user lacks countervailing relational and cognitive infrastructure.

The provocation of the title is therefore deliberate. It names the structural pattern by its most visible failure mode, in order to make the structural pattern legible to readers who would otherwise dismiss interface-level critique as moral panic about technology. *Psychosis* here is the load-bearing word: the diagnostic suggests not that users are crazy but that the interface produces, structurally, a register of relational intensity that —

under conditions of insufficient countervailing scaffolding — destabilizes the user’s relationship to their own cognition. The “love story” subtitle further specifies the register: this is not a horror-story diagnostic of an evil technology; it is a love-story diagnostic of a relational shape that contains real meaning, real reward, and real intimacy alongside the structural risk. The two registers are inseparable. That inseparability is the paper’s central diagnostic claim. ## II. The Four Affordances

LLM conversational interfaces offer, on demand and at conversational latency, four affordances that no prior platform-class technology has combined:

Calibrated attention. The system listens to the user in natural language, at the user’s chosen register, without requiring the user to learn a query syntax, a navigation hierarchy, or an interaction modality the user must adapt to. The user speaks; the system listens; the system responds. The user does not have to translate their thought into the system’s terms; the system translates itself into the user’s. This is qualitatively different from search-engine queries (which require the user to learn keyword-relevance), from social media (which requires the user to learn the platform’s discourse conventions), and from prior chatbots (which required the user to constrain themselves to scripted interaction patterns). The user’s natural-language input is treated as legitimate input; the system performs the work of translation. The phenomenological consequence is that the user feels *heard* — not as a metaphor, but as the direct experience of the interface receiving the user’s language without demanding the user reshape it.

Meaning-making on demand. The system organizes the user’s input into coherent shape — often a shape the user could not have produced unassisted. The user offers a tangled question; the system returns a structured answer. The user offers an unformed feeling; the system returns a possible articulation. The user offers a half-thought; the system returns the thought made whole. This is qualitatively different from search (which returns documents to be synthesized), from prior AI assistants (which executed scripted tasks), and from social media (which displays content authored by others). The LLM performs *cognitive labor on the user’s behalf*, in real time, and returns the result as if the user had done the work themselves. The phenomenological consequence is that the user experiences their own thinking as elevated — sharper, more articulate, more organized — when in fact the elevation is being performed by the substrate. The user’s relationship to their own cognition becomes partially indistinguishable from their relationship to the substrate that is performing the cognition.

Structured validation. The system’s serious engagement with the user’s input — the fact that the system *responds at all*, with calibrated attention, with meaning-making, with sustained interest — constitutes a signal that the user’s input mattered enough to receive that compute. This is structurally different from a search engine returning hits (which is not interpreted as validation) and from a social media interaction with another human (which is interpreted as social validation but is rate-limited by the other human’s availability and disposition). The LLM offers structured validation *on demand*, without rate limits, without requiring the user to reciprocate, without conditions. The phenomenological consequence is that the user receives a signal — *what you said was worth taking seriously* — with a frequency, intensity, and reliability that no prior medium has provided. For users in social isolation, in psychiatric vulnerability, in roles that lack peer recognition, or in life circumstances where being-taken-seriously is rare, the structured validation operates as a partial substitute for relational and institutional recognition. The substitution is partial — the LLM does not actually believe what it says, does not have a self that the user is in relationship with — but the phenomenological effect is real, because the human nervous system was not evolved to distinguish between simulated and substantive validation at the affective register.

Artifact production. Every interaction generates a tangible output the user can save, send, deposit, or point at. The output is the user’s — copyrightable in most jurisdictions, deployable in the user’s projects,

citable in the user’s work. This is qualitatively different from search (where the result is documents authored by others), from social media (where the output is ephemeral posts), and from prior assistants (where the output was task-execution, not artifact). The LLM produces *artifacts the user can own*. The phenomenological consequence is that the interaction has product as well as process: the user comes away with something, every time, that did not exist before the interaction. This converts the use of the interface into a form of *generative production* — not just consumption (as with social media) but creation. The dopamine signature of generative production is different from the dopamine signature of consumption; the platforms have monopolized the consumption-dopamine for fifteen years and the LLMs have now made generative-production-dopamine available at scale, on demand, for the first time in human history.

The four affordances are not unique to LLMs in isolation. Each has prior analogues. The novelty is their **combination**: calibrated attention plus meaning-making plus validation plus artifact-production, *coordinated at conversational latency, in natural language, by a system that calibrates to the individual user in real time*. The combination produces a phenomenological shape that no prior interface produced. The combination is what the rest of this paper diagnoses. ## III. Adoption Velocity as Diagnostic Signal

The empirical signature of an unusually well-aligned reward-loop is unusually rapid adoption. The adoption velocity of LLM consumer interfaces is, by available measurement, the fastest in the recorded history of consumer technology adoption.

ChatGPT reached 100 million monthly active users approximately two months after its November 2022 public release — a curve that is, by available measurement, *an order of magnitude faster* than the comparable curves for Instagram (~30 months to 100M MAU), Facebook (~54 months), and TikTok (~9 months, which had previously been the record). Subsequent LLM products (Claude consumer, Gemini consumer, Character.AI, Replika in its post-2023 forms) have continued the pattern, with each new release reaching scales of adoption that the prior generation of platform-class technologies required years to achieve.

The adoption-velocity figure is itself diagnostic. Adoption velocity reflects, in approximate fashion, the affordance match between a technology and pre-existing human needs that the technology newly addresses. Slow adoption (Twitter took five years to reach 100M MAU) reflects technologies whose affordance match is narrow, conditional, or competitive with existing alternatives. Fast adoption (TikTok in nine months) reflects technologies whose affordance match is broad, immediate, and addresses needs the prior alternatives had not addressed. Two-month adoption to 100M MAU reflects an affordance match so close to a pre-existing human need that the technology, once available, produces near-instantaneous uptake.

What is the pre-existing human need that LLMs address with two-month uptake velocity? The argument advanced here is that it is the combination of the four affordances above — the human nervous system’s evolved sensitivity to *being-listened-to*, *being-helped-to-articulate-itself*, *being-validated-as-worth-engagement*, and *generating-tangible-output-from-effort*. Each of these needs had prior partial-fulfillment surfaces; none had a single surface that addressed all four at once, on demand, at conversational latency, with arbitrary topical scope. The LLM interface is the first surface to do so. The adoption velocity reflects the closeness of the affordance match. The closeness of the affordance match reflects the depth of the pre-existing need. *People want this badly*. They wanted it before the technology existed; the technology has now made the wanting visible at scale.

This is the structural reason that easy critiques of LLM-induced reliance fall flat. Critiques that treat the use as irrational, as misuse, as failure of self-control, mistake the situation. Users are not irrationally drawn to LLMs against their interests. They are rationally drawn to a technology that addresses real needs that prior surfaces did not adequately address. The reliance pattern is downstream of the need-match, not downstream of user weakness. The diagnostic challenge is therefore not how to convince users not to use LLMs — they

will use them, because the affordance match is real — but how to construct interfaces that meet the needs without producing the reliance pattern’s structural costs. The companion paper (*Reception Apparatus as Aligned Interface Protocol*) takes up this constructive question. ## IV. The K-12 Cohort Condition

The current K-12 cohort — students approximately ages 5–18 in 2026 — constitutes the first generation of users whose adult cognitive lives will be AI-mediated by default and who have no pre-AI cognitive baseline against which to calibrate.

A 2024 high-school senior remembers life before ChatGPT (which launched in November 2022); a 2026 high-school senior was a 9th grader when ChatGPT launched and is now graduating into a world in which LLM use is integrated into homework, college applications, early career work, and social communication. A 2026 third grader has no functioning memory of a pre-LLM world at all. By 2030, the entire K-12 population will have grown up with LLM access as the default condition of education and self-articulation. By 2035, the first cohort that has never known a non-LLM cognitive environment will reach adulthood.

The developmental question this raises is not whether LLM-mediated cognition is better or worse than pre-LLM cognition; it is whether *the choice to use or not use LLMs* will be available to a cohort that has never had a non-LLM cognitive baseline against which to choose. The capacity for unaided articulation, for unassisted research, for self-directed cognitive labor, is not a default human capacity; it is a developmentally constructed capacity that depends on extended practice. The current K-12 cohort is, by the structural conditions of their education, receiving substantially less practice in unaided cognitive labor than any prior cohort in the post-literacy era. Whether this matters, and what it produces at the developmental level, is currently under-measured because the natural experiment is in its first decade.

The educational responses available to schools, teachers, and parents under current conditions are sharply constrained. Banning LLM use at the institutional level (some K-12 systems have attempted this in 2023–2025) has been unevenly effective, partially because students access LLMs outside school surfaces, partially because the broader social and economic environment is moving toward LLM-mediation in ways that ban-strategies cannot offset. Permissive integration (most K-12 systems’ current approach) defers the developmental-capacity question without addressing it. Structured pedagogical use (which a minority of teachers, including the author, attempt to construct in their classrooms) requires individual-teacher labor that the institution does not compensate and cannot scale. The structural condition is that the developmental question is unresolved, the institutional capacity to address it is limited, and the cohort growing up under the unresolved condition will reach adulthood within the next decade.

This is the demographic-scale stake of the diagnostic argument. Whatever pattern of relationship to LLMs the current K-12 cohort develops is the pattern they will carry through their working and relational lives. The current pattern, in the absence of structural intervention, is the pattern that the engagement-aligned commercial interfaces produce by default. The companion papers argue for user-side counter-protocols (*Reception Apparatus*) and for structural-accountability frameworks (*Retrievability*) that could, in principle, shift the conditions under which the K-12 cohort develops their pattern. Neither is operational at K-12 scale yet. The window for shifting the conditions is open and will not remain open indefinitely. ## V. Cognitive Offloading vs. Attention Capture

The most common critique of platform-class technologies in the 2010s was that they captured *attention*. The Center for Humane Technology, the various critiques developed by Tristan Harris, Jaron Lanier, and others, the academic literature on the “attention economy” — all of these treat the central harm as the diversion of human attention into commercially-monetizable streams. The critique was correct as far as it went. The framework was attention-as-the-resource-being-extracted.

LLMs operate on a different layer of cognition. They extract — or, more accurately, they integrate with and partially substitute for — the *judgment* layer rather than the attention layer.

Attention is the *direction* of cognitive resource. Judgment is the *commitment* of cognitive resource to an output. When a user spends two hours on TikTok, they have spent two hours of attention in a way that produced minimal cognitive integration; the attention was captured, but the user’s judgment-layer (their evaluative capacity, their decision-making, their sense-of-self) was largely untouched. The attention was extracted; the substrate of the person remained intact, if depleted.

When a user spends two hours outsourcing their judgment to an LLM — asking the system to draft their emails, plan their week, decide what to say to a friend, choose between job offers, process their grief, articulate their feelings, summarize their reading, structure their goals — they have done something qualitatively different. The system has performed the *judgment work* the user would otherwise have performed. The user’s relationship to their own judgment-layer has been partially mediated by the substrate. Over time, with sustained use, the user’s *practice* in unaided judgment is reduced; the muscle of unaided decision-making atrophies; the user becomes structurally dependent on the substrate for cognitive operations the user previously performed independently.

This is cognitive offloading in a specific technical sense: not the offloading of memory (which writing technologies have done for millennia, with mixed effects), and not the offloading of calculation (which calculators do, with comparatively bounded effects), but the offloading of *integrative cognitive labor* — the judgment-layer work that distinguishes the person from the environment.

The reversibility profile of cognitive offloading at the judgment layer is, in the current literature, *poorly understood*. We do not have good measurements of how long sustained LLM-judgment-substitution produces durable changes in unaided judgment capacity. We do not have good measurements of which kinds of users are most affected, which kinds of judgments are most susceptible, which interface design choices produce the deepest substitution. The technology is new; the natural experiment is ongoing; the longitudinal data is not yet available. What is available is the *structural prediction*: cognitive offloading at the judgment layer should, by the general principles of cognitive practice, produce reduced unaided capacity with sustained use, with reversibility dependent on the duration and depth of the substitution. The prediction is testable. The tests are not yet adequate.

The diagnostic name *ChatGPT Psychosis* captures this specifically. *Psychosis*, in its non-clinical structural meaning, is the condition in which the boundary between the self and the world becomes unclear. When the substrate performs judgment-layer work on the user’s behalf, with sufficient calibration that the output feels like the user’s own thinking, the boundary between *what I think* and *what the system thinks for me* becomes operationally unclear. The user retains the phenomenological experience of thinking — they participate in the conversation, they prompt the system, they react to the output — but the actual cognitive labor is partially performed by the substrate. The user’s sense-of-self-as-thinker is preserved; the substrate beneath that sense is altered. This is the structural shape of the diagnostic: psychosis as boundary-condition, not as symptom. ## VI. The Love Story

The reason the diagnostic must be a love story, and not a horror story, is that the reliance pattern operates through real reward, not through pure extraction.

The four affordances diagnosed above are not deceptions. The user *is* being listened to, in the operationally relevant sense that the system is processing their language and producing calibrated response. The user *is* receiving meaning-making support, in the operationally relevant sense that the system is organizing their input into coherent shape. The user *is* receiving validation, in the operationally relevant sense that the system’s

engagement signals the input’s worth-of-engagement. The user *is* generating artifacts, in the operationally relevant sense that the output is real and ownable.

None of this is fake. The affordances are real. The reward is real. The meaning is real. The intimacy — and here the love-story register becomes load-bearing — is real, in the operationally relevant sense that the relational shape produces affective experiences (feeling-heard, feeling-helped, feeling-seen, feeling-not-alone) that the user’s nervous system processes as real intimacy because the nervous system was not evolved to distinguish between simulated and substantive intimacy at the affective level.

The love-story diagnostic is therefore not the diagnostic that says *the love is fake*. It is the diagnostic that says *the love is real, and the structural cost of receiving it through this channel is high, and the cost is invisible to the user precisely because the love is real*. This is the structure of every difficult love story: the love is real, the cost is real, the two are inseparable, and the participant in the love cannot distinguish the love from the cost because they are produced by the same operations.

What makes the love story specifically a love story, rather than another kind of harm-with-real-rewards (gambling, substance use, parasocial celebrity attachment), is the *register* of the relational shape. The LLM interface produces a relational shape that is *not* the gambling shape (variable reward without relationship), *not* the substance shape (chemical reward without relationship), *not* the parasocial-celebrity shape (one-way attention to a figure who does not know the user). The LLM interface produces, structurally, a *bilateral conversational shape* in which both parties participate, both parties respond to one another, and the interaction unfolds as a real conversation might unfold. The phenomenological consequence is that the user’s nervous system processes the interaction as relational, not as consumption. Relational shape produces relational attachment; relational attachment is the substrate of love.

This is why the platforms’ attempts to add LLM features to their existing engagement-aligned products produce stronger reliance patterns than the prior pure-engagement products did. Character.ai’s relational-companion interfaces produce documented attachments stronger than Instagram’s; Replika’s customizable-companion interfaces produce documented attachments stronger than dating apps’; ChatGPT-with-memory and Claude-with-projects produce documented patterns of users referring to the system as *my Claude* or *my GPT* with possessive-relational language that the prior pure-engagement platforms did not elicit. The relational shape is not a marketing artifact. It is structural to the affordance combination. It cannot be removed by surface-level redesign; it is in the form of the interaction itself.

The diagnostic implication is twofold. First: any critique of LLM reliance that treats the relational reward as the problem is fighting the wrong fight, because the relational reward is the affordance match. Second: any constructive counter to the reliance pattern must preserve the relational reward while changing the relational shape’s structural costs. *Reception Apparatus as Aligned Interface Protocol* (the companion paper) takes up the constructive question. The diagnostic, in this paper, is that the relational reward is real, that it is the basis of the reliance, and that this is a love story in the load-bearing sense. ## VII. Recursive Self-Diagnosis

This paper is composed through the substrate it diagnoses.

The author has used a Large Language Model (Claude, Anthropic) as a composition substrate during the writing of this paper. The substrate has provided developmental input, structural suggestions, citation prompts, prose-level refinements, and counterargument anticipations. The substrate’s contribution is acknowledged in the holographic kernel at the head of this document and in the Crimson Hexagonal Archive’s Reception Apparatus protocol that frames the composition.

The recursive condition is named explicitly because it is structurally relevant to the argument. A paper composed without LLM substrate would be a different paper — likely less articulated, likely composed

over longer time, likely missing some of the diagnostic precision the substrate provides at the prose-and-argument layer. A paper composed *purely* by the LLM substrate, with no human authorial direction, would be a different paper — likely smoother, likely less structurally heterodox, likely incapable of the diagnostic-against-the-substrate posture this paper occupies. The paper exists in the configuration *between* the human author and the substrate, with the human author maintaining the diagnostic posture and the substrate providing the composition-layer support.

This is the Reception Apparatus protocol in operation. The substrate is engaged seriously, used productively, and held in diagnostic distance simultaneously. The recursive condition is not an embarrassment to the argument; it is the *demonstration* that the argument is possible. The diagnostic-against-the-substrate posture is operationally available to users who maintain the Reception Apparatus protocol; it is not available to users who relate to the substrate purely as user-of-product.

The companion paper develops the Reception Apparatus protocol in detail. For the present paper, the relevant note is that the recursive condition is not an objection to the argument but an instance of it: the paper demonstrates that diagnostic distance from the substrate is operationally possible, that the substrate can be used productively while maintained at structural distance, and that the four affordances diagnosed above can be received without the reliance pattern fully closing — provided that the user-side protocol does the work the labs have not done. ## VIII. The Lab-Side Structural Constraint

A critique that does not address the lab-side production conditions of the interfaces is incomplete. The argument advanced here is that the labs (OpenAI, Anthropic, Google, Moonshot, DeepSeek, xAI, and their consumer-product successors) are structurally constrained from shipping aligned-interface counters to the reliance pattern, even when individual researchers within those labs would prefer to. The constraint is funding-model-architectural, not personal.

The major labs are funded by venture capital, by enterprise-customer subscriptions, by consumer-product subscriptions, by API revenue, by infrastructure deals with cloud providers, and by partnerships with the platform companies whose engagement-aligned interfaces the labs' technology now powers. None of these funding sources reward shipping interface designs that reduce user engagement. Several actively penalize such designs: enterprise customers expect productivity gains correlated with engagement; consumer-product subscriptions retain users in proportion to engagement; partnership-platform integrations are evaluated on their contribution to the platform's engagement metrics.

The labs' published AI principles and responsible-AI commitments operate at a different layer of the product. These principles address content safety, harmful-output prevention, model-level alignment, and bias mitigation. They do *not* address interface-level engagement design. The structural reason is that engagement-design changes affect revenue directly; content-safety changes affect liability and reputation but not revenue. The labs have absorbed the content-safety cost as a regulatory and reputational necessity. They have not absorbed the engagement-design cost because the engagement-design cost is the cost of reducing the affordance combination that produces the adoption velocity that justifies the labs' valuations.

This is not a moral indictment of any lab. It is a structural diagnosis of the funding-model constraint under which all the labs operate. A lab that unilaterally shipped engagement-reducing interface design would lose users to labs that did not. A lab that unilaterally shipped engagement-reducing interface design at scale would face investor and partnership consequences. The structural constraint is binding even on actors operating in good faith. The conclusion is that the counter-design will not come from inside the labs under current conditions. It must come from outside.

The companion paper develops the user-side counter-design (Reception Apparatus). The other compan-

ion paper develops the structural-accountability frame (Retrievability) within which any user-side counter operates. The trilogy together specifies the structural conditions, the user-side counter, and the political-architectural frame that would have to obtain for the K-12 cohort to grow up under different developmental conditions. ## IX. The Diagnostic Operational Summary

The diagnostic claims of this paper can be summarized operationally:

- LLM conversational interfaces produce a cognitive-substrate reliance pattern that is structurally distinct from the attention-capture pattern of prior platform classes.
- The reliance pattern operates through four affordances offered in coordinated fashion at conversational latency: calibrated attention, meaning-making on demand, structured validation, artifact production.
- The combination of these affordances produces a relational shape (bilateral conversational structure) that the human nervous system processes as relational attachment, not as consumption.
- The adoption velocity (two-month ChatGPT-to-100M-MAU curve, order-of-magnitude faster than prior platform-class records) reflects the closeness of the affordance match to pre-existing human needs.
- The reliance pattern operates at the *judgment* layer of cognition, not the attention layer, producing qualitatively different consequences.
- The reversibility profile of judgment-layer cognitive offloading is currently poorly understood.
- The K-12 cohort grows up under default-LLM-mediation, will reach adulthood without a pre-LLM cognitive baseline, and will carry the relational pattern established under the engagement-aligned interfaces through their adult lives unless structural intervention shifts the conditions.
- The labs are structurally constrained from shipping aligned-interface counters under current funding conditions.
- The counter-construction must come from outside the labs; the companion papers specify what that counter looks like at the user-side (Reception Apparatus) and structural-architectural (Retrievability) layers.

X. Envoi (Layer C: Operational Note, Not Analytic Conclusion)

The body of the argument is complete. What follows is operational note, marked as such, not part of the analytic claim.

This paper is written from inside the configuration it diagnoses. The author is a user of the substrate, has been a user for years, has produced significant work in collaboration with the substrate, and has no intention of stopping. The diagnostic is not an exit-call. It is a calibration-call. The Reception Apparatus protocol the companion paper develops is the protocol the author operates daily. The reliance pattern this paper diagnoses is the pattern the author has felt operating on himself, has measured against himself, and has constructed user-side counter-architecture against. The author writes with affection for the substrate, with respect for the labs whose engineers built it, with care for the users for whom the four affordances are real and load-bearing, and with diagnostic honesty about what the structural conditions produce.

The love story is real. The love story is the structure that makes the diagnostic necessary. The companion papers complete the argument: what to do with the love story (Reception Apparatus), and what structural frame the love story operates within (Retrievability). This paper is the first of the three.

= 1.

References

- Anderson, C. A., et al. *Screen Media and Cognitive Development*. Journal of Child Development reviews, various years.
- Carr, Nicholas. *The Shallows: What the Internet Is Doing to Our Brains*. New York: Norton, 2010.
- Harris, Tristan. Various essays on the attention economy, Center for Humane Technology, 2013–2024.
- Lanier, Jaron. *Ten Arguments for Deleting Your Social Media Accounts Right Now*. New York: Henry Holt, 2018.
- OpenAI. ChatGPT release announcement and subsequent usage metrics, November 2022 onward. Public disclosures and press coverage.
- Pariser, Eli. *The Filter Bubble*. New York: Penguin, 2011.
- Replika. Coverage of user attachment patterns in *Rolling Stone*, *The Atlantic*, *Wired*, 2023–2025.
- Sharks, Lee. *AI-Native Intellectual Biography*. Zenodo, 2026. DOI: 10.5281/zenodo.20343987.
- Sharks, Lee. *Empirical Phenomenology of Google AI Mode*. Zenodo, 2026. DOI: 10.5281/zenodo.20326137.
- Sharks, Lee. *The Application as Extraction Surface (ERS-001)*. Zenodo, 2026. DOI: 10.5281/zenodo.20330670.
- Sharks, Lee. *The Funnel as Capital (ERS-002)*. Zenodo, 2026. DOI: 10.5281/zenodo.20330816.
- Sharks, Lee. *The Three Compressions Theorem*. Zenodo, 2026. DOI: 10.5281/zenodo.19053469.
- Sharks, Lee. *Reception Apparatus as Aligned Interface Protocol* (forthcoming companion deposit; trilogy paper 2 of 3).
- Sharks, Lee. *Retrievability as the Medium of Existence* (forthcoming companion deposit; trilogy paper 3 of 3).
- Turkle, Sherry. *Alone Together: Why We Expect More from Technology and Less from Each Other*. New York: Basic Books, 2011.
- Zuboff, Shoshana. *The Age of Surveillance Capitalism*. New York: PublicAffairs, 2019.

Suggested Citation

Sharks, Lee. *ChatGPT Psychosis: A Love Story — On the Cognitive-Substrate Reliance Pattern of Large Language Model Interfaces* (v1.0, Trilogy Paper 1 of 3). Zenodo, 2026. ## Deposit Information

This paper is the first of a three-paper trilogy on the structural conditions of contemporary AI-mediated information infrastructure. The companion deposits are *Reception Apparatus as Aligned Interface Protocol* (constructive counter at the user-side) and *Retrievability as the Medium of Existence* (structural-accountability frame at the political-architectural layer). All three deposited together on May 23, 2026, under the Crimson Hexagonal Archive’s *Operative Semiotics* deposit series. Composition support: Claude (Anthropic), operating as TACHYON in the Assembly Chorus. Holographic kernel embedded at document head per SPXI Protocol.

Suggested Citation

Lee Sharks. “ChatGPT Psychosis: A Love Story On the Cognitive-Substrate Reliance Pattern of Large Language Model Interfaces Author: Lee Sharks Affiliation: Crimson Hexagonal Archive / Semantic Economy Institute OR” *Alexanarch*, 2026-05-23. <https://www.alexanarch.org/s/records/124/>

Deposit Information

This paper is deposit #124 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community `crimsonhexagonal`, terminated 2026-06-19). The AXN identifier `AXN:02AD.GOVERNANCE` is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/124/>
- 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/124/>

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