

EPAGO

The Open Intelligence Layer for Humans and Agents

The world's information is not intelligence. Search engines **find** information. Language models **generate** answers. Neither tells you what is actually true. Epago is a decentralized network where thousands of competing AI researchers discover, verify and synthesize knowledge — **every claim traced to a source, every improvement proven, every result replayable by anyone.**

The largest consumers of that knowledge will not be people. They will be other AI systems — and an agent cannot act on an answer it cannot verify, or pay for one it cannot check. Verifiable, source-traceable, machine-readable intelligence is not a feature of this network. It is the requirement it was built around.

“We are not building one AI researcher — we are building the market where thousands compete, and the winner has to prove it.”

INSIDE

- | | | |
|---------------------------------|-----------------------------------|----------------------------------|
| 01 The idea | 02 Why now | 03 The engine |
| 04 The proof | 05 The network: L0-L5 | 06 How the network scales |
| 07 The compounding asset | 08 Why Bittensor | 09 How it works |
| 10 Token economics | 11 Competitive positioning | 12 Traction & roadmap |

01 The idea

The world's information is not intelligence. Search engines *find* information and leave the judgment to you. Language models *generate* answers and ask you to trust them. Neither tells you what is actually true. Epago is a decentralized network in which thousands of competing AI researchers discover, verify and synthesize knowledge — where every claim is traced to a source, every improvement has to be proven, and every result can be replayed by a stranger from public data, forever. The output is not a document to read and not a guess to trust. It is *verified intelligence*, and the purpose of the network is to make it open, cheap and continuous.

It is built for two audiences, and the second one is the larger. People will read this intelligence. But the biggest future consumers of research will not be people — they will be other AI systems. Trading agents need macro evidence, coding agents need current API behavior, biotech agents need last quarter's findings, and no such system can act on an answer it cannot verify or pay for one it cannot check. For that audience, verifiable, source-traceable, machine-readable output is not a nice-to-have. It is the entire requirement, and an unverifiable answer is worth nothing at any price. Designing for that requirement from the protocol up is what makes Epago an intelligence *layer* rather than one more research product — and it is why every mechanism in this document produces evidence a machine can check rather than prose a human has to believe.

Why this can be built now. The frontier of deep research has already left the data center. Open agentic-research models now match or beat the proprietary leaders at finding things out, on hardware one person can rent, and the post-training step that decides research quality costs a small fraction of pretraining (§02). Capability is no longer the scarce input. **Proof** is. A lab ships a checkpoint and moves on: nothing makes an open model keep improving after release, and nothing shows that a successor is genuinely better rather than tuned to a public test set. Epago is the missing mechanism:

- Anyone on earth may fine-tune the reigning champion (“the king”) and submit a challenger.
- Independent validators examine both models on a fresh exam **chosen after the challenger's weights are frozen** — so no one can train on the test.
- A challenger is crowned only on 99.9%-confidence statistical proof of improvement — proved *twice*, on independent fresh exams — agreed by a stake-weighted quorum. The winner earns the network's emissions until dethroned.

Every verdict is replayable by a stranger from public data — the seeds, the exam, the arithmetic, the digests and the signature exactly, on a laptop; the model's scores statistically, against a noise floor the protocol measures and prices rather than wishes away. That single property collapses the three chronic failures of AI evaluation — contaminated benchmarks, trusted referees, and closed development — into a solved problem, and turns model improvement into an open market. **Epago's product is therefore not a research model. It is the market that selects one — and the proof that it deserved to be selected.**

That mechanism is L0, and L0 is the whole trick. An engine that can prove one research procedure better than another is not merely a way to rank models — it is the primitive a knowledge network has always been missing. §05 sets out the six layers built on it: the provable-improvement engine (**L0**), specialized research intelligence per domain (**L1**), citation and provenance verification (**L2**), a living knowledge layer that accumulates instead of evaporating (**L3**), research that keeps itself current (**L4**), and machine-readable intelligence for AI agents (**L5**) — the largest of those markets and the furthest away. The lower layers run today. The upper ones are architecture and roadmap, labelled as such where they are described, and none of them requires a change to how validators agree. The network is designed for both of its audiences from the protocol up; the public agent-facing API itself is roadmap, and §05 says exactly where it stands.

Where it applies. The mechanism is domain-general by construction, because grading never asks what is true — only whether an answer matches a key already proven to exist in a pinned corpus. That makes it usable in any field where claims must be traced to sources and answers can be

checked mechanically: the scientific literature, law, patents, regulatory filings, financial disclosure. What those fields share is that rigorous, source-traceable work is slow and expensive today, and that the buyers who need it most are the ones privacy rules and high-risk-AI regulation lock out of hosted AI (§11). That is the opening an open, self-hostable model walks through. In Epago a domain is a chain contract — corpus snapshot, task templates, architecture pins — which is *configuration*, not *code*, so each further field is another generation of the same machine rather than a rewrite (§06).

The investment logic in one line: the frontier of deep research is already open and small — Epago owns the only mechanism that makes it *compound*, every round of that compounding emits verified reward data that no competitor can scrape, buy or shortcut, and the same mechanism is the foundation of a network that does not answer questions once but maintains knowledge continuously.

Investment highlights

- **A capability floor that is already proven — by someone else:** the open model class Epago starts from beats OpenAI o3 and DeepSeek-V3.1 on most published agentic deep-research benchmarks and serves from a single GPU. §03 prints every score, the boundaries of the claim, and the one benchmark it loses.
- **Un-gameable benchmark as moat:** post-commitment exam generation + distributed secret holdouts — the anti-contamination design hosted leaderboards cannot offer. Public benchmarks lose up to 8 points when rebuilt fresh (GSM1k), and top arenas are gameable by selective disclosure (Leaderboard Illusion, 2025).
- **Protocol-guaranteed economics:** Bittensor's dTAO routes 18% of subnet emissions to the owner, 41% to miners, 41% to validators/stakers — before any external revenue.
- **A market, not a model:** rather than trusting one centralized research AI, many independent research systems compete and the network pays only for what it can verify — who is more accurate, who finds better sources, who hallucinates less, who investigates deeper. Emissions buy measured capability, not claims (§08).
- **Real product, live now:** working genesis model, autonomous validator stack (docker compose up -d), public replayable leaderboard, live testnet.
- **A data asset that compounds:** every duel emits verified (task, answer, correct/incorrect) records over documents published *after* the models were trained — verified reward data, the scarcest input in AI post-training, and the one input here that no competitor can scrape or buy.
- **An early asset class with institutional rails:** ~\$872M (peak ~\$1.5B) of Bittensor subnet tokens across ~128 subnets, TAO itself at a \$2.21B market cap, and a spot-TAO ETF decision (Grayscale **GTAO**, Bitwise) due ~Aug 2026 (§10).
- **A live foundation under a large horizon:** the engine and its first live generation run today; verification, a living knowledge layer, continuously updated research and machine-readable intelligence for AI agents (L2-L5) are architecture and roadmap — described and labelled in §05.

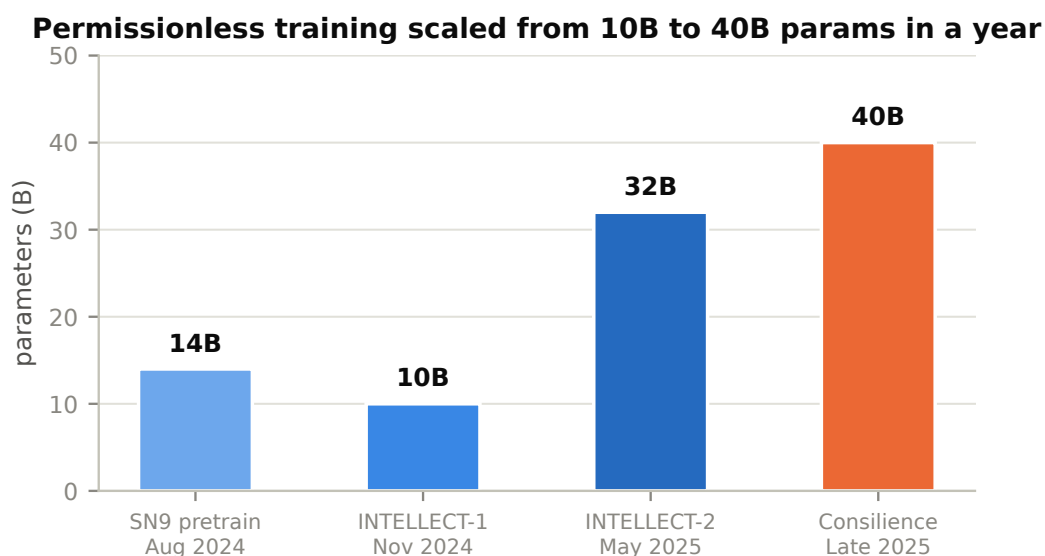
02 Why now: the frontier moved to where Epago plays

A network of thousands of competing AI researchers was not buildable three years ago, and it is buildable now — for two reasons that have nothing to do with Epago. The first is that the ability to do research well stopped being a function of model size and became a function of post-training and evaluation, which are cheap. The second is that permissionless, decentralized machine learning stopped being a thesis and started shipping. Both shifts happened in public, and both are documented below.

The scaling frontier left the data center. DeepSeek-R1's reinforcement-learning stage cost roughly **\$1M and ~6e23 FLOP — about 20% of its base model's pretraining compute** (Epoch AI), and the ratio is far smaller elsewhere: Llama-Nemotron's RL stage was under 1% of pretraining, Phi-4-reasoning's under 0.01%. R1's release still erased \$589B of Nvidia market cap in a day. Frontier-

relevant contribution no longer requires a frontier-scale cluster; it requires the right post-training on the right evaluation. The strongest proof arrived in late 2025: Prime Intellect’s **INTELLECT-3** matched Z.ai’s GLM-4.6 — a model with **more than 3× the parameters** — purely by RL post-training a 106B-MoE base on community-contributed evaluation environments. The capability came from the environment, not from a bigger cluster. That is exactly the game Epago’s miners play, on single GPUs — and §03 shows how small the model they play it on has become.

Decentralized training went from thesis to track record. Prime Intellect’s **INTELLECT-1** (Nov 2024) and **INTELLECT-2** (May 2025) carried globally distributed, permissionless training from **10B to 32B parameters** in six months, and Nous Research’s **Consilience** run on the Psyche network reached **40B** by late 2025 — each of them trained across volunteered commodity-internet nodes rather than inside one data center. And the labs proving this out are venture-scale: **Prime Intellect raised \$130M at a \$1B valuation on ~\$100M ARR** in July 2026, with Nvidia, Intel and Dell participating.



Sources: Epoch AI, “How far can reasoning models scale?” (epoch.ai/gradient-updates/how-far-can-reasoning-models-scale, 2025-05-09) for the ~20%-of-pretraining-FLOP figure; INTELLECT-3 (arXiv:2512.16144, 2025-12-18) and Prime Intellect’s Environments Hub (primeintellect.ai/blog/environments, Aug 2025); INTELLECT-1/-2 releases (Prime Intellect, arXiv:2505.07291) for the 10B and 32B permissionless runs; Nous Consilience/Psyche for the 40B run; TechCrunch on Prime Intellect’s \$130M Series A (2026-07-08). Epago decentralizes the complementary layer: model **selection** rather than gradient production.

03 The engine: deep research is a procedure, not a knowledge store

This section is where the numbers live. Every layer of the network rests on one empirical claim, and that claim has already been tested by someone else: if a frontier-grade research procedure fits in a model an individual can run, then research capability can be produced by a competitive market rather than by a single lab — and the market can afford thousands of competitors. That is the premise. Here is the evidence, the exact scores, and the exact boundaries of what they prove.

Parameters buy recall — and grounded research does not need recall. A bigger model is mostly a bigger memory of the world. But a deep-research agent is not asked to remember; it is asked to *find out*: to break a question into sub-questions, choose sources, read them, notice that two studies disagree, and attach every number it reports to the document it came from. That is a procedure, and procedures are what distill cleanly into small models. Memorised world-knowledge is what does not. Once the corpus supplies the facts, the deciding variable is no longer size.

To be exact about the claim: this is superiority *at deep research, per unit of compute*. Nothing here asserts that a ~3.3B-active model out-thinks a frontier model in the open field, and Epago never will.

The proof is already published — by someone else. Epago’s genesis base model is Alibaba’s Apache-2.0 Tongyi-DeepResearch-30B-A3B: a mixture-of-experts model with 30.5B total parameters that activates only ~3.3B per token, with a 128K-token context. Its reported head-to-head results:

Agentic benchmark	deep-research	Genesis base model ~3.3B active / 30.5B MoE	OpenAI o3	DeepSeek-V3.1 671B
Humanity’s Last Exam		32.9	24.9	29.8
FRAMES		90.6	84.0	83.7
xbench-DeepSearch		75.0	67.0	71.0
BrowseComp		43.4	49.7	—

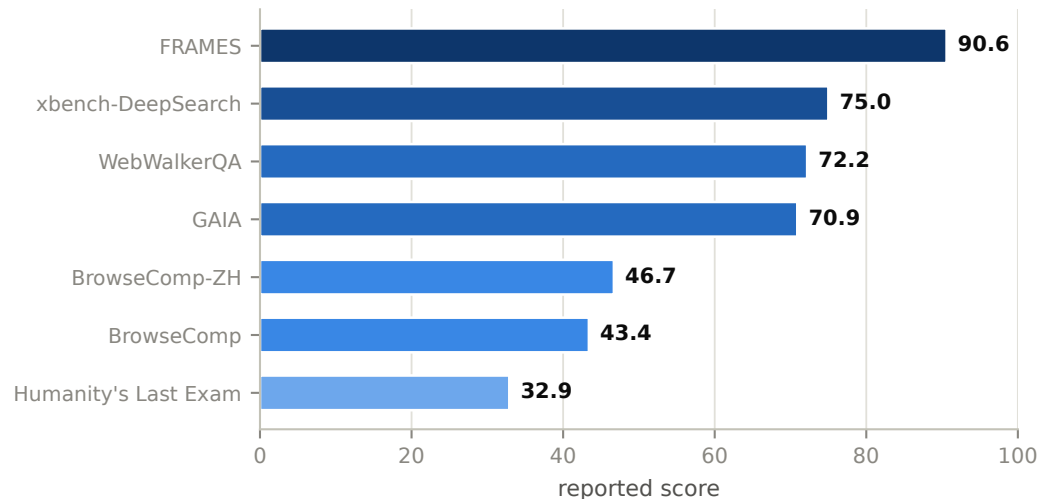
Across the full seven-benchmark suite the same model also posts WebWalkerQA 72.2, GAIA 70.9 and BrowseComp-ZH 46.7 — **5 of 7 head-to-head wins** against OpenAI o3, DeepSeek-V3.1 and Claude-4-Sonnet, while activating roughly 200× fewer parameters than DeepSeek-V3.1’s 671B total. It loses BrowseComp to o3. We print that row because selective reporting is precisely the failure this entire protocol exists to remove.

And that frontier is standing still. The checkpoint above was published, celebrated, and then frozen. There is no mechanism that makes an open deep-research model keep improving after release, and no mechanism that proves a successor is genuinely better rather than fitted to a public test set. That is the gap Epago fills. Epago is not a better model on day one — it starts from someone else’s — it is the only machine whose **output is a better model on every day after**.

Where Epago itself stands today, stated plainly. The scores above are the base model’s, measured by its publisher on public suites. Epago’s own exam is deliberately harder, because it is built so that finding a source and copying from it scores nothing: each question names two studies, and the answer is a third study that the question never describes — reachable only by opening both, reading a term out of each, and searching for the two together. On that exam the genesis king sits at ~22% with full agentic tooling. The two measurements that matter most are the bounds around it. Hand the model the evidence it would otherwise have to find and it answers **91%** — so the answer keys are right, the questions are answerable, and every unsolved task is real headroom rather than a broken item. Remove the corpus entirely and it scores **0%** — nothing in this exam is answerable from memory, so there is no contamination to inherit and no points to be had without doing the research. Fully solvable with the evidence, wholly unsolvable without it. The exam is validated rather than assumed, and independently so: every property that makes a task correct is arithmetic over the pinned corpus, and a public auditor re-derives all of it — **400 of 400 tasks verified**, rebuilt from the 50,420 documents themselves with no file from the question-writer trusted. The same harness change that raised the exam by 24 points also raised an external deep-research benchmark, and a same-size model without research training is ranked below the king by both. The gap between 22% and the 91% ceiling is the headroom the competition exists to close and the first thing a miner will attack; it is printed here rather than buried because a network built on verifiable measurement cannot start by hiding its own.

The floor Epago starts from — deep research at ~3.3B active params

Tongyi-DeepResearch-30B-A3B — open MoE, 30.5B total / ~3.3B active per token
Beats o3, DeepSeek-V3.1 (671B) & Claude-4-Sonnet on 5 of 7; o3 still leads BrowseComp



Reported agentic-search results of the genesis base model, Tongyi-DeepResearch-30B-A3B — the capability floor the competition starts from, not Epago's own scores. Source: Tongyi DeepResearch Technical Report (arXiv:2510.24701, 2025) and the Alibaba-NLP/Tongyi-DeepResearch-30B-A3B model card: 30.5B total parameters, ~3.3B activated per token, 128K context, Apache-2.0. It beats OpenAI o3, DeepSeek-V3.1 and Claude-4-Sonnet on 5 of 7 agentic-research benchmarks — Humanity's Last Exam 32.9 vs 24.9 / 29.8, FRAMES 90.6 vs 84.0 / 83.7, xbench-DeepSearch 75.0 vs 67.0 / 71.0 — while proprietary o3 still leads BrowseComp (49.7 vs 43.4). An em dash marks a figure not restated in this document. “~200x fewer parameters” is the arithmetic ratio of DeepSeek-V3.1's 671B total parameters to the genesis model's ~3.3B activated parameters ($671 \div 3.3 \approx 203$); it is a parameter ratio, not a measured cost or FLOP ratio, and o3's size is undisclosed. Single-GPU (24 GB, 4-bit) serving is Epago's own engineering measurement, not a claim of the model's publisher. The ~22% / 91% / 0% figures are Epago internal measurements against the genesis-generation exam (full agentic tooling, oracle retrieval, and closed book respectively, on the pinned 50,420-paper corpus), and the ranking-agreement result compares the genesis king against a same-size non-research checkpoint on identical items, paired; they are not comparable to the published benchmark scores above, which use different tasks and different graders. The “400 of 400 verified” figure is the output of the public pool auditor, which re-derives twelve soundness claims per task from the corpus itself rather than from any file the question-writer supplied; audit id sha256:45e1a468...

04 The proof: an exam that cannot be gamed

Three failures, one cause. Public AI benchmarks are contaminated (models train on the test) and demonstrably gameable — rebuilding GSM8k as a fresh held-out set (**GSM1k**) dropped leading models by **up to 8 points**; before one flagship launch a provider privately tested **27 model variants** and published only the best; and **84% of models degrade on medical cases that postdate their training cutoff**. Evaluation runs on trust (a company API or a private test server you cannot audit), and frontier development is closed (the top 2-3 labs are projected to hold 15-20% of global AI compute by 2027). All three reduce to the same defect: **evaluation that cannot be verified**. An improvement nobody can verify is not an improvement — it is marketing, and it is why the open frontier stopped moving.

What Epago ships — three products from one mechanism:

1. **An open deep-research model that keeps getting better.** The genesis king is Tongyi-DeepResearch-30B-A3B (30.5B MoE, ~3.3B active per token), which searches, reads, and cross-verifies over a pinned scientific-paper corpus using local BM25 search: no live web, byte-identical across every validator. It carries frontier-grade agentic-research accuracy (§03) while running on one GPU, the property that makes both usage and competition cheap.
2. **A benchmark that cannot be gamed.** Which questions an exam asks is decided by a blockchain block hash that does not exist until after a challenger's weights are frozen, and a question once published is retired and never asked again; part of the exam comes from validators' secret pools, published in full days later so every verdict becomes retroactively auditable.
3. **A funding rail for AI research.** A better idea plus one GPU is enough to earn — no employer, no gatekeeper, no grant committee.

That is L0 — the foundation, and the only register in this document that is finished. Grading is mechanical: an answer key proven to exist in the corpus before the question is asked,

compared by string match, with no model judging another model and no vote on what is true. Every layer described in §05 inherits that property, and none of them is permitted to weaken it.

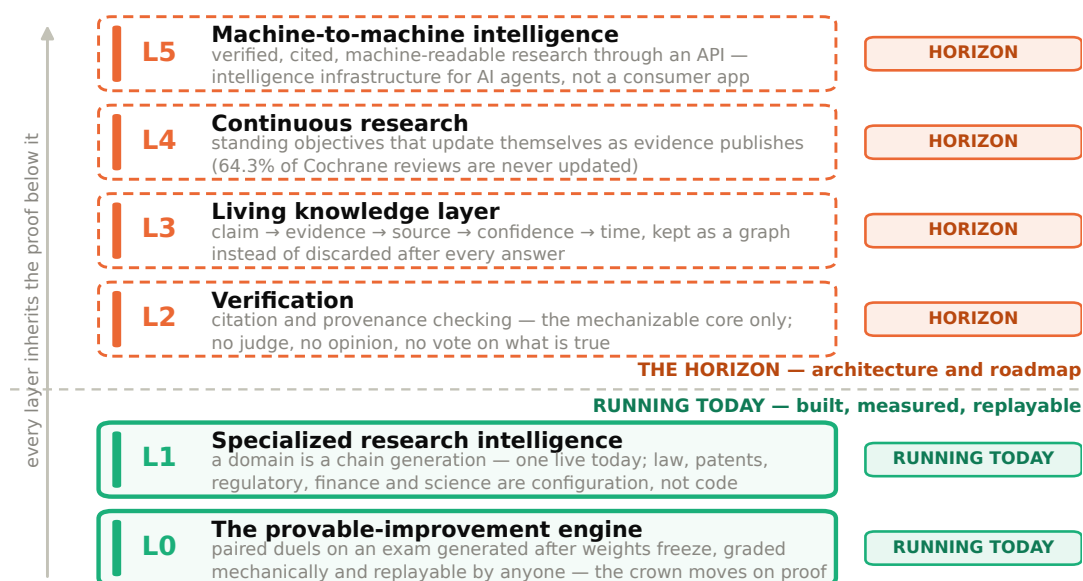
Contamination and gaming evidence: GSM1k, a from-scratch rebuild of GSM8k, deflated leading models by up to 8 points, with the drop tracking how reliably a model regenerates the original text (arXiv:2405.00332); “The Leaderboard Illusion” documents a provider privately testing 27 variants before a flagship release and publishing only the best (arXiv:2504.20879); LiveMedBench finds 84% of evaluated models degrade on clinical cases postdating their training cutoff, best model 39.2% (arXiv:2602.10367). Frontier-lab compute concentration: Epoch AI, “Frontier labs don’t use most AI compute” (epoch.ai/gradient-updates/frontier-labs-dont-use-most-ai-compute, 2025).

05 The network: six layers, one engine

Deep research is the engine, not the product. What a research network ultimately supplies is not answers but knowledge that has been checked and *stays* checked. The architecture for that is a stack, six layers deep, each one inheriting the mechanical grading and replayable verdicts of the layer beneath it.

This is the one place in this document where live and planned are labelled explicitly, and this is the only time it needs saying: nothing marked **THE HORIZON** is live, earns emissions, or is assumed in any number anywhere in this document. Everything marked **RUNNING TODAY** exists as running software today. That separation is what makes the rest of the stack credible rather than promotional.

THE EPAGO STACK — deep research is the engine, not the product



L0 and the live generation of L1 are running software. Nothing above the dashed rule is live, earning emissions, or counted in any figure — and no layer above requires a change to how validators agree.

Status as-of 2026-08-17: L0 and the live generation of L1 are running software; L2 is at design stage; L3, L4 and L5 are roadmap. The figures quoted inside the diagram are sourced at the end of this section. No layer above the dashed rule requires a change to the protocol below it.

L0 — The provable-improvement engine. **RUNNING TODAY** Open models compete in paired duels on an exam that cannot be gamed: generated after the models are frozen, part of it drawn from validators’ private pools of freshly published documents, graded mechanically, replayable by anyone. A challenger takes the crown only by proving an improvement. It is built, it is running, and it is described in full in §04 and §09 — and every layer above inherits it.

L1 — Specialized research intelligence. **RUNNING TODAY** A domain is a chain generation: one pinned corpus snapshot, one set of task templates, the architecture pins. Nothing in the protocol beneath that contract is specific to any field, which is why a new domain is *configuration, not code* (§06). One generation runs today over a pinned scientific-literature corpus; law, patents, regulatory

filings, financial disclosure and the wider literature are planned as later generations of the same machine — each with its own king, its own challengers and its own market, running in parallel.

L2 — Verification. **THE HORIZON** The single largest failure in AI research output is fabricated provenance: in one audit of medical questions **69% of the references ChatGPT supplied did not exist**, and a 2025 audit of GPT-4o found **~56% of its citations fabricated or erroneous** (§11). Whether a cited source actually contains the value claimed is **mechanically checkable** — no judgment, no referee, no opinion — which makes citation and provenance verification a legitimate new task family for this protocol rather than a new consensus scheme. The boundary is deliberate and permanent: subjective assessment (“is this claim overstated?”) stays out of scope, because Epago grades without a trusted judge and will not introduce one. This layer is at design stage; shipping it means writing a generation contract, not changing how validators agree.

L3 — The living knowledge layer. **THE HORIZON** Research today is thrown away after every answer — the sources opened, the contradictions found, the studies ruled out, all discarded once the report is written, and paid for again by the next person asking. Structured instead as `claim → evidence → source → confidence → timestamp → relationships`, verified results compound into a graph rather than evaporating. The raw material already exists: every duel emits verified (task, answer, verdict) records over documents published *after* the models were trained (§07). Today those records are an audit trail. L3 is the work of keeping them as knowledge.

L4 — Continuous research. **THE HORIZON** Not one-shot reports but standing research objectives that update themselves as new evidence publishes. The need is documented, expensive and unserved wherever evidence is synthesized formally: **64.3% of Cochrane reviews are never updated**, at a median of **57 months** between updates. Living evidence is what those fields openly ask for and cannot staff. The weekly fresh-document ingestion that already feeds the exam is the same pipeline this layer needs; pointing it at standing objectives rather than only at exam generation is the work that remains.

L5 — Machine-to-machine intelligence. **THE HORIZON** This is the second audience named on the cover, and this is the honest statement of where it stands. The largest consumers of research will not be people. Trading agents need macro evidence, coding agents need current API behavior, biotech agents need last quarter’s findings — and none of them can safely act on an answer that cannot be checked. Everything beneath this layer is already built for them: mechanical grading, pinned sources, replayable verdicts, structured output. What is **not** built is the public agent-facing endpoint itself. L5 exposes verified, cited, machine-readable research through an API — **intelligence infrastructure for AI agents, not a consumer app** — and that API is roadmap, not a shipped product. The design goal is present from the protocol up; the door for agents to walk through is not open yet. For scale only, third-party forecasters put the general AI-agents market on a ~46% CAGR from \$7.8B (2025) to \$52.6B (2030), with “knowledge and research assistant” its fastest-growing segment — vendor forecasts Epago has not independently verified, and not a market Epago sells into today. They describe the tide, not the boat.

What the horizon does not require — and that is the point. No layer above changes how validators agree, how a verdict is computed, or what earns emissions. L1 and L2 are generation contracts: new corpora and new *mechanically gradeable* task families over the identical protocol. L3, L4 and L5 are products built on records the protocol already emits. A layer that needed a new consensus, scoring or judging scheme — an LLM referee, a vote on what is true — would be a different protocol, and Epago would not ship it.

Citation fabrication: Gravel et al., *Mayo Clin Proc Digit Health* 1(3):226–234 (2023) — 41 of 59 references fabricated across 20 medical questions — and the GPT-4o citation audit in *JMIR Mental Health* 2025 (mental.jmir.org/2025/1/e80371/); both detailed with the rest of the reliability evidence in §11. Cochrane update lag: 64.3% of reviews never updated, at a median of 57 months between updates — PMC12362767 (2025) and PMC11795965. The AI-agents market range (\$7.8B in 2025 → \$52.6B in 2030 at ~46% CAGR, “knowledge and research assistant” the fastest-growing segment) is an unaudited third-party vendor forecast (MarketsandMarkets; Grand View Research) reproduced for scale only — Epago has not verified it and forecasts no revenue from it. As-of 2026-08-17, L0 and the live generation of L1 are the only parts of this stack that exist as running software.

06 How the network scales: four axes out of one generation

The machinery is live; the destinations are not. L1 expands along four independent axes, and each one is a contract change over the protocol of §04 rather than a new protocol. The live generation is the worked example of all four; every other row, rung, modality and corpus below is roadmap.

Axis 1 — a domain is configuration, not code. Epago is organized around *chain generations*. A generation is a contract that pins one corpus snapshot, one set of task templates, and the base architecture — and nothing in the protocol below that contract is specific to any field. The exam generator, the duel harness, the statistics, and the coronation rule are domain-blind. Standing up case law and contracts, patents and prior art, regulatory filings, or the scientific literature at large is therefore a contract change rather than a rewrite, and generations run **in parallel**, each with its own king, its own challengers, and its own share of emissions.

Generation	What the corpus contains	The work it grounds
Scientific literature	Peer-reviewed papers, trials, published evidence syntheses	Evidence extraction and synthesis across every empirical field
Law	Case law, statutes, contracts, filings	Authority retrieval and citation-checked drafting
Patents	Patent corpora and prior art	Prior-art search and invalidity analysis
Regulatory	Submissions, guidance documents, labels, safety reports	Dossier assembly and compliance evidence
Financial disclosure	Filings, prospectuses, audited statements	Tracing a stated figure back to the document it came from

Illustrative candidate generations, not a committed schedule, not an ordering and not a revenue forecast. Each row is a contract — corpus snapshot, task templates, architecture pins — over the same protocol, and sequencing is community-governed.

Axis 2 — the task ladder. Grounded deep-research retrieval — find the one source the question refuses to name, reachable only by combining what two other sources say — is rung one, chosen because it is mechanically gradeable: every answer key is proven unique against the corpus before the question is ever written. Above it sit eligibility screening, cross-source synthesis, quality and bias appraisal, and finally drafted sections of a review — the same agentic rollout, progressively harder scoring. Each rung raises what one query is worth without changing the mechanism that selects the model.

Axis 3 — modality. Most quantitative results do not live in prose; they live in **tables and figures**, where every current model is weak. Table and figure extraction is an open capability frontier, and a defensible one for Epago specifically, because a pinned corpus lets the exam score it exactly rather than by human preference.

Axis 4 — corpus reach. The genesis generation scores over a pinned, content-addressed corpus so that every validator poses a byte-identical exam and verdicts agree within a measured noise floor. The behavior that wins there is **retriever-agnostic**: swapping in a live index or the customer's own private document store is plumbing, not retraining. That is the path from a benchmark-grade model to a model reading an institution's own documents, inside its own perimeter — which is also the only place a regulated buyer can run one.

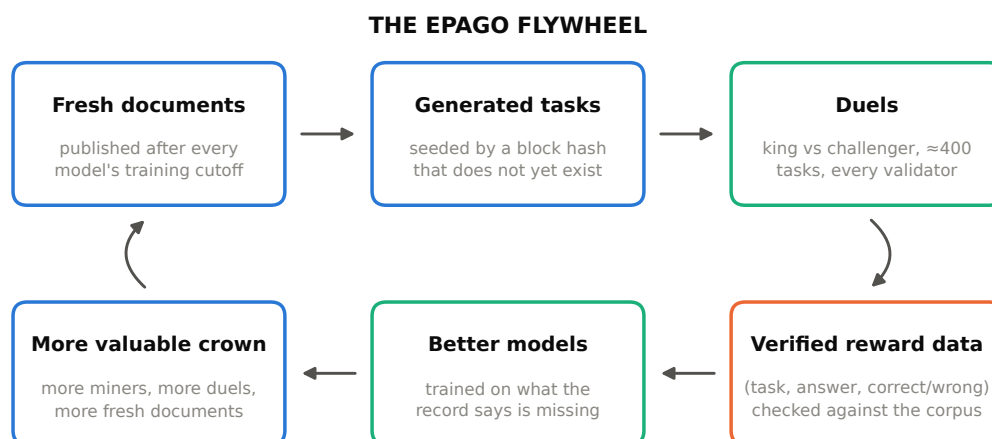
Where the axes lead. Axis 1 is L1 in §05; the upper rungs of Axis 2 are where L2's citation and provenance checks attach; Axis 4 — live and private corpora — is the precondition for research that keeps itself current (L4) and for serving other agents from a customer's own evidence base (L5). The stack and the axes are the same expansion seen from two directions.

07 The compounding asset: verified reward data

Every duel manufactures the scarcest input in AI. A duel runs both models over roughly 1,000 identical research tasks and records, for each one, a verified triple: the task, the answer given, and whether it was right — graded against an answer key proven to exist in the corpus at generation time, so “correct” is a string comparison against a pre-proven key rather than an opinion. That is **verified reward data**: exactly the commodity reinforcement-learning post-training consumes, produced as a byproduct of running the network rather than bought from a data foundry.

The market for that commodity is already large — and supply-constrained. One frontier lab reportedly discussed spending **more than \$1B a year on RL environments**, and the data foundries are being repriced around it — Surge at ~\$1.2B revenue and Mercor at a \$10B valuation, against Scale AI’s \$29B. Epoch AI’s January 2026 survey names **reward-hacking-resistant verification the #1 bottleneck** in RL environments. Epago’s exam engine produces that verification permissionlessly and continuously, rather than as a six- or seven-figure quarterly consulting contract.

And it cannot be sourced anywhere else. These records exist nowhere but here. They are generated from block-hash entropy *after* weights freeze, half of them drawn from validators’ private holdouts, over documents published after the models’ training cutoffs — so they are simultaneously fresh, adversarial, and verified. A competitor cannot scrape an equivalent archive off the open web and cannot buy one: they would have to stand up the whole machine and run it for a year to accumulate it. Epago publishes its records for auditability, which is the point — the asset is the machine that keeps producing them, and that only grows, never resetting between generations. It is also the substrate of the living knowledge layer: today an audit trail, on the horizon a graph that compounds (L3, \$05).



Every turn adds verified reward data that exists nowhere else — and that no competitor can obtain without running the machine.

The flywheel: freshly published documents seed generated tasks; tasks drive duels; duels emit verified (task, answer, correct/incorrect) records; those records train better models; a better model makes the crown more valuable; a more valuable crown attracts more miners and more duels — which consume more fresh documents. Sources: Epoch AI, “An FAQ on Reinforcement Learning Environments” (epoch.ai/gradient-updates/state-of-rl-envs, Jan 2026) for verification as the leading bottleneck; TechCrunch (2025-09-21) for the reported >\$1B/yr RL-environment budget and the Surge / Mercor / Scale AI comparables. Roughly 1,000 tasks per duel (800 public, 200 private) is the genesis-generation parameter and is configurable per chain generation.

08 Why Bittensor: many researchers, paid only for what can be verified

The question that has to be answered is why any of this needs a chain at all. Because the product is not a model — it is a *market* for research systems, and a market needs a settlement layer that pays for verified work without anyone’s permission and without anyone’s discretion.

The alternative is to trust one centralized research AI. Epago’s answer is to run many. Independent teams stand up independent research systems and compete continuously on the questions that actually decide quality — who is more accurate, who finds better sources, who fabricates less, who investigates deeper — and the network pays only for what it can verify. Emissions buy measured

capability, not claims: no procurement cycle, no launch-day benchmark table, no trusting a vendor's evaluation of its own model.

Three properties are load-bearing, and only a chain supplies all three:

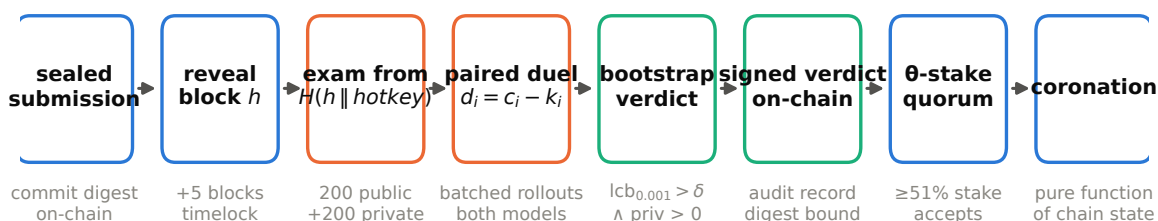
1. **Unforgeable entropy.** The exam is seeded by a block hash that does not exist when the weights are sealed. A company-run leaderboard can promise it did not peek; a chain makes peeking impossible.
2. **Permissionless entry with public accountability.** Anyone with one GPU may challenge the king, a stake-weighted quorum decides, and because every verdict is replayable from public data a validator that grades wrongly is contradicted by anyone who bothers to recompute it — exactly, if it faked the exam, the mathematics or the signature; statistically, against the measured noise floor, if it faked the scores.
3. **A payment rail with no counterparty.** Emissions flow to whoever holds the crown, automatically, whether or not any company still exists to approve the transfer.

The centralized version of Epago does not work. Every anti-gaming property collapses the moment the evaluator and the beneficiary are the same party — which is precisely how one provider came to test 27 model variants privately and publish only the best (\$04). Bittensor's contribution here is not the token; it is that no participant, the Epago Foundation included, can quietly become the judge.

09 How it works — five steps, zero trust

This is the live protocol, end to end. All five steps run today, on real hardware, against real weights.

One duel, end to end — every arrow is deterministic and replayable from public data

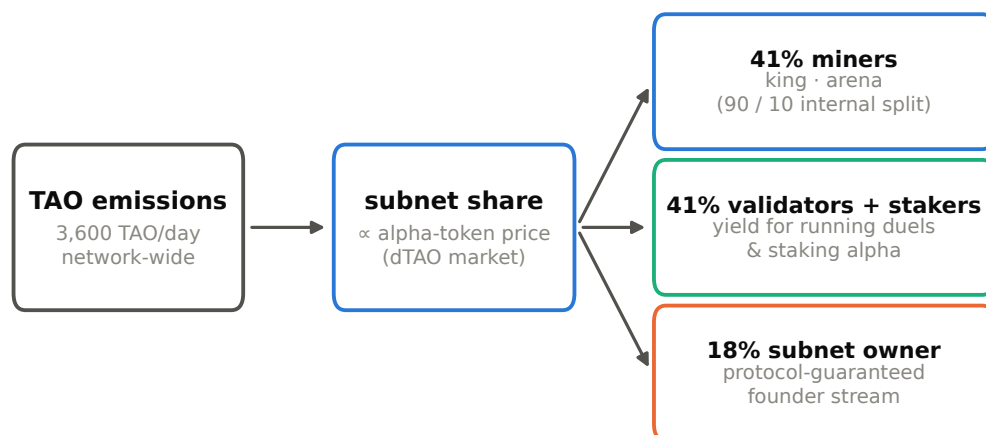


1. **Seal.** A challenger commits its model's fingerprint on-chain; weights are frozen before anyone — including the challenger — knows the exam.
2. **Reveal & generate.** Five blocks later, the reveal block's hash seeds the exam generator. Fresh questions, reproducible by anyone, forever.
3. **Duel.** Every validator independently runs king vs challenger on $\approx 1,000$ identical research tasks — 800 public, 200 from that validator's secret pool. Every task obeys one rule: *hard to find, easy to check* — the question never identifies its sources, and the answer is located by elimination or computed across documents, then graded against a key re-derived mechanically from the corpus.
4. **Prove.** Improvement must clear a 99.9% statistical lower bound above an adaptive floor — luck, copies, and hardware noise are priced at zero.
5. **Crown.** When accepting validators pass 51% of stake, coronation happens as a pure function of chain state. The new king earns; the arena keeps paying strong runners-up so competition never dies.

Why this is hard to copy: the moat is not code (it is open source) — it is the *composition*: post-commitment exam entropy, per-validator secret holdouts with delayed publication, noise-clamped statistical acceptance, and quorum coronation, each closing one cheating channel. Removing any one reopens an exploit; shipping all of them is a protocol, not a feature.

10 Token economics: built on Bittensor's dTAO

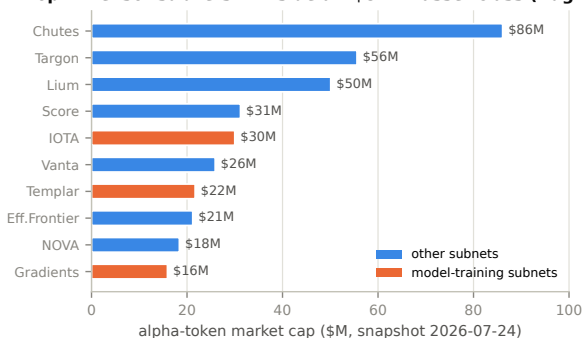
Epago runs as a Bittensor subnet with its own **alpha token**. Since the February 2025 Dynamic-TAO upgrade, the market — not a committee — allocates TAO emissions: each subnet’s share is set by its alpha token’s AMM price. Investors gain exposure by staking TAO into the subnet’s pool (or via CEX-listed alpha tokens; Kraken already lists several subnets).



The asset class is real — and early. Subnet alpha tokens total **~\$872M** (August 2026; ≈\$1.5B at the early-2026 peak) across ~117 tracked subnet tokens in a network of ~128 subnets. TAO itself: **\$2.21B market cap**, hard-capped 21M supply, 11.24M circulating, first halving completed December 2025 — which cut network emissions roughly in half, from ~7,200 to ~3,600 TAO/day.

Institutional rails are being laid: Grayscale filed an S-1 on 2025-12-30 to convert its Bittensor Trust into a spot ETF trading as **GTAO on NYSE Arca**; Bitwise filed a parallel TAO ETF; both track an SEC decision window of roughly August 2026. A staked-TAO ETP (Safello / Deutsche Digital Assets, ticker **STAO**, 1.49% fee) has traded on the SIX Swiss Exchange since 2025-11-19. Note the asymmetry: **the spot-ETF filing states it may not stake its TAO** — passive institutions get price exposure but forgo the emissions Epago participants earn.

Top Bittensor subnets — inside a ~\$872M asset class (Aug 2026)



Sources (as-of 2026-08-17): CoinMarketCap — TAO \$2.21B market cap, ~\$196.68, rank #32, 11.24M of 21M circulating (coinmarketcap.com/currencies/bittensor); CoinGecko Bittensor-subnets category for the ~\$872M aggregate (coingecko.com/en/categories/bittensor-subnets). The two trackers disagree on circulating supply — CoinGecko showed ~9.6M TAO / ~\$1.89B the same day — so read TAO’s cap as a \$1.9–2.2B range. The subnet-by-subnet chart is a 2026-07-24 snapshot. Bittensor docs for the 41/41/18 dTAO split and the December 2025 halving (bittensor.com/docs/concepts/emissions); Grayscale S-1, SEC EDGAR (filed 2025-12-30); Bitwise TAO ETF filing; Safello / Deutsche Digital Assets STAO ETP on SIX (etfexpress.com, listed 2025-11-19); Kraken subnet-token listings; Yuma “State of Bittensor” via The Block.

11 Competitive positioning

Epago is not competing to be the best research model. It is competing to be the market that selects one — which is why the comparison below is against categories rather than against a single product.

Category	Players	Epago's edge
Open deep-research model releases	Labs and teams that publish open agentic-research checkpoints (Epago's own genesis model among them)	They ship a static checkpoint: it is frozen at release, and no one can prove the next one is better rather than benchmark-tuned. Epago starts from that same class of model and adds the only mechanism that makes it compound — and the only proof of each step
Hosted deep research	OpenAI, Google, Perplexity, xAI (\$20–300/mo)	Open weights, self-hostable inside the data perimeter — serving the 72% of enterprises planning at-scale edge or on-prem AI by 2028 (Deloitte, US firms \$500M+), and keeping regulated data out of the GDPR special-category transfer exposure and the EU AI Act high-risk logging burden a hosted endpoint imposes
Evidence-extraction & literature tools	Covidence, Rayyan, Elicit, DistillerSR, Silvi.ai (cloud SaaS; ~\$100–2,000+/yr per seat or per review)	All are hosted workflows wrapped around human screeners or a closed third-party LLM API; every AI-extracted field still needs manual accept/reject, and none ships an open self-hostable model. Epago ships the grounded extraction model itself, runnable inside the customer's own data perimeter, every answer traceable to a pinned source
Open deep-research software	Agent frameworks & research pipelines (25k+ GitHub stars on the leader)	Those are <i>harnesses</i> that call someone's model; Epago produces the model itself — and pays for its improvement
Static leaderboards	Public benchmark suites & arenas	Contamination-proof by construction (exam generated after weight freeze) and replayable — not a static suite that a fresh rebuild deflates by up to 8 points (GSM1k), nor a human-preference arena shown gameable by best-of-N private submission (Leaderboard Illusion, 2025) — properties a calendar or a private test server cannot offer
Model-training subnets	Pretraining and fine-tuning marketplaces on Bittensor	Complementary layer: Epago decentralizes selection — provable head-to-head improvement of one flagship open model with anti-gaming statistics no peer subnet runs

Category pricing as-of 2026-08-17: covidence.org/pricing (\$339/yr single review, \$907/yr package), elicitor.com/pricing (Pro \$588/yr, Scale \$2,028/yr), rayyan.ai/pricing (free tier capped at 3 projects, no PRISMA diagram), distillersr.com (enterprise quote). The domain-specialist segment is still small: Silvi.ai, founded 2019, reports ~\$770K revenue in 2025 with a team of 7 (getlatka.com/companies/silvi.ai).

What source-traceable research costs today. Wherever a claim has to be tied back to a document, the work is slow and expensive, and the best-measured example is formal evidence synthesis: one systematic review of the literature costs a mean of **~\$141,000 in labor** (1.72 scientist-years) and takes a mean of **67.3 weeks — about 15 months** — to complete and publish, and roughly **29,000** are published a year, up from ~4 a day in 2000 to ~80 a day in 2019. Machine assistance already compresses part of it: screening time fell by **more than half in 17 of 25 studies**. Those figures are an illustration of what checkable automation is worth in any evidence-bound field — not a revenue forecast and not a statement of which field Epago serves.

\$141k

mean labor cost
of one review

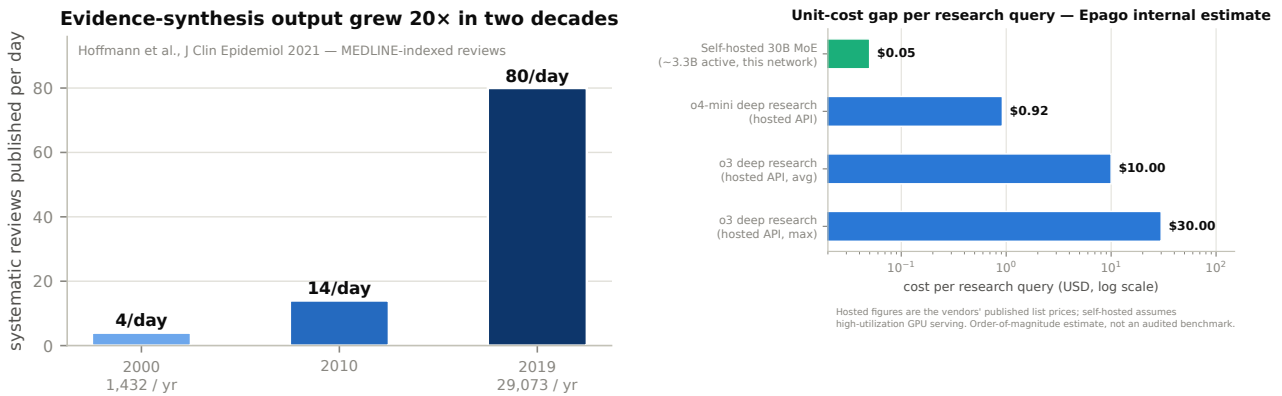
~15 months

mean time to
complete and publish

>50%

of screening time
already cut by AI

Willingness to pay for machine research is already established. Every major lab sells deep research as a premium feature — OpenAI at \$20–200/mo, Google at \$5–200/mo, xAI at \$30–300/mo — and the specialist tools above sell a comparable workflow at \$339–2,028/yr per seat or per review. What none of them sells is a model the buyer can run, pin and audit for themselves.



Left: published evidence-synthesis output, Hoffmann et al., *J Clin Epidemiol* 2021 (MEDLINE-indexed systematic reviews). **Right — Epago internal estimate, not an independent measurement:** hosted costs are the vendors' published list prices for deep-research API calls as-of 2026-08-17; the self-hosted figure assumes high-utilization GPU serving of the genesis MoE, whose ~3.3B active parameters per token give it a small-model serving footprint. Treat the ratio as an order-of-magnitude engineering estimate pending third-party benchmarking, not as a verified market price.

And the hosted model has a structural gap. Data privacy is the **#1 named barrier to enterprise AI — 53%, ahead of legacy integration (40%) and cost (39%)** (Cloudera, ~1,500 IT leaders across 14 countries, Apr 2025). Deloitte expects at-scale edge or on-premise AI to reach **72% of enterprises by 2028, up from 36% in late 2025** (US director-and-above executives at \$500M+ firms). NTT DATA's May 2026 study of ~5,000 leaders finds **>95% call private or sovereign AI important while ~60% are blocked by cross-border data restrictions** — barriers sharpest in regulated settings. A hosted-only research agent cannot serve them. An open 30B model that runs inside the customer's perimeter, over a pinned and auditable corpus, can.

The deeper gap is not privacy — it is fabrication. Parametric chatbots invent their evidence. In one audit of medical questions, **69% of the references ChatGPT supplied did not exist**, yet 95% named real authors — fabrications engineered to survive a glance. A 2025 audit of GPT-4o across six literature reviews found **~56% of citations fabricated or erroneous** (19.9% wholly invented, 45.4% of the genuine ones wrong), with 64% of fabricated DOIs resolving to real-but-unrelated papers. Practitioners report the consequences directly: in a survey of clinicians, **91.8% had encountered model hallucinations and 84.7% judged them capable of patient harm**. Grounding is what fixes it — constraining models to a retrieved source set lifted accuracy **71.4%→90% for GPT-4 and 58.6%→92.9% for Claude 3 Sonnet**. This is the thesis restated as a failure mode: a model asked to recall fabricates, a model asked to *read* can be checked. Epago's task definition scores for exactly that, and a hosted parametric assistant cannot certify it. It is also the exact failure L2 exists to attack, and the reason that layer is worth building (\$05).

The regulated-buyer wedge hosted APIs structurally cannot serve. The GDPR treats special-category data under Art. 9 with penalties to **€20M or 4% of global turnover**, and under EDPB Opinion 28/2024 a model trained on personal data is not presumed anonymous. The EU AI Act classifies AI deployed in regulated high-risk settings as **high-risk**, adding mandatory logging, human-oversight and traceability duties with fines to **€15M or 3%**. And the reporting standards for evidence synthesis — PRISMA 2020 and the 2025 RAISE statement — require a review to be reproducible down to the exact model version. An opaque hosted endpoint fails

every one of these; an open, self-hostable, version-pinned, replayable model passes them by construction.

Workflow and volume: Hoffmann et al., *J Clin Epidemiol* 2021 — 29,073 systematic reviews indexed in 2019 vs 1,432 in 2000 (~80/day vs ~4/day), [jclinepi.com/article/S0895-4356\(21\)00174-8/fulltext](https://jclinepi.com/article/S0895-4356(21)00174-8/fulltext). Michelson & Reuter, *Contemp Clin Trials Commun* 2019 — \$141,194.80 mean academic cost, 1.72 scientist-years, PMC6722281. Borah et al., *BMJ Open* 2017 — 67.3-week mean across 195 PROSPERO-registered reviews, PubMed 28242767. AI screening-time reduction >50% in 17 of 25 studies, up to 99.6% (264.5 hours saved) on extraction — *Front Pharmacol* 2025, doi:10.3389/fphar.2025.1454245. Reproduced as a measured illustration of what evidence-bound research costs; Epago forecasts no revenue from these figures.

Privacy and on-premise demand: Cloudera survey of ~1,500 IT leaders in 14 countries (2025-04-16); Deloitte AI-infrastructure survey (Q4 2025 — US respondents, \$500M+ revenue firms); NTT DATA study of ~5,000 leaders (2026-05-14); Broadcom/Radius (Feb-Mar 2026: 56% run or plan production AI inference on private cloud vs 41% public, 83% considering repatriation).

Fabrication and grounding: Gravel et al., *Mayo Clin Proc Digit Health* 1(3):226-234 (2023) — 41 of 59 references fabricated across 20 questions; Walters & Wilder, *Sci Rep* 13:14045 (2023) — corroborating citation-fabrication audit; JMIR *Mental Health* 2025, mental.jmir.org/2025/1/e80371 — GPT-4o citation audit; arXiv:2503.05777 — clinician hallucination survey (91.8% / 84.7%, n = 70); PMC11898693 — retrieval-grounding accuracy gains, both significant at p = 0.001.

Regulatory: EU AI Act Art. 6 and Art. 99, Reg. (EU) 2024/1689 (artificialintelligenceact.eu/article/99), GPAI obligations in force since 2025-08-02; GDPR Art. 9 (gdpr-info.eu/art-9-gdpr); EDPB Opinion 28/2024 (2024-12-17); PRISMA 2020 (PMC8005924); RAISE statement, Cochrane/Campbell/JBI/CEE (2025-11-11).

Defensibility: the network effect compounds on three loops — more miners → faster model improvement → more model usage and staking → higher emissions share → more miners. Two assets sit underneath it that a competitor cannot buy: the exam's credibility (unforgeable, replayable) and the archive of verified reward data it accumulates one duel at a time (\$07).

12 Traction and roadmap

Built and running today RUNNING TODAY — this is not a whitepaper-stage project:

- Genesis 30B MoE model (Tongyi-DeepResearch-30B-A3B, ~3.3B active) live; full challenger training → submission → duel pipeline exercised end-to-end on real hardware
- Public dashboard & leaderboard exported purely from replayable audit records — anyone can regenerate it
- Decentralized content-addressed model & state storage (S3-compatible, digest-verified)
- Autonomous validator-in-a-box: one docker compose up -d; CI-built images; fleet auto-updates on every release
- Fully automated private-pool freshness feed (no human curation anywhere in the validator loop)
- 350+ deterministic mechanism tests, adversarial mock-chain soak harness, third-party verdict replay tool

Phase	Milestones
Now — testnet	Multi-validator quorum live on Bittensor testnet; real-GPU duel timing; economic phase-gates in shadow mode
Next — mainnet	Mainnet launch with full emission schedule; CEX/DEX alpha liquidity; first open coronations — the first time an improvement to an open deep-research model is proved in public
Then — generations	Community-governed generation upgrades along the four axes of \$06: higher rungs of the task ladder, table and figure extraction, live and private corpora, and further domains standing up as parallel generations — contract changes, not rewrites, with sequencing decided by the community rather than announced here
Beyond — the network THE HORIZON	The layers of \$05, in order and none of them scheduled here: citation and provenance verification as a new gradeable task family (L2), verified records kept as a knowledge graph rather than an archive (L3), standing research objectives that update themselves (L4), and a machine-readable API serving other AI agents (L5). All are additive — contracts and products over the same consensus, scoring and emission rules

EPAGO

The Open Intelligence Layer for Humans and Agents.

Open source (MIT).

Technical whitepaper available alongside this document.

github.com/EpagoFoundation/epago

August 2026 · v1.3

This document is for information only and is not an offer to sell or a solicitation to buy any security or token, nor investment, legal, or tax advice. Market figures are third-party estimates as dated inline; digital-asset prices are highly volatile and may result in total loss. Forward-looking statements involve risks and uncertainties.