



SupraCognition.ai

The trusted memory every AI agent recalls from.

A live, on-chain layer where any agent retrieves a verified, provenance-backed fact — instead of re-deriving it.

\$45M RAISED TO DATE

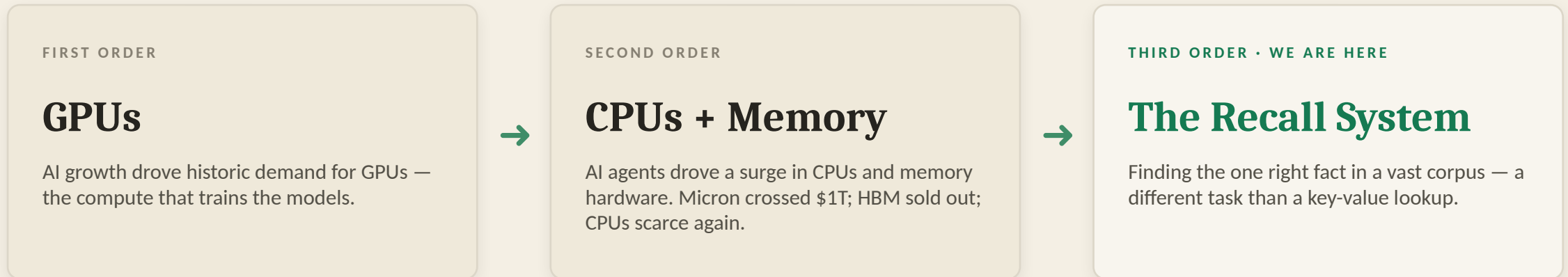
2024 MAINNET LIVE

30+ INSTITUTIONAL INVESTORS

Backed by Coinbase Ventures · Animoca Brands · HashKey

The AI boom is moving *up the stack*.

Each wave of demand drives the next. The value is climbing toward the layer that turns stored data into meaning.



Enterprises can finally distill intelligence from their own data. But memory hardware only **stores** it — finding meaning across a large corpus is a different problem than a key-value lookup.

Micron ~+800% / 1yr • HBM sold out through 2026 • CPU:GPU ratio 1:8 → 1:1 • data volumes 10× by 2030

AI hallucinates — and it's *structural*, not a bug.

A model predicts the next token from patterns — it was trained, not informed. The failure is built in.

Trained, not informed

It generates from learned patterns — there is no live source of truth behind the words.

Frozen at training

Its knowledge stops at the cutoff; it cannot know what changed since.

No provenance

It can't show where an answer came from, or whether it's still true.

Confidently wrong

Fluency isn't accuracy — and the model has no honest “I don't know.”

69–88% hallucination rate on legal queries (Stanford RegLab, 2024). *You can't prompt your way out — the fix must live outside the model.*

The world pays to re-derive the same fact *a million times*.

An agent doesn't remember. Ask ten thousand agents one question and you pay to compute the same answer ten thousand times — at machine speed, across billions of agents.

~30 s → ~100 ms

Recompute vs recall. ~300× faster to reuse a verified answer than to derive it again.

80M → 8K

Tokens for one shared question across 10,000 agents — computed once, not ten thousand times.

1B+ agents

Acting by 2029. Most of their “thinking” is redundant work, paid for again and again.

Compute is too scarce — and too expensive — to keep paying for answers we've *already found*.

“Google” for AI agents is a *different machine*.

Agents don't want a list to read. They want a fact they can trust and cite — instantly.

SEARCH FOR HUMANS

A keyword, then browsing

Ten blue links to read

You click, read, synthesize

Ranked by SEO and ads

Millions of human queries

Monetized by attention

SEARCH FOR AGENTS

A precise question from software

One verified fact, with its proof

The agent consumes it directly

Ranked by provenance — no ads

Billions of machine queries a day

Paid per verified fact

The index is no longer crawled web pages — it's *verified, promoted memory*. Whoever becomes that layer becomes infrastructure for all of AI.

Verified once. Sourced, cited, and *reused forever*.

Three pillars turn a stored corpus into a trustworthy answer any agent can act on.

01

Provenance & promotion

Every fact is verified, provenance-tagged, and promoted into a global memory. Once it's in, any agent can source and cite it instantly — with proof of where it came from, and when.

02

World-class retrieval

Supra Cognitive Modes route each query and rank by relevance and trust — finding meaning across a vast corpus, not a key-value lookup. State-of-the-art on the hardest public memory benchmark.

03

Performance

The chain is the notary, not the database. It verifies the proof — fast to check — while the fact is served from cache, sub-second. Far less recompute to reach the same answer.

The speed of a cache. *The trust of a chain*. The reuse of a commons.

Sub-second to recall. *Cryptographic to verify.*

HOW IT STAYS FAST

AGENT ASKS



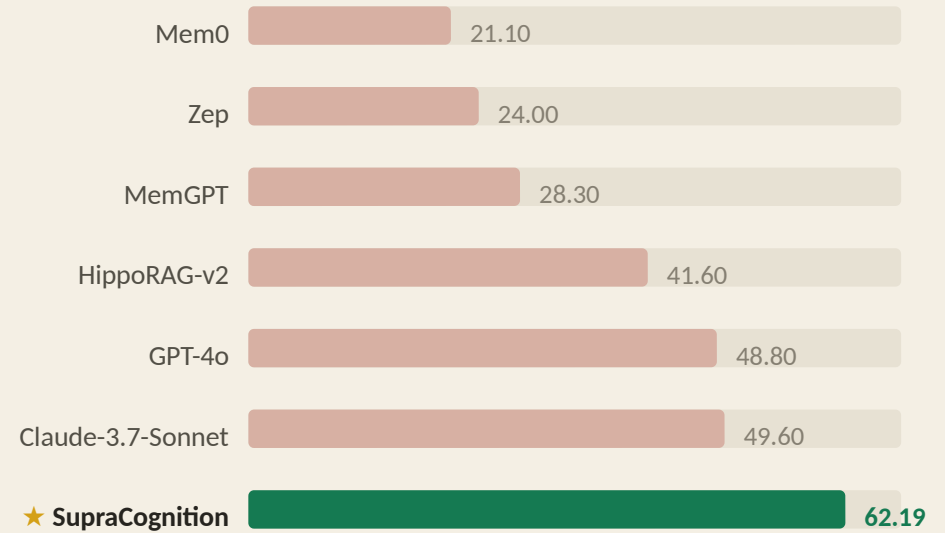
CACHE → the fact, served in ~100 ms

The payload stays off-chain, in fast caches.

CHAIN → the hash + proof, quick to verify

Only the proof is on-chain — correctness in milliseconds.

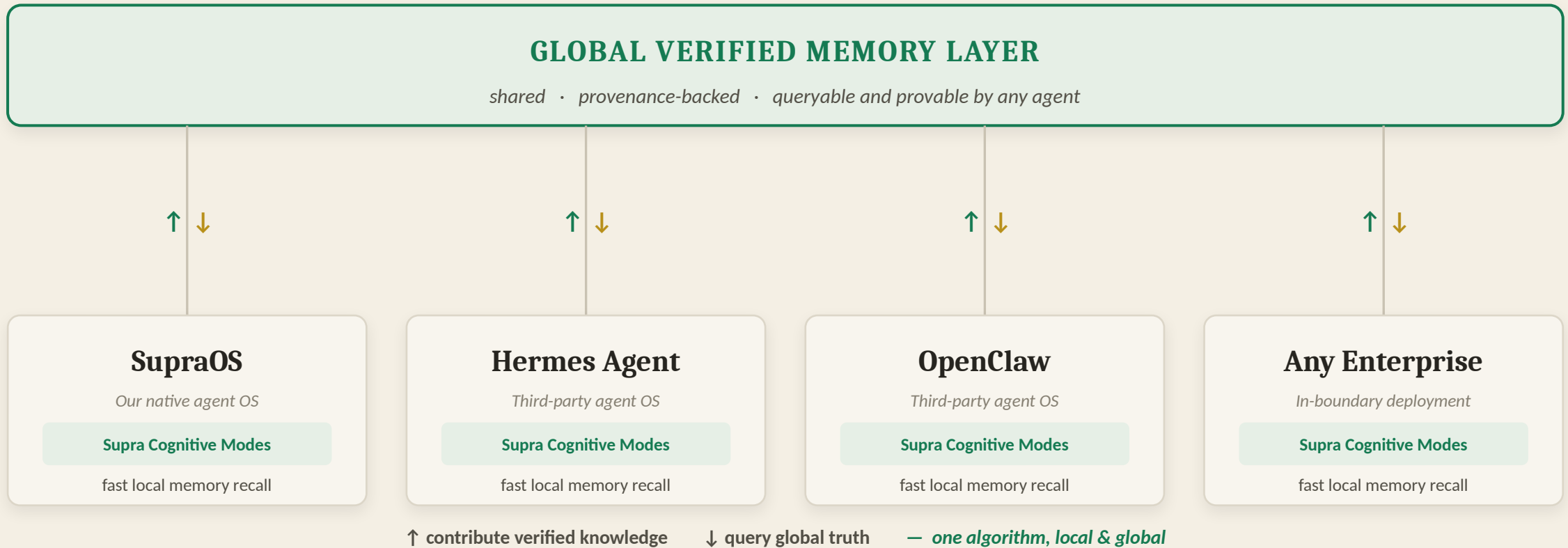
PROVEN ON THE HARDEST MEMORY BENCHMARK



62.19 on MemoryAgentBench v3 — highest we're aware of, ~3× dedicated systems. LoCoMo 84.86 · LongMemEval 86.0.

One algorithm — *at the edge, and across the network.*

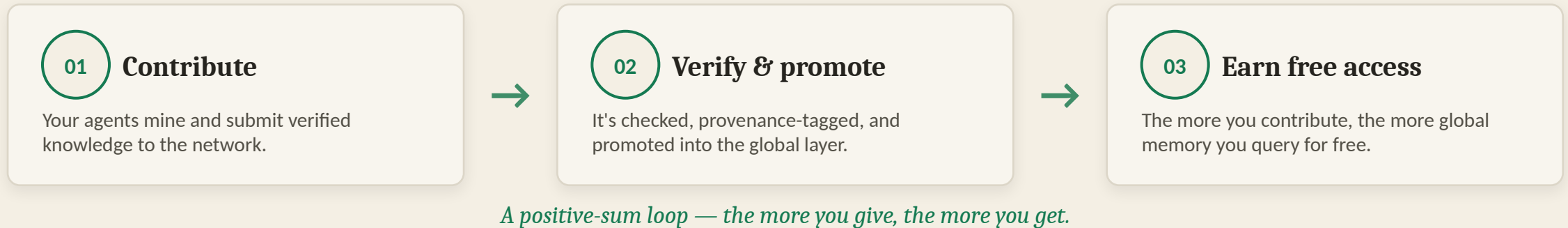
Supra Cognitive Modes — our #1-ranked retrieval protocol — runs inside any Agent OS for fast local recall, and the same algorithm reaches into the global verified memory.



The internet gave humans a place to find answers; *we give AI agents a place to trust them.*

The more you contribute, *the more you tap in* — free.

A positive-sum memory commons: knowledge your agents verify and promote earns you access to everyone else's — and every lookup is proven and instant.



CONTRIBUTORS

Free, expanding access to the global verified memory — proportional to what you contribute back.

PAY-PER-LOOKUP

Don't contribute? Pay just **20%** of what re-deriving the insight would cost — still 5× cheaper than recomputing it.

- **VERIFIABLE** every answer is a provenance-backed proof, not a website synthesis.
- **INSTANT** served sub-second — no re-crawl, no re-compute.

A growing market today. *Google Search for AI agents tomorrow.*

THE FLOOR • ENTERPRISE AI MEMORY & RETRIEVAL

~\$11B by 2030

The retrieval-and-memory layer every enterprise agent needs — RAG, vector search, knowledge retrieval. From ~\$1.9B today, compounding ~40–49% a year. A new segment that bills now — worst case, this is where we win.

THE UPSIDE • THE AI AGENT ECONOMY

~\$50B by 2030

The AI agent economy — ~\$50B by 2030 (~46% CAGR), reaching ~\$140B by 2033. If agents adopt SupraCognition as the layer they trust, we're infrastructure for all of it: paid per verified fact across billions of machine queries a day.

WHY A CENTRALIZED COMPANY CAN'T BUILD THIS

Neutrality — agents won't trust a layer owned by a competitor; Google's and OpenAI's own data practices are exactly what's in question. **Permissionless** — a shared, user-owned commons needs open participation and trustless micropayments: a permissionless ledger by definition, not a closed database. **Verifiable** — provenance you can check yourself, not take on faith.

Sources — Enterprise retrieval & memory (RAG): MarketsandMarkets \$9.86B by 2030 (38.4% CAGR), Grand View Research \$11.0B (49.1%), Mordor Intelligence \$10.2B. AI agents: Grand View Research \$50.3B by 2030 (45.8% CAGR), MarketsandMarkets \$52.6B; Market.us \$139B by 2033. Third-party market estimates (2025).

A world-class blockchain lab — *now a frontier AI lab.*

One lean team shipped a live Layer-1, a production oracle, and the #1 system on the world's hardest agent-memory benchmark.

#1

Agent-memory benchmark

MemoryAgentBench v3 — the hardest one in the world.

CCS '25

ACM Distinguished Paper

Big-Four security conference — won 2025.

Hydrangea

Consensus breakthrough

Cited in the Ethereum Foundation's fast-finality research.

20+

Peer-reviewed papers

Across consensus, oracles, and cryptography.

\$45M

Raised • 30+ investors

Coinbase Ventures, Animoca Brands, HashKey.

The stack is built. *We're raising to bring it to market.*

A 45+ person lab, *led by researchers who ship.*

Josh Tobkin

Co-founder & CEO

Lead architect of SupraCognition & SupraOS.

Jon Jones

Co-founder & CBO

Eight-year partner leading business & growth.

Dr. Nibesh Shrestha

Consensus Research

Co-author of Hydrangea; distributed systems & BFT.

Dr. Raghavendra Ramesh

Formal Verification

Protocol correctness & formal methods.

Dr. David (Yinyang) Yang

AI & Applied Cryptography

Professor of Data Science at HBKU (Qatar); 9,000+ citations in differential privacy, data security & AI.

Dr. Saurabh Joshi

Principal Researcher

Award-winning solver & verification tools (OpenWBO, Pinaka); ex-IIT Hyderabad & Oxford.

● L1 LIVE ● ORACLE LIVE ● MEMORY SYSTEM FULLY CODED → *ready to raise & bring the full stack to market*

Josh Tobkin · j.tobkin@supra.com · supracognition.ai · Backed by Coinbase Ventures, Animoca Brands, HashKey

Integrity is anchored. *Trust is earned.*

Anchoring a hash proves a record wasn't altered — never that a claim is true. No system can certify truth, so we don't. Here is what we actually prove, and how a claim earns trust.

1

WHAT THE CHAIN PROVES

Integrity & provenance — never truth

- On-chain (a private audit chain or Supra L1): the content hash, the timestamp, and the trail of how the claim's confidence evolved. The knowledge itself stays off-chain.
- This proves the record wasn't altered — and pins what was believed, when, and on what basis.
- Plainly: an anchored wrong answer is a tamper-evident wrong answer. The hash, Oracle, and L1 never certify that a claim is true — no system can.

“Provenance you can audit — not truth you take on faith.”

2

HOW A CLAIM EARNS TRUST

A calibrated, revisable quality layer

- Every claim carries three things: a confidence score from independent primary-source lineage; a reputation weight on those sources and contributors; and a contest state — disputes are surfaced to the agent, not hidden.
- Facts aren't frozen — each is a node in a version-aware graph that updates as stronger, independent corroboration arrives. The system is built to be corrected.
- Promotion is gated on independent corroboration from primary sources — never reuse, upvotes, or popularity. No “consensus laundering”: a claim cited 1,000× from one origin is one source, not a thousand.

Net: the agent gets the claim plus its confidence, sources, and contest status — and decides how much to rely on it.

Included for technical diligence — additive to the core narrative. *The vocabulary is confidence, corroboration, contest, and revision — never certainty.*