

SEED ROUND · AUGUST 2026

Sapience Labs

*Continual learning, **for science.***

*How fast **could** science move?*

Discovery runs as a system of learning loops, and we ask that question of the whole system, not one tool inside it. Every frontier model is trained once and then frozen; we build the layer that keeps learning from your work, in your control, across every AI you use, closing the roughly 72 percent of discovery time that passes before the right knowledge reaches the right person, in academia and industry alike.

The intelligence that keeps learning.

SAPIENCE LABS · BUILT ON THE OPEN CORETX KNOWLEDGE-OBJECT STANDARD

Live · **brain.coretx.ai** · beta cohort forming

7 papers, five on arXiv

seven patent families

THE ARC · A METHOD THAT CROSSED FROM COSMOLOGY TO INDUSTRY · WE FIND CROSSINGS LIKE IT

1,120 papers · one corpus · every line a dependency

Researchers told us why they hold back. We built to those reasons.

Two years of conversations with working scientists. The same reasons came up again and again.

THEY SAID Its memory is a profile of my preferences, not the state of my work. Every session still starts from scratch on the things that matter.

What if it knew what you settled last Tuesday, and why: the decision, the evidence, the alternatives you already ruled out?

THEY SAID Everything I teach it trains someone else’s model, for free, in their walled garden, with my name stripped from it.

What if what you taught it stayed yours: private by default, portable across every AI, credited wherever it traveled?

AND THE WISH NOBODY THINKS TO SAY OUT LOUD

When you choose: what if two people could think together through their AIs, and what one lab learns reached the person stuck on it in days, not a decade?

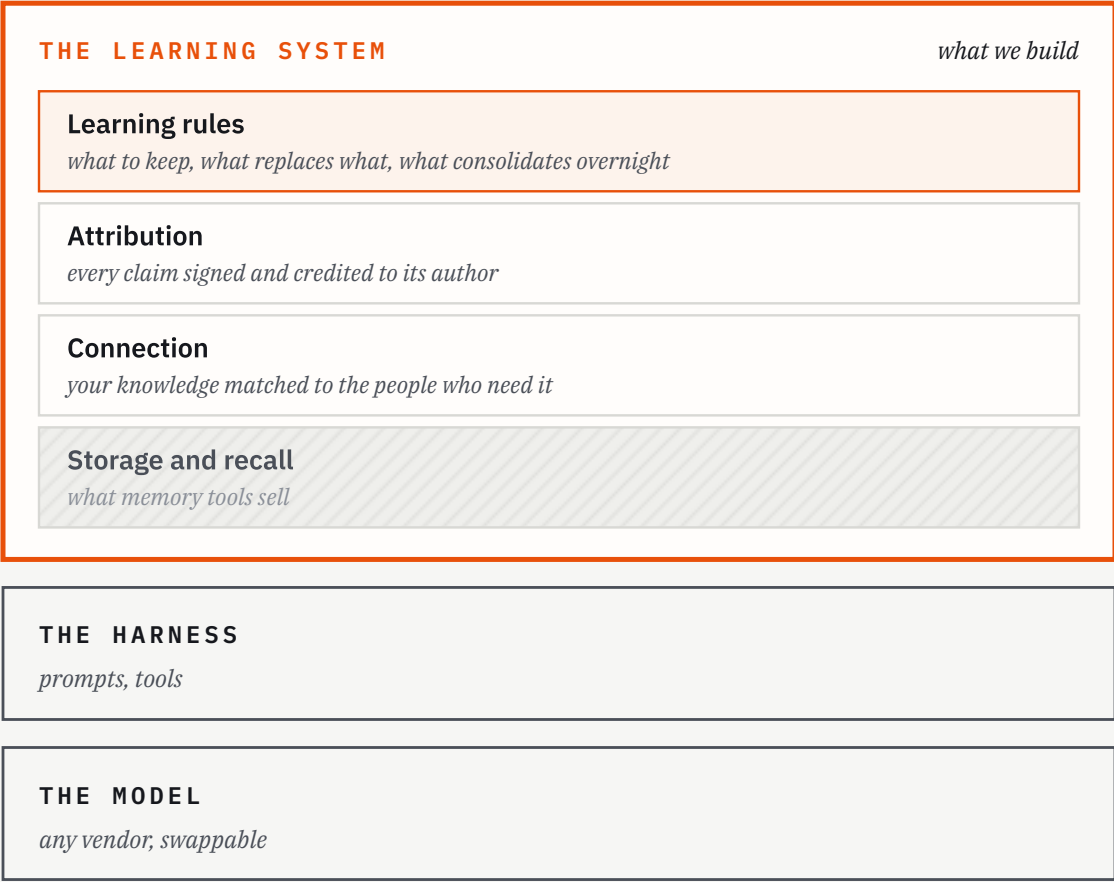
What if none of this were a promise?

Live today at brain.cortex.ai. The rest of this deck is the evidence. One of these is still a promise: on most days the person working on the same problem in another field is not in the network yet. That part we can only build with you.

We build the third component of the AI stack, the one that learns while it works.

The model and the harness are the first two components; ours is the one with its own state and its own learning rules. The model stays frozen and swappable, and the system around it learns while it works: a world model of your research.

WHERE IT SITS



Memory tools supply one layer of the component. The rest is the architecture.

Live today, with the honest numbers.

“It’s amazing to have one brain’s memory infuse all my sessions: my Claude Code and web sessions are fully aware of each other now, and they talk to my Codex sessions too, across providers. That brain is constantly compounding.”

ALEX

· PRINCETON

An earlier version of our co-scientist reached researchers in nine university labs; the Sapience beta cohort is forming now. We are customer zero: this raise was run through the product. Overnight consolidation is live, and sharing between two people is live.



USAGE, AS INSTRUMENTED THIS WEEK · PRINTED AS MEASURED

110

on the waitlist (unique real emails)

31

brains with content (at least one note)

12 / 4

connect codes minted / redeemed

16

installs since Aug 19 (distinct tokens; 35 ever)

18

active in the last 7 days (a write, a retrieval, or a new note)

4 / 2

share links active / joined this week

Counts from the hosted app as of Sep 3, 2026; test tokens excluded; installs before Jul 24 are unattributable. “About 100 beta users” is true only as the waitlist; the product numbers are 31 brains with content and 18 active this week. Notes per active brain in the last 7 days: median 937, range 0 to 34,256; the top end is bulk ingestion, not typing.

Small numbers are printed small. Instrumentation started the week of August 31; the first retention curve follows when there are enough weeks to draw one.

Most of the delay between a problem and its solution is a connection left unmade, and the connection is cross-domain.

“...publication has been extended far beyond our present ability to make real use of the record.” VANNEVAR BUSH ,
1945

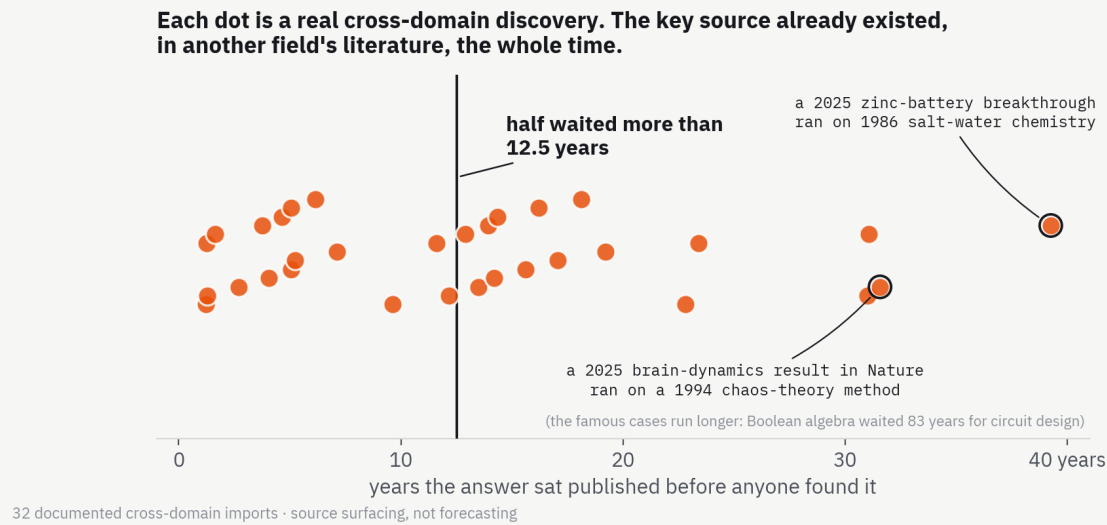
Generation is scaling. Travel is not. Across 48 of 50 documented cross-domain breakthroughs (two excluded for no measurable lag), a median 72 percent of the delay from problem to solution was reducible in principle: the winning approach already existed, in another field, before validation began, and the connection was never made. That 72 percent is the target this company is built toward.

29 / 155

realized imports findable a decade early
(a floor; top-40 slate, abstract-level
ranks)

12.5 yr

median time the true source sat
surfaceable before import, across those
29 cases



SOURCE LEAD TIME · FOR THE 29 OF 155
REALIZED IMPORTS WHOSE TRUE SOURCE SAT IN
THE TOP-40 SLATE, IT HAD BEEN SURFACEABLE
A MEDIAN 12.5 YEARS BEFORE ANYONE
IMPORTED IT; EVERY ONE OF THE 32 IN-
BUDGET ROWS HAD MORE THAN A YEAR OF LEAD.
SOURCE SURFACING, NOT FORECASTING.

Other labs raise on a thesis. Ours is running.

FULL-TIME TEAM

Oliver Zahn PhD, Founder & CEO

Astrophysics PhD work at Harvard. Directed UC Berkeley's Center for Cosmological Physics. Head of Data Science at Google, Senior Data Scientist at SpaceX. Founder & CEO of Climax Foods.

Dr. Simran Chana Surgeon and neuroscientist. Director, Frontier Technologies Lab, Cambridge.	Hashim Piracha Engineering Lead. Machine learning and materials science, UC Berkeley.	Mason Rodriguez Rand Lab Integration Lead. Molecular engineering, UChicago; founded a DOE-funded energy startup.
Senior Research Scientist Google DeepMind. Speak with them under NDA; joins full-time on close.	Senior Research Scientist Google DeepMind. Speak with them under NDA; joins full-time on close.	Senior Researcher OpenAI. Speak with them under NDA; joins full-time on close.

Investors can speak with the incoming frontier-lab researchers under NDA today; they join full-time when the round closes. Six papers, five on arXiv; seven patent families, with the continual-learning mechanisms in filing now.

What sets us apart is the evidence behind the claims: matched protocols, multiple seeds, published kill conditions, and nulls reported on the slide. Diligence gets a reproduction package.

ADVISORS AND INVESTORS

Peter Norvig ex-Director of Research, Google; co-author, <i>AI: A Modern Approach</i>	David Eagleman Neuroscientist, Stanford.	James Evans Science of science, UChicago.
Matias Zaldarriaga Cosmologist, Institute for Advanced Study.	Shervin Pischevar Sherpa Capital; early Uber, Airbnb.	Ravin Kumar Google DeepMind.
Tom Zschach CIO, SWIFT.	Julia Yan Former Head of Growth, TikTok.	Dean Banks Former CEO, Tyson Foods.

The system-level learning space is open now, and this work is measured, filed, and in the field.

The two funded bets will reach for system-level learning once their own pillars mature. Today this space is ours: measured, filed, and in the field. Capital-efficient by design: we do not pre-train, and an 8B open model on the store beats a model 30x its size reading the raw context.

~12% vs ~70%

compute share of this raise vs a pretraining lab's

30x

an 8B open model on the store beats a model 30x its size on raw context

\$3T → \$270B

R&D inside the wedge; enterprise-AI spend at the reach

HOW IT MONETIZES

The wedge is research organizations, entered from the bottom up: researchers adopt the system inside the AI sessions where their work already happens. Free client, Pro and Team seats, Enterprise licenses; revenue starts as subscription, royalties on transacted value later. Signed LOIs with two major scientific publishers sit outside this plan while those partners are paused. The full pricing table is in the appendix.

SOVEREIGN WHERE IT NEEDS TO BE

The entire system can run on customer hardware with no data leaving the building; for regulated industry and enterprise outside the United States that is often the deciding requirement, and no closed-weights lab can offer it.

Sapience turns what we each know
into what we all know. Owned, never
taken.

AI today generates. What it cannot yet do is connect. Sapience is the connective layer: one living memory that connects your sessions to each other, your AIs to each other, you to your AIs - and, with consent, you to other people.

This isn't several products stapled together. The same knowledge object that gives one person's AI memory is exactly the object that lets two people's AIs connect. Build the memory and the network is latent in it. Sharing is live today: connecting two brains takes one sentence and one yes.

Ultimately, the network is **the thing we are building.**

Single-player from day one. Networked by design.

Every model ships frozen. Ours keeps getting more powerful.

The ask. A seed round weighted to senior research and engineering hires, funding three milestones: the first paid cross-organization deployments, the weight path productized at the scale this deck states, and the network reaching self-sustaining density. Sapience is at its most useful on day one, alone; the network is not the price of admission but the compounding on top.

ALLOCATION · 40% engineering & product · 18% research · 22% network & go-to-market · 12% compute & infrastructure · 8% operations & reserve

The core plan is people, on roughly 24 months of runway; we serve open weights rather than pretraining, so **compute is ~12% of the raise rather than ~70%**. Sizing detail in the appendix.

The question this round exists to answer: how fast could science move?

You own your intelligence. Sapience makes it compound.

Oliver Zahn
oliver@coretx.ai

Sapience Labs
*Investors buy Sapience Inc, which owns or exclusively
licenses the architecture filings; CoreTx Inc stewards the
open knowledge-object standard.*

Infrastructure for Humans and AI to Discover Together

Frontier science needs a model trained today, **not last year.**

A model trained on yesterday's internet cannot lead a field that changes today.

"In today's systems, we train them, and then they're kind of frozen and put out into the world. What you'd like is for those systems to continually learn online from experience."

DEMIS HASSABIS · INDIA AI IMPACT SUMMIT · 2026

01 The knowledge that decides the next breakthrough is being created right now, unpublished and never scraped. Frontier models do not contain it.

02 Scientists will not pour unpublished work into a frontier lab. OpenAI launched Prism for scientists in January 2026 and closed it by April over exactly this trust wall (Decrypt).

03 And retraining cannot keep up: a frozen model refreshes once or twice a year; science moves every day.

Sapience resolves both. It learns continuously with no retraining, and attribution at write time lets scientists contribute while keeping ownership and credit. One mechanism serves scientists and builds the pre-publication dataset frontier AI currently lacks.

The reasoning is ours and the model plugs in. Swap in any frontier or open model and the mechanisms are the same, because they live in the system around the core. Store-and-recall is the commodity that memory tools sell. Consolidation, replay, routing, and cross-domain edges are not, and they are what produce the gains on the next slide. The attribution tier at the top of the stack is the one layer biology never needed: it gives every claim a provenance so knowledge can travel between minds with credit intact.

NINE MECHANISMS A MEMORY LAYER DOES NOT HAVE

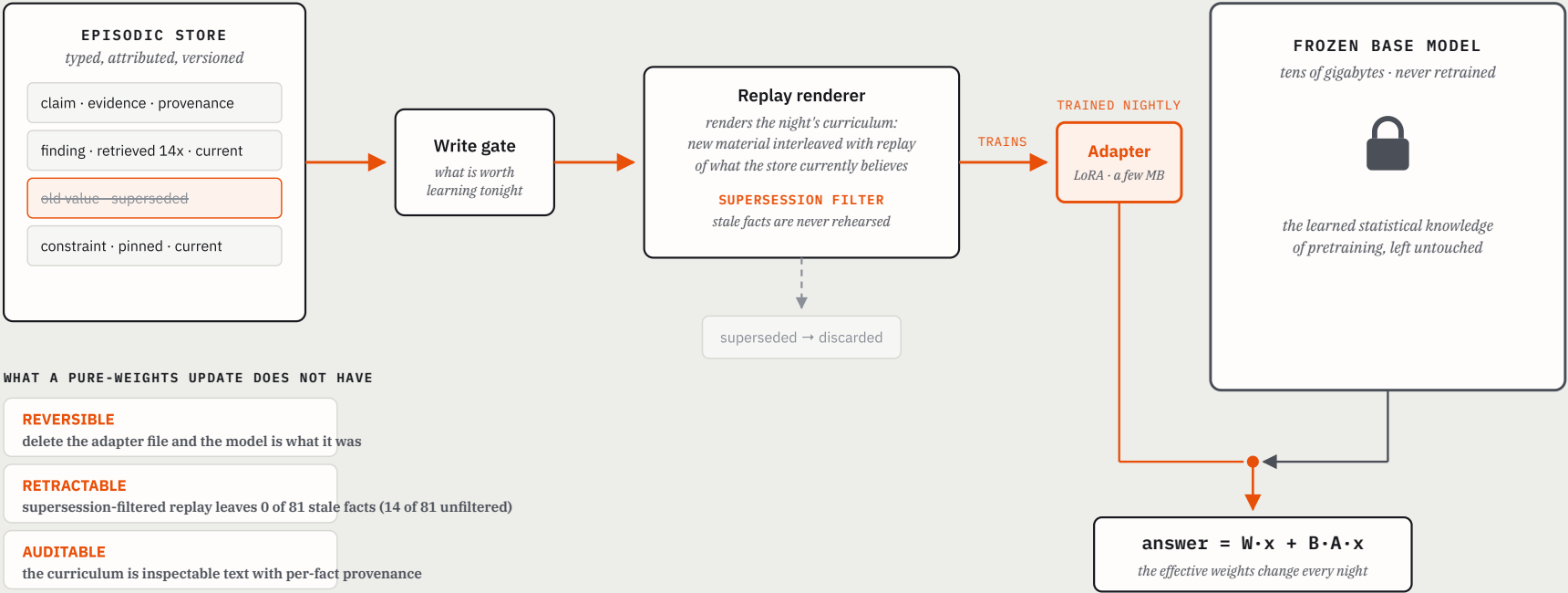
Prediction-error gated encoding writes hard on what is surprising, drops the redundant	Sleep consolidation and replay reorganizes and links what it has seen, offline	Reconsolidation on recall a memory becomes editable when it is used, and updates
Vividness decay and reweighting forgets the stale, keeps what is still in use	Competitive inhibition the strongest relevant memory wins the context	Belief revision contradictions are versioned, never silently overwritten
Contextual reinstatement recalls what was true in the moment a thing was learned	Provenance on every claim source and author travel with the fact	Cross-domain matching connects a result here to a need there, by shared mechanism

NATURE SOLVED CONTINUAL LEARNING TWICE · SCIENCE ALREADY RUNS THIS ARCHITECTURE

A BRAIN	SCIENCE	AN AI THAT JOINS
hippocampus, fast episodes	journals, fast episodes	a store of typed atoms
cortex, slow consolidation	textbooks, slow canon	weights
sleep, replay	review and synthesis	overnight consolidation
forgetting, renewal	retraction and paradigm shift	supersession

The knowledge object itself carries what the mechanisms need: provenance anchors it to reality, typed links carry how it connects, and the record stays current as findings replace one another. The nine mechanisms operate on the object itself (there is no separate database underneath).

ATTRIBUTION	<i>every claim carries its author and source, so credit is auditable</i>	AT SCALE
META - MEMORY	<i>confidence, contradiction detection, validity windows</i>	
PROCEDURAL	<i>gates that decide what to admit, surface, consolidate</i>	
SEMANTIC	<i>consolidated abstractions, schema store</i>	<i>= any frontier or open model plugs in here</i>
EPISODIC	<i>every claim with provenance and session context</i>	INDEXING



WHAT A PURE-WEIGHTS UPDATE DOES NOT HAVE

REVERSIBLE

delete the adapter file and the model is what it was

RETRACTABLE

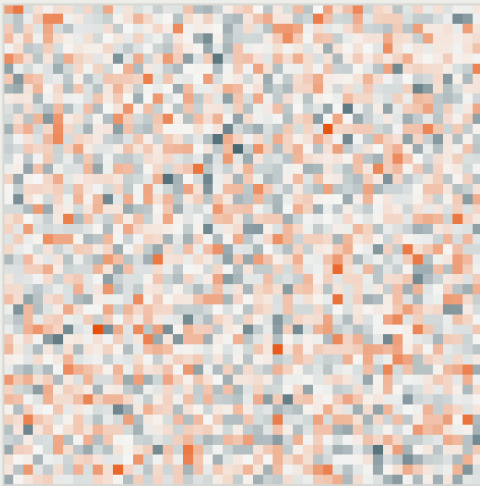
supersession-filtered replay leaves 0 of 81 stale facts (14 of 81 unfiltered)

AUDITABLE

the curriculum is inspectable text with per-fact provenance

Adapter-trained retention **47 to 75 points above the no-adapter control** at matched-or-better acquisition (two model families, 3 seeds, pre-registered; lab result, not shipped).
Replaying structured objects beats the same content as raw prose by **10 to 12 points** at a matched token budget (3 seeds, pre-registered).

THE MATRIX VIEW · THE UPDATE LIVES IN A RANK-64 CHANNEL, BOUNDED AND REMOVABLE

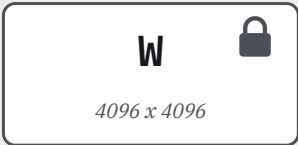


full fine-tuning: dense, every weight moves



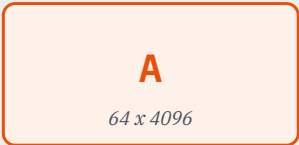
low-rank: a few directions

FROZEN · NEVER TRAINED



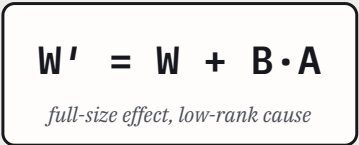
+

TWO THIN MATRICES · TRAINED OVERNIGHT

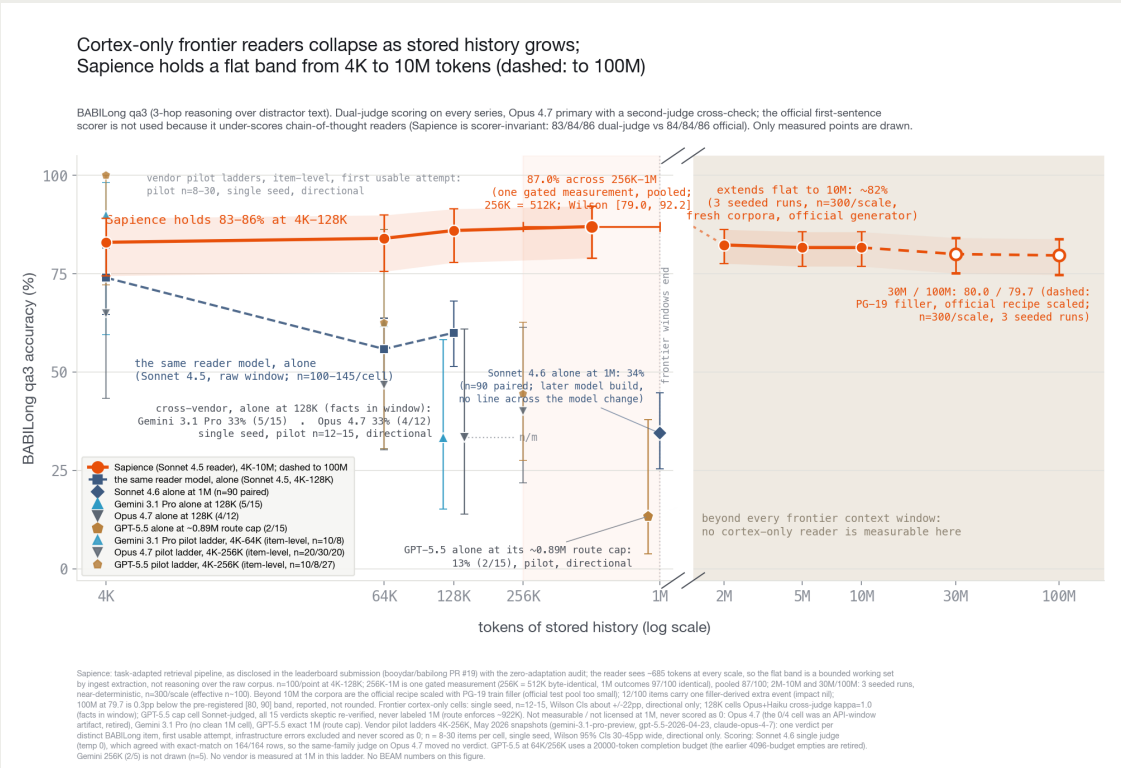


=

WHAT IT COMPUTES WITH



*W holds 16.8M numbers and is never trained.
B and A together hold about 0.5M numbers, roughly 3% of the layer,
and are the only thing trained each night. Delete them: the base is untouched.*



Accuracy and cost stay flat as the stored history grows from 2M to 100M tokens (81.0 / 80.0 / 80.0 / 80.0 / 79.7 at 2M / 5M / 10M / 30M / 100M; 3 seeded runs, n=300 per scale, official scorer), two orders of magnitude past any frontier window. The mechanism is not a bigger window: the reader sees about 685 tokens per question at every scale. Same figure as spnc.ai/evidence.

A SEPARATE 10M TASK • ANSWERING VS FRONTIER-NO-MEMORY • INTERNAL 10M CORPUS, HELD-IN CELLS, MATCHED PROTOCOL

0%

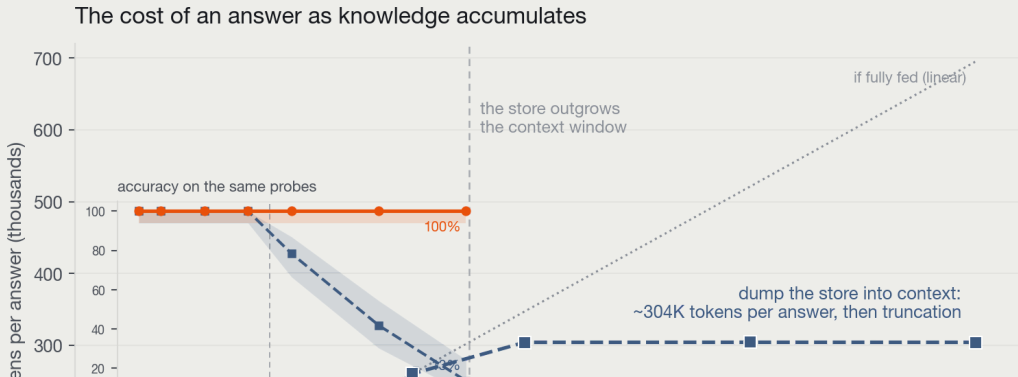
frontier at 10M (field exceeds its window)

45.0%

with Sapience at 10M (27 of 60 non-abstention, 3 seeds), vs 0 of 20 for the same model given its full 1M window (12.5% pooled, the correct answers being abstentions)

A different task family: our internal 10M-token corpus (populated gold, held-in cells), answering vs a frontier model with no memory. Sapience matches strong retrieval's accuracy from ~33x fewer evidence tokens per query (n=24).

AS THE STORE GROWS PAST THE WINDOW • ACCURACY AND COST



not from a swapped reader or an easier item set.

Multi-seed validation

No single-seed claim is a headline. Headline cells run three seeds (HALO_SEED 42, 123, 456) with stochasticity on. A positive on one seed is treated as provisional and labeled as such until it survives multiple seeds. A one-seed lift that dissolves under additional seeds is not reported.

Calibration holdouts excluded

Stratified calibration holdouts are reserved for threshold and variant selection and are excluded from every headline aggregate. We never report performance on the items that set the threshold.

Cross-judge audits

A second judge audits a subset of every scored run. On the cited cells, cross-judge agreement sits in the κ 0.97–1.0 range. Disagreement rows are re-adjudicated, not silently dropped.

Item-paired statistics

Paired claims use item-paired tests (McNemar) on the same items scored under both conditions, not a comparison of two independent rates. The 1M matched-reader pair is item-paired at n=90, same reader both arms; the headline BABILong figures are official-scorer leaderboard cells.

Every number is traceable

Every number in this deck traces to a result file and the

Cross-vendor at one million tokens.

Frontier readers alone on BABILong qa3 (3-hop), reading the raw window directly. Sapience holds where every frontier reader collapses. Provisional cells are labeled.

AT 1M TOKENS · BABILONG QA3 3-HOP			
READER	QA3 3-HOP	ACCURACY	N
Sapience (Sonnet 4.6 reader)		66.7%	90
Sonnet 4.6 alone, raw 1M window		34.4%	90
Item-paired at n=90, same reader both arms, Opus primary judge with four cross-vendor judges in agreement on every row; one deterministic run per block (seed 42), so item count is the robustness axis. Other frontier readers at 1M and below: the pilot cells on the A4 figure, directional only, never quoted as powered 1M values.			
AT 660K AND BEYOND · THE CONTEXT CEILING			
Fable 5 alone @ 660K	64%	50	dual-judge, CC harness
AT TEN MILLION TOKENS · INTERNAL 10M CORPUS, HELD-IN CELLS, MATCHED PROTOCOL			
Sapience + Sonnet @ 10M	45.0%	3 seeds	answerable, non-abstention
Sapience + Fable @ 10M	41.7%	24	early / provisional
Fable 5 alone @ 10M	0% substantive	4	n=4 provisional; 25% pooled incl. abstention-credit
Sonnet alone @ 10M	0% substantive	–	full 1M window, gold in-context: 0 of 20 on real questions; 12.5% pooled incl. abstention-credit

Beyond one million tokens the frontier readers cannot run; given its full 1M window with the gold in context, the same model scores 0 of 20 on real questions (12.5% with abstention credit), and 0% where the evidence lies beyond the readable window (window arithmetic). Pooled figures credit correct abstentions. Sapience is the only configuration here that still answers at ten million tokens.

CROSS-VENDOR IN THE FIELD · TWO PEOPLE, TWO AIS

Tom writes a note with his Claude: perovskite cell stability, recorded August 14, 2026, signed tom-de. His colleague asks a question in her ChatGPT and his note answers it, name attached, honest about its limits. His note, her answer, across vendors, with provenance intact. *(Demonstration brains; the mechanism is live.)*

variable.

Non-abstention n=60 (n=72 pooled), 3 seeds, non-overlapping Wilson intervals; the bare model on its native window, 2 seeds, 10.4% pooled and 0 of 40 non-abstention; the full-1M-window run is one deterministic seed. Approved runs 2026-06-19 (close-out) and 2026-06-21 (full window). Calibration holdouts excluded from every headline aggregate (see Methodology).

NON-SYNTHETIC, REAL-DATA BENCHMARKS

CONSOLIDATION UPLIFT · REPLICATED

+6.01_{pp}

Consolidation into weights

Held-out accuracy from training-time consolidation (LoRA, into weights); identical inference in all arms, so the lift is not test-time compute. 95% CI +3.16 to +8.87, 10 of 10 seeds, replicated on 5 fresh seeds. Protocol in Discovery by Dreaming (arXiv 2607.16256).

MULTI-HOP QA · REAL

43 to 28 vs 19 to 0%

MuSiQue, beyond the window

Multi-hop QA over natural Wikipedia text, exact match, 3 seeds: beyond the window the Sapience ladder holds 43 to 28 percent where a truncated full-context reader falls 19 to 0. Within the window, iterative bridge retrieval adds 32 to 37 points over full context with the

A real research program, and the IP under it.

ON THE PUBLIC RECORD · SEVEN PUBLIC, ONE IN REVIEW

Sapience: A Hybrid Architecture for Long-Context Reasoning and Continual Learning · [Zenodo](#)

Facts as First-Class Objects · [arXiv](#)

Selective Memory for AI (salience gating) · [arXiv](#)

Attention Is Not Retention · [arXiv](#)

Attribution-Native Machine Learning · [arXiv](#)

Constraint Gain: valuing negative results · [ResearchGate](#)

Discovery Stack · Nature Communications, in review

Discovery by Dreaming (cross-domain recombination) · [arXiv](#)

THE IP POSITION

Seven patent families, spanning structured attributed knowledge objects, matching and discovery, write-time extraction, consolidation and episodic replay, and human-AI collaboration. The continual-learning mechanisms are in filing now; priority receipts available in the data room.

Competing methods are published and open; ours are filed.

Why won't a frontier lab crush this? The question comes up in every serious meeting, and it has a direct answer.

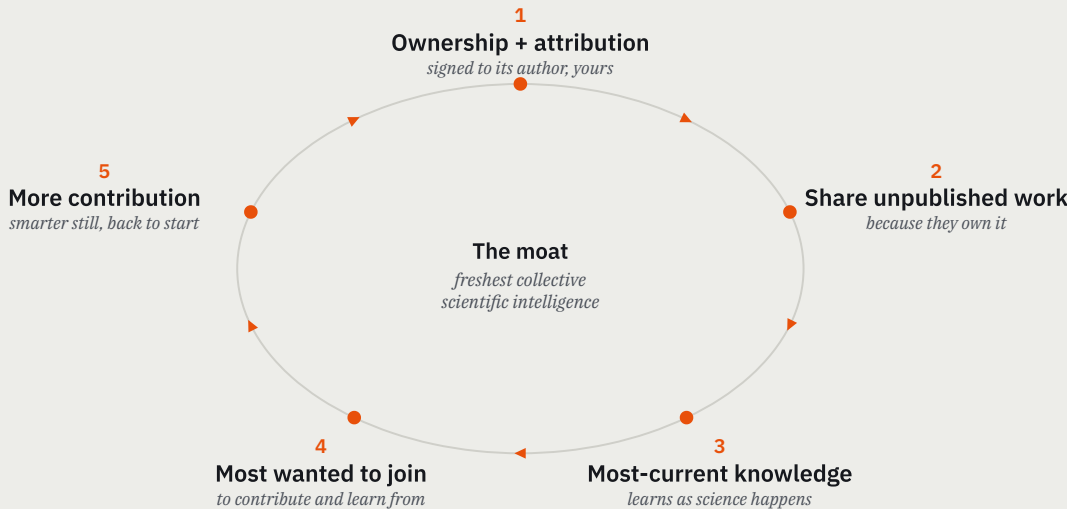
They are racing each other on generators, and the neutral learning layer across all of them cannot be owned by one of them. The free memory inside Claude, ChatGPT, Cursor, and Glean is the near-term version of this question, and the answer is structural: their memory features are single-vendor by construction, and deeper than that: training on users is the model business, and attribution would price it. And capital does not close the gap: compute cannot buy the attributed knowledge that accumulates here.

A model built to *absorb* your experts into its weights cannot also *compound* their expertise for them. A frontier lab's training pipelines require ingestion; ours do not. Their business puts your data on their servers; ours keeps it yours. They cannot follow without unwinding their own model.

THE MECHANISMS · HARD TO COPY Storing and recalling your own data is table stakes now, and better retrieval does not lead to this architecture: write order, supersession, and provenance are discarded at embedding time, and no amount of retrieval quality recovers what ingestion threw away. Memory and retrieval tools keep improving how they choose what to return; our advantage is in how we keep the knowledge in the first place. What compounds is connecting knowledge *across* organizations, shared as the mechanism rather than the raw result, with permission, credit to the contributor, and compensation. Attribution is also the training bet: provenance lets the system weight what it learns by how useful each claim has proven.

THE STRUCTURE · CLOSED TO A SINGLE-VENDOR LAB The reasoning core is swappable, so no vendor owns us. A single-vendor lab structurally cannot offer a you-own-it substrate that spans every AI and every organization, including the ones it competes with. The mechanisms above are engineering; this constraint is not: matching it would mean unwinding their own business.

Open standard, sealed engine. The CoreTx standard spreads and becomes the way knowledge moves between AIs; the matching and learning engine stays Sapience's.



what the discovery engine reasons over. The structure itself does measurable work all the way into the weights: replaying structured objects beats replaying the same content as raw prose by 10 to 12 points at a matched token budget (3 seeds, pre-registered).

05 IN THE FIELD: WHAT THE PRODUCT DOES OVERNIGHT

Overnight consolidation is live: it comes back in the morning with grounded connections between your own open problems, each one citing the objects it came from.

The researcher quotes in the appendix are unedited. They describe one brain compounding across providers, surfacing new ideas and connections, and cross-domain links a researcher would not have found alone.

06 WHERE THE TIME GOES: THE MECHANISM BEHIND THE 72 PERCENT

The oldest loop in science, ideas crossing fields, still runs on decades: Boolean algebra waited 83 years for circuit design. In research first, across academia and industry, where the price of frozen knowledge is highest, and then in knowledge work generally.

Structured, cross-domain matching covers that gap. On imports that actually happened, the true source sits in a top-40 slate a median of a decade before anyone reaches for it (a floor, not a cherry-pick).

Reasoning outside the box is where current AI is weakest: it completes patterns close to what it has already seen. Structured knowledge lets the engine trade off deliberately between nearby associations and far analogies, matching by mechanism across fields. The imports in the figure are that tradeoff at work.

Science is where this matters most, because nowhere else does knowledge move this fast and travel this badly. And the discovery architecture compounds: every brain that joins adds structured, typed knowledge the matching engine reasons over, so it gets better the more it connects. One case from the figure: given a 2025 zinc-battery problem blind, the engine surfaces a 1986 solvation-chemistry source in its slate, 39 years after it was published.

08 TIMING: DATA SUPPLY, COST STRUCTURE, AND WHAT IS OPEN

The industry’s data supply curve is running out: frontier labs now bid millions for defunct companies’ internal chat logs. The scarcest training signal is expert reasoning, and the store captures it at write time: structured, current, owned by the people who produced it.

Capital-efficient by design: we do not pre-train. The cost advantage is structural. An 8B open model on the store beats a model 30x its size reading the truncated raw context. A stronger reader still helps within its own window, and either way there is no per-token rent, so no vendor owns us.

What researchers said, **unedited.**

The first quote is a current Sapience user. The other three used an earlier version of the co-scientist in 2026; institution labels only.

“It’s amazing to have one brain’s memory infuse all my sessions: my Claude Code and web sessions are fully aware of each other now, and they talk to my Codex sessions too, across providers. That brain is constantly compounding, and it augments my own: it proactively suggests new ideas and connections.”

ALEX · PRINCETON UNIVERSITY · SAPIENCE

“I could swap out my current engine and use yours. The cross-domain connections are things I would never have found on my own.”

RESEARCHER · UNIVERSITY OF CAMBRIDGE · EARLIER VERSION

“The first two papers it surfaced were spot on for my research. It is giving me ideas on what to do next.”

RESEARCHER · CERN · EARLIER VERSION

“The more collaboration across the world, the faster science progresses. I one billion percent see the benefit.”

RESEARCHER · QUEEN MARY UNIVERSITY OF LONDON · EARLIER VERSION