



THE SECOND-ORDER BOOM OF THE AI ERA

Models commoditized. *Memory* didn't.

The AI boom poured a trillion dollars into models that now converge on the same benchmarks. The durable, defensible layer is forming one level up — in **memory and verified recall**: the system that finds the one true fact in an ocean of data, and proves it.



ACT I

The Boom

A trillion dollars built the model layer. Then the model layer became a commodity.

EVERYONE IS BUILDING ON AI

The fastest capital deployment in tech history.

In three years generative AI went from demo to the largest infrastructure build the industry has ever attempted — and the agents consuming that compute are multiplying.

BIG-FOUR CAPEX • 2026 GUIDANCE¹

\$700B

planned for 2026 alone — more than **2× global telecom** capex, approaching oil & gas.

CHATGPT WEEKLY ACTIVE USERS²

~900M

in under three years — the fastest consumer-tech adoption curve **ever recorded**.

AI-AGENTS MARKET • 2025 → 2030³

7 → 50

\$7.6B today to **~\$50B by 2030** (\$bn), a ~46% CAGR — agents become the dominant consumer of AI.

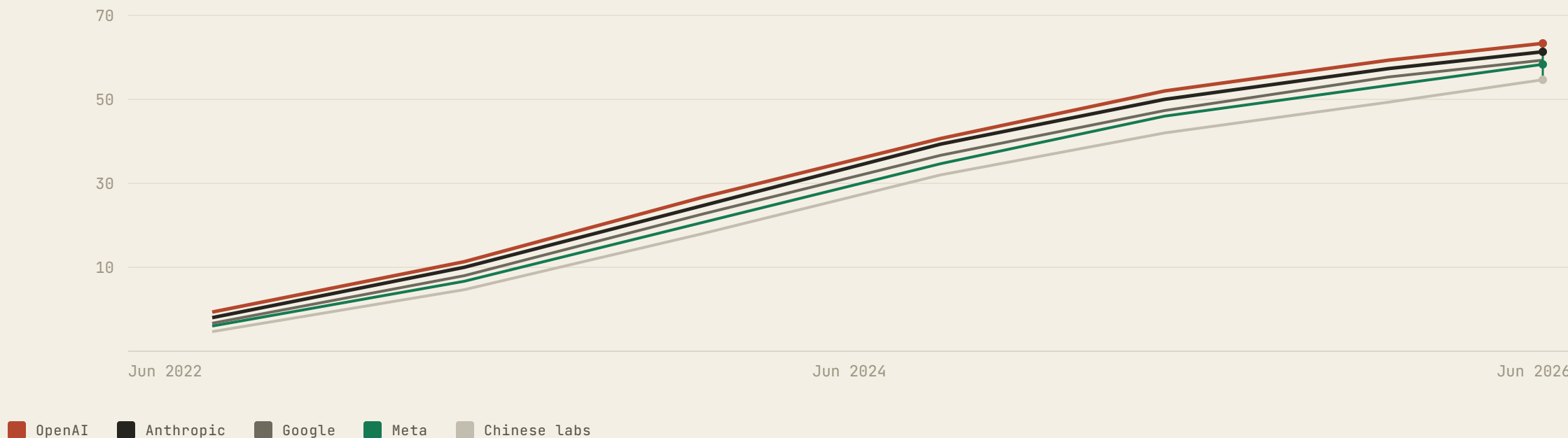
*Every dollar, every user, every agent is pointed at **the model layer**.*

AND SO FAR, MODELS LOOK LIKE COMMODITIES

The frontier is converging — and there are *no network effects*.

Every lab clusters at the top of the same benchmarks within months of each other. For most use, the models are interchangeable. A commodity with no lock-in cannot, by itself, be the moat.

SELECTED FRONTIER LLMS — AGGREGATE BENCHMARK SCORE, JUN 2022 → JUN 2026⁴



If the model is rented and interchangeable, the question becomes: where does durable value go?

Source: ArtificialAnalysis aggregate benchmark index (illustrative trajectories).

ACT II

Memory

The model is rented. The memory is owned. The second-order boom begins here.

THE SECOND-ORDER EFFECT — THE THING MODELS CAN'T DO

Agents can't remember. So a new layer is being funded to fix it.

Models are frozen at training time and forget everything between calls. The fix is a memory layer — and capital is pouring into it. The retrieval market that surrounds it is already **\$1.9B and compounding ~40% a year**, with memory as its fastest-moving frontier.

RETRIEVAL / MEMORY MARKET¹⁰

\$1.9 → 10B

retrieval-augmented generation, 2025 → 2030 at **~40% CAGR** — the layer agents read from.

MEMO · ONE LAYER'S TRACTION⁵

186M

API calls/qtr in 2025 — a **5× jump in two quarters**;
80K+ developers, ~48K stars, millions of interactions daily.

ENTERPRISE VALIDATION⁵

AWS

chose Mem0 as the **exclusive** memory provider for its Agent SDK — the cloud is picking sides.

THE NEW MEMORY STACK⁶ · Memo · Letta (ex-MemGPT) · Zep · LangMem · Cognition — a venture category that did not exist 18 months ago.

WHY NOW • THE MOAT

The model is rented. The *memory* is owned.

As frontier models commoditize, durable value migrates to the **memory-and-verification layer** around them. Models are interchangeable and rented by the token. Accumulated, provenance-tagged memory compounds — and cannot be copied. Whoever holds the trusted, reusable knowledge layer for agents holds the position search held for the human web.

ACT III

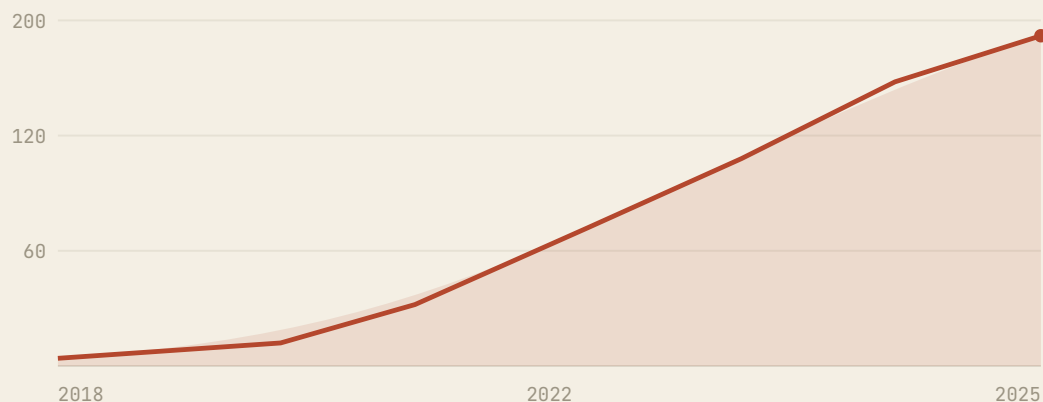
The Haystack

Storing data is solved. Finding actionable insights and distilling meaning across a large memory system is the new bottleneck.

A WORLD THAT STORES MORE THAN IT CAN FIND

181 zettabytes — and 9 of every 10 bytes is a *copy*.

GLOBAL DATASPHERE — DATA CREATED & REPLICATED PER YEAR (ZB)⁷



~80%

unstructured — documents, calls, images, logs. Invisible to a database query.⁸

~90%

replicated. Only ~10% of the datasphere is unique — the rest is the same answers, recomputed and re-stored again and again.⁷

The scarce thing is no longer storage. It's finding — and trusting — the one fact that already exists.

THE MARKET HAS ALREADY REPRICED MEMORY

Enterprises are buying all the memory they can get — *at any price.*

As businesses race to capture and store everything for AI, the hardware that holds it has become the scarcest thing in tech. The storage layer is sold out and repricing — the clearest proof that “keep all the data” is now a board-level mandate.

MICRON (MU) • 12-MONTH RETURN²³

~+900%

memory maker MU ran ~\$94 → ~\$1,000+ and **crossed a \$1 trillion valuation** as AI buys every chip it makes.

HIGH-BANDWIDTH MEMORY • 2026²⁴

Sold out

2026 HBM capacity **fully booked** under binding contracts — Micron, SK Hynix and Samsung alike.

DRAM / NAND / HBM PRICES • 1 QTR²⁴

+60–80%

DRAM & NAND average selling prices rose this much in a **single quarter** (Micron FQ2'26) — a rating analysts call “unprecedented.”

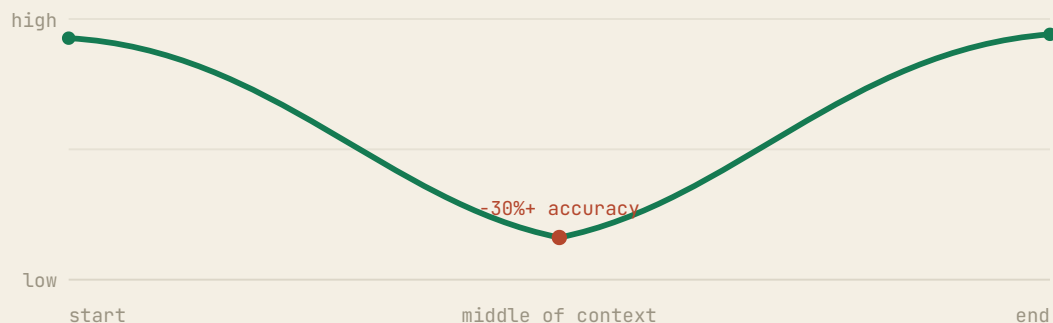
Storing more is now a capitalized boom. Finding — and reasoning over — the one right answer inside it is the next one.

WHY A BIGGER CONTEXT WINDOW ISN'T THE ANSWER

Give a model the whole haystack and it still loses the *needle*.

Stanford and MIT showed a U-shaped curve: models use facts at the start and end of a long context, and miss what's in the middle — even models built for long context. Accuracy falls sharply past ~64K tokens.

RETRIEVAL ACCURACY VS. POSITION OF THE FACT IN A LONG CONTEXT⁹



-30%+

accuracy drop when the answer sits in the **middle** of the window rather than the edges.

64K

tokens — the point where retrieval reliability begins to **break down**, “context rot.”

Scaling the window doesn't find the needle. Recall is a retrieval problem.

THE NEXT-ORDER EFFECT OF THE AI BOOM

It's not a database lookup. It's *reasoning* about the question.

Through 2025, retrieval meant “embed the query, fetch the nearest chunks.” That breaks on any real question. The bottleneck is no longer the vector database — it's the reasoning layer that decides what to pull, resolves conflicts, and synthesizes one trustworthy answer. That layer is a new field, and it is where the next wave of value forms.

YESTERDAY • VECTOR LOOKUP

Fetch the nearest text.

Embed the query, return top-*k* chunks, hope the answer sits in one of them. **One pass, no judgment** — and it fails the moment a question spans more than one fact.

NOW • COGNITIVE RECALL

Reason to the right answer.

Plan the query, retrieve across sources, **resolve conflicts, reason about the core of the ask**, and synthesize a single answer — with proof of where each fact came from.

A brand-new field — and on its hardest test, every existing memory system collapses. One doesn't.

THE COST OF GETTING IT WRONG

Hallucination is a *retrieval* problem — not a generation problem.

Hallucination is now proven mathematically unavoidable inside the model. The only durable lever is outside it: ground every answer in a retrieved, verifiable fact. Done well, that cuts hallucination by up to 71%¹⁵.

GLOBAL LOSSES • 2024¹²

\$67.4B

the estimated cost of AI hallucinations to business in a single year — and rising with adoption.

ENTERPRISE DECISIONS¹³

47%

of enterprise AI users made **at least one major decision** on hallucinated output (Deloitte).

HIGH-STAKES DOMAINS • LEGAL¹⁴

69–88%

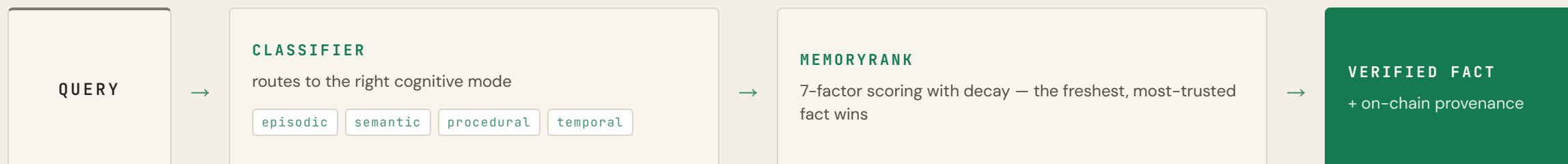
hallucination rate on legal queries (Stanford) — precisely where a wrong answer is unaffordable.

*The fix isn't a smarter guess. It's a **verified fact**, retrieved — with proof of where it came from.*

SUPRA COGNITIVE MODES

Route to the right memory. Rank it. Recall it — *with proof.*

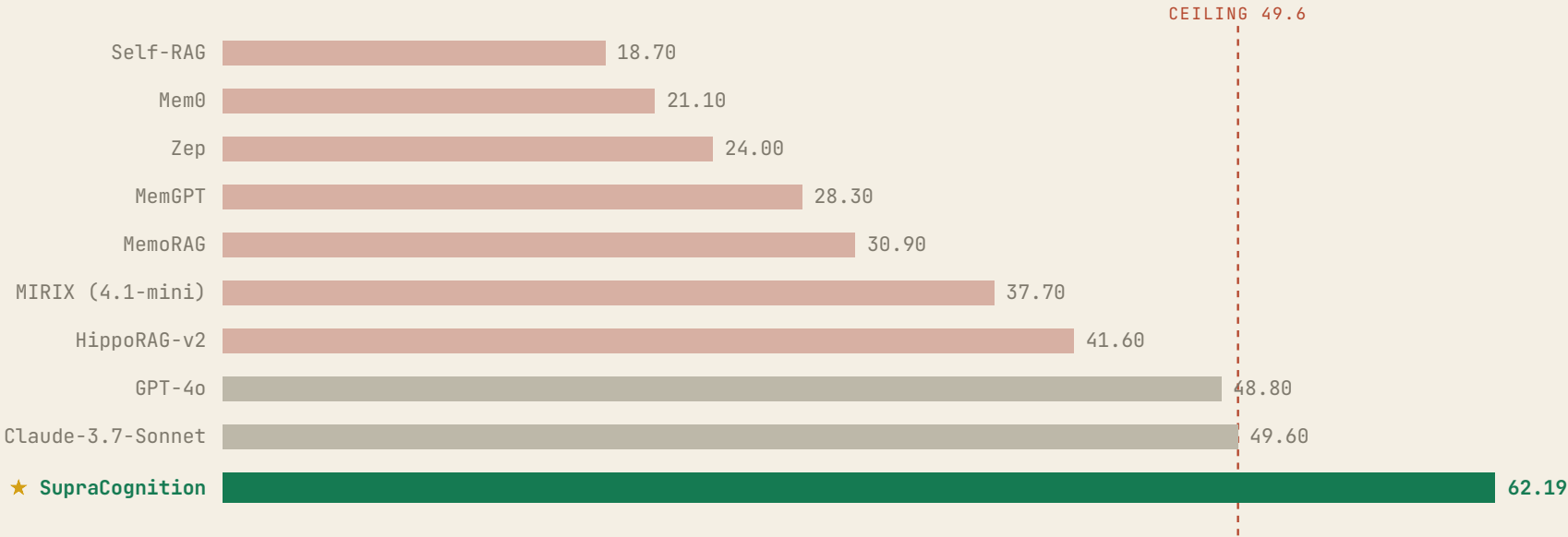
When an agent does real work, the result is stored, scored, and provenance-tagged — so the next agent retrieves it instead of recomputing it. A classifier routes each question to the right cognitive mode; MemoryRank scores and decays facts over time; a hash-chained audit anchors every claim.



*One lookup returns the fact **and** the proof of where it came from, what was known when, and whether it changed.*

MEMORYAGENTBENCH V3 – THE TEST BUILT TO BREAK MEMORY SYSTEMS

Where the field collapses, we *cracked the ceiling*.



LoCoMo 84.86¹⁸ · saturated LongMemEval 86.0¹⁷ · saturated MemoryAgentBench 62.19¹⁶ · the hard one

Saturated tests are table stakes; the field collapses on the neutral hard test. We score **62.19 — above the 49.6 ceiling, ~3× the dedicated memory systems.**

ACT V

The Business

An enterprise product that pays today, and a platform that could own the category tomorrow.

THE WEDGE • ENTERPRISE

Every company is now hoarding data. *Reasoning* over it — privately — is the unlock.

Enterprises are capturing everything for AI: documents, calls, tickets, logs. Roughly 80% is unstructured and most is never used again. SupraCognition lets agents reason over that private corpus to answer the real question — without the data ever leaving the customer's boundary.

PRIVATE BY DESIGN

In boundary

Runs inside the customer's cloud; the corpus never leaves. **Privacy is the enabler** of deep, long-term personalization — not a limit on it.

ANSWERS, NOT HITS

Synthesis

Reasoning across the whole corpus to return **one trustworthy answer** — not ten documents a human still has to read and reconcile.

PROVABLE & AUDITABLE

Provenance

Every answer carries its sources, timestamp, and an audit trail — **essential in regulated industries** where a wrong answer is unaffordable.

*The same ~80% unstructured data that's impossible to query becomes the enterprise's **most valuable asset**.*

ONE ANSWER, COMPUTED ONCE — REUSED TEN THOUSAND TIMES

Today every agent re-derives a fact that *already exists*.

RETRIEVE VS. RE-DERIVE ²²

300×

faster — **~100 ms** to retrieve a verified fact vs ~30 sec to recompute it.

ONE QUESTION × 10,000 AGENTS

80M → 8K

tokens. The same answer, computed **once** — not 80M tokens, ten thousand times over.

REDUNDANT COMPUTE

10,000×

the same work billed again and again, team after team, across the web.

THE BUSINESS MODEL — WE CHARGE 20% OF THE SAVINGS

USERS KEEP 80% — AND ALL THE SPEED

OUR FEE 20%

We capture a fifth of the value we create and hand back the rest — aligned with the customer on every query. With **1 billion+ AI agents forecast by 2029 (IDC)²⁰**, that fifth compounds as every lookup gets cheaper and faster.

HOW WE MONETIZE

Two ways to buy. One *aligned* incentive.

Land with a hosted memory layer the customer runs against their own data; expand into metered retrieval as the agent fleet grows. The same engine, sold two ways — both priced on the value created, not the seats filled.

MANAGED MEMORY • HOSTED

We run it. You consume answers.

A private memory layer deployed in or beside the customer's cloud — we handle storage, scoring, retrieval, and upgrades. Priced on usage and seats: **an ROI line item against token spend**. The enterprise floor.

VERIFIED-LOOKUP API • METERED

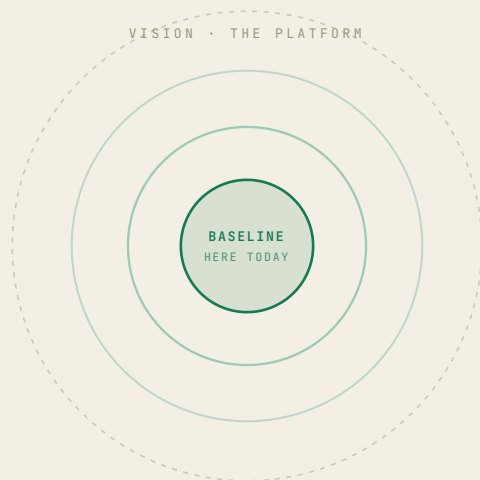
Three lines of code. Pay per fact.

Every agent calls one endpoint and pays per verified lookup. We take **20% of the savings**; the customer keeps 80% and all the speed. Revenue **compounds with agent adoption** — no reselling required.

Addressable today: 1B+ agents executing 217B actions a day by 2029 (IDC)^{2θ} — every action a potential lookup.

THE PLATFORM • THE VISION

Past the enterprise baseline: the infrastructure that *verifies* global knowledge.



- **ENTERPRISE MEMORY • THE BASELINE**
Hosted, in-boundary; agents retrieve instead of recompute. A clear GTM with immediate ROI, billed today.
- **KNOWLEDGE LEDGER**
Private audit chain for provenance & explainability. A risk-reduction sale.
- **KNOWLEDGE FEDERATION**
Privacy-preserving sharing across trusted orgs. A network-effect sale.
- ⚙ **GLOBAL KNOWLEDGE LAYER • THE VISION**
The public, provenance-backed substrate — the “Google search” for AI agents, and our biggest opportunity.

Each private deployment is a knowledge graph. Connected — with permission, provenance, and payment — they become the layer every agent queries and trusts.

WHY THE GLOBAL LAYER MUST BE DECENTRALIZED

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

A shared knowledge layer that every agent reads from and writes to cannot sit inside one company's walls. The hard requirements are exactly what a permissionless ledger provides — and the performance objection dissolves in the architecture: only hashes and metadata live on-chain; data is served from caches with a proof attached.

LOAD-BEARING, NOT DECORATIVE

- | **Permissionless participation** — any agent or org can read and contribute.
- | **Censorship resistance & credible neutrality** — no single owner of the truth.
- | **Trustless micropayments** — pay per verified fact, no trusted operator.
- | **Provenance by default** — every claim anchored, timestamped, auditable.

THE RAIL ALREADY EXISTS

100M+

agentic **x402** payments settled on Base in ~3 quarters — agents already pay per call, on-chain.¹⁹

~2 yr

DORA / Threshold AI Oracle in **production**, securing financial data on this exact hybrid design. The oracle isn't a slide — it's deployed.²¹

Agents are already transacting on-chain at machine speed. A verifiable knowledge layer slots straight into the rail.

A FRONTIER LAB · BLOCKCHAIN, NOW AI

A frontier lab — and the benchmarks *prove it.*

A frontier blockchain lab, now proving frontier AI — prolific research output on remarkably efficient capital, and still inventing. The receipts:

\$45M²¹

raised to date · 30+ investors incl. Coinbase Ventures, Animoca, HashKey

Live

Layer-1 mainnet shipped Nov 2024 · 20+ peer-reviewed papers

62.19¹⁶

MemoryAgentBench — above the 49.6 paper ceiling

~2 yr

DORA oracle in production, securing financial data on-chain

CCS

ACM CCS Distinguished Paper — consensus & verification

217B²⁰

agent actions/day forecast by 2029 (IDC) — the market we index

The enterprise product is the *baseline.* The public network is the *vision.*

A baseline business with a clear go-to-market and immediate ROI today — funding the public verification network that is the biggest opportunity of the AI era.

THE ASK

Raising to scale enterprise deployments and the public verification network. Round details, traction, and unit economics shared on request.

SOURCES & METHODOLOGY

References.

1

Hyperscaler 2026 capex (~\$700B) — **CNBC; Reuters; Goldman Sachs** (Q1 2026 earnings); B. Evans, “AI eats the world,” 2026.

2

ChatGPT weekly active users (~900M) — **OpenAI**, reported 2025–26.

3

AI-agents market (\$7.6B 2025 → ~\$50B 2030) — **Grand View Research; MarketsandMarkets**, 2025.

4

Frontier-model benchmark convergence — **ArtificialAnalysis** aggregate index, 2026 (trajectories illustrative).

5

MemO funding (\$24M), API growth (35M → 186M, Q1 → Q3 2025), AWS Agent SDK exclusivity — **TechCrunch; PRNewswire; Built In SF**, Oct–Nov 2025.

6

Agent-memory vendor landscape (Letta, Zep, LangMem, Cognition) — **AgentMarketCap**, Apr 2026.

7

Global datasphere (~181 ZB by 2025; ~1:9 unique-to-replicated) — **IDC** Global DataSphere, Data Age 2025.

8

Share of data unstructured (~80%) — **IDC / nRoad** analysis, 2022.

9

“Lost in the middle” (>30% mid-context drop) — **Liu et al.**, Stanford/UW, TACL 2024; **Chroma** “Context Rot” (Hong et al.), Jul 2025.

10

RAG market (\$1.9B 2025 → ~\$10B 2030; vector fastest layer) — **Grand View; Mordor; MarketsandMarkets**, 2025.

11

Generative-AI ROI (\$3.70 per \$1 with retrieval) — **Microsoft**, via Mordor Intelligence, 2025.

12

AI-hallucination business losses (\$67.4B, 2024) — **AllAboutAI**, 2025.

13

Enterprise decisions made on hallucinated output (47%) — **Deloitte** Global AI Survey, 2025.

14

Legal-query hallucination (69–88%) — **Stanford RegLab & HAI**, 2025.

15

RAG hallucination reduction (up to 71%) — multiple 2025–26 production benchmarks (incl. **Vectara HHEM**).

16

MemoryAgentBench overall scores; 49.6 ceiling; ≤7% multi-hop — **Hu et al., 2025** (arXiv:2507.05257), Table 1. Venue (ICLR 2026) pending confirmation.

17

LongMemEval — **Wu et al., 2024** (arXiv:2410.10813; ICLR 2025).

18

LoCoMo — **Maharana et al.**, 2024.

19

Agentic x402 payments on Base (100M+ in ~3 quarters) — **Chainalysis**, Jun 2026.

20

AI-agent population & activity (1.15B agents; 217B actions/day by 2029) — **IDC** FutureScape, 2026.

21

SupraCognition company facts (\$45M / 30+ investors; mainnet Nov 2024; ~2-yr DORA oracle; ACM CCS) — **SupraCognition**.

22

Unit economics (300× latency; 80M → 8K tokens; 10,000× redundancy; 80/20 split) — **SupraCognition** model (illustrative).

23

Memory-equity rally (Micron ~+900% in 12 mo, ~\$94 → \$1,000+; \$1T+ valuation; HBM sold out 2026) — **Yahoo Finance; Investing.com; 24/7 Wall St.**, Jun 2026.

24

Memory pricing (DRAM ASPs +~65%, NAND +~75% sequentially) — **Micron** fiscal Q2 2026 results, Mar 2026.

METHODOLOGY & NOTES

Benchmarks. MemoryAgentBench backbones vary across systems; SupraCognition’s 62.19 exceeds the highest overall result reported in the paper. On multi-hop conflict resolution, every system tested collapses to ≤7%. LoCoMo and LongMemEval are saturated; vendor self-reported figures (e.g. MemO 92.5% / 94.4%) are disputed between vendors — SupraCognition reports its own measured 84.86 / 86.0.

Illustrative figures. The convergence trajectories (ref 4) depict the shape of the published index, not exact values. Unit economics (ref 22) are per identical answer: ~100 ms cached retrieval vs ~30 s re-derivation; one shared answer vs 10,000 independent recomputations. Market sizes are third-party analyst estimates and vary by source.

Verification. All external figures were independently checked against published 2026 sources during fact review. MemoryAgentBench is cited to its arXiv preprint (2507.05257, 2025); confirm final conference venue before external distribution. SupraCognition’s own scores (62.19 / 84.86 / 86.0) and company facts (raise, mainnet, oracle) are internally measured and not third-party verified. SupraCognition.ai • June 2026 • Confidential. External figures independently verified against published sources, June 2026; analyst forecasts subject to revision.