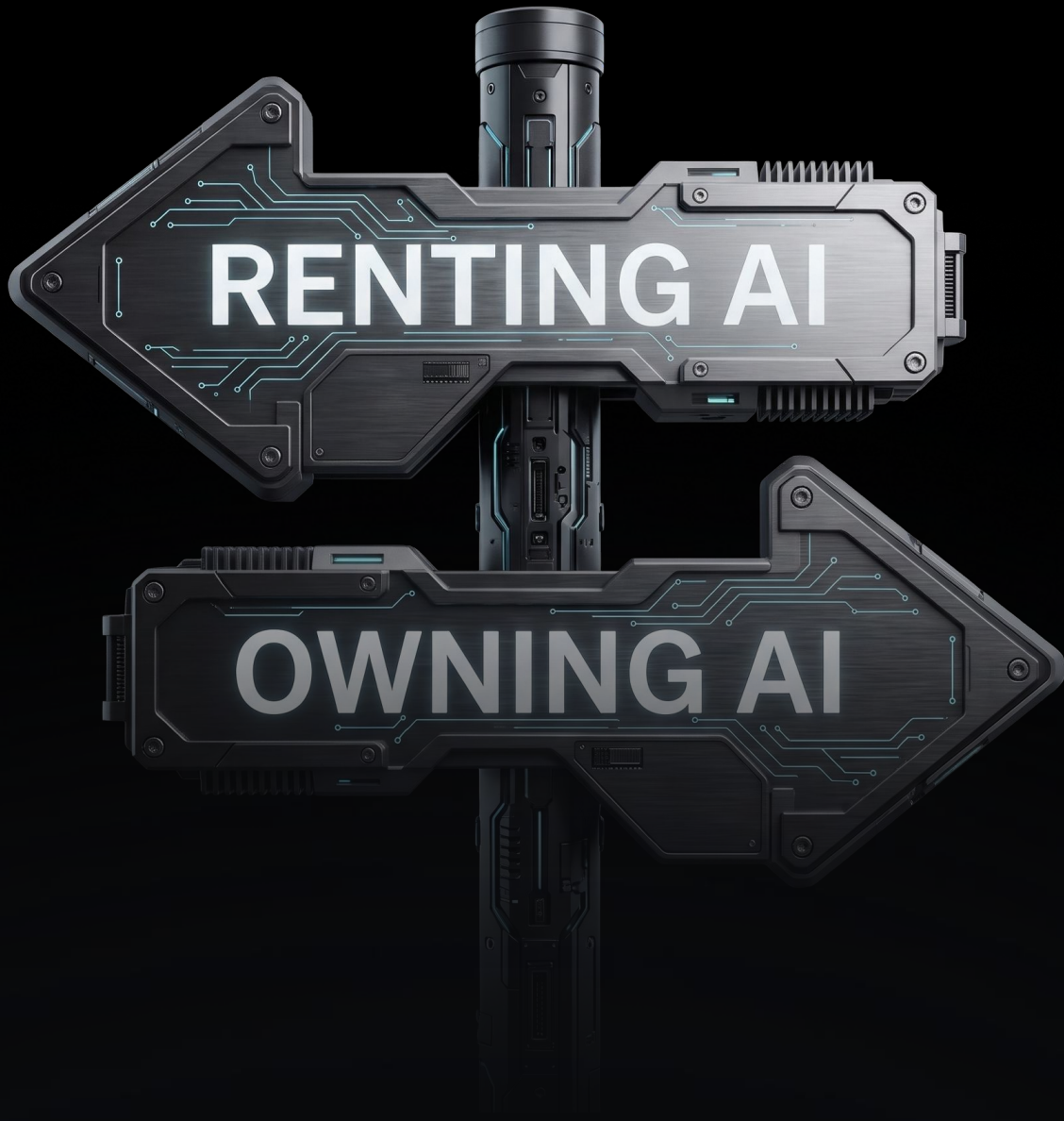




The AI Ownership Threshold

01 Introduction

“ How many developers does it take to justify owning your AI?

A two-year cost model built on GPT-5.6 Sol, Claude Opus 5, and the open-weight models that now rival them.

There is a line where renting frontier AI and owning it cost exactly the same, and as of July 2026 it sits at roughly 40 developers running agentic coding tools.

This paper calculates the relative costs of renting vs. ownership over a two-year time horizon. Forty developers at current agentic-tooling intensity consume about 8.4 billion tokens per month. Routed through the frontier flagships, OpenAI's GPT-5.6 Sol or Anthropic's Claude Opus 5, that volume bills roughly \$1.37 million over two years. Served instead on a self-hosted open-weight model, GLM-5.2 on two 8x H200 nodes, it costs roughly \$1.35 million over the same two years, all-in: hardware purchased outright, power, colocation, a dedicated two-person platform team, and credit for what the hardware resells for at month 24. At 40 developers, the two bills are the same dollar figure. That is the ownership threshold.

Everything else in enterprise AI economics is a position relative to that line. Below it, renting wins: a ten-developer team pays about \$327,000 over two years in API fees against \$1.13 million to own. Above the line, the advantage compounds: at 100 developers, owning wins by 1.9x to 2.0x; at 500 developers, roughly 105 billion tokens per month, renting bills \$17.2 to \$18.3 million over two years against \$5.4 million owned, a 3.2x to 3.4x advantage.

This paper derives the threshold, states every assumption behind it, and places the economics inside the larger framework that should govern the decision, including the two thresholds that have nothing to do with finances. The companies that make the best AI decisions over the next decade won't necessarily use the smartest models. They'll make the smartest ownership decisions.

02 The threshold, derived in humans

Nobody budgets in tokens. Organizations budget in seats and headcount, so the model starts there, with the two workload archetypes that dominate enterprise consumption. The conversion assumptions are stated so you can substitute your own telemetry.

Chatbot assumption: a typical enterprise chat exchange carries about 2,000 input tokens (system prompt, retrieved context, conversation history) and 500 output tokens, roughly 2,500 tokens per inference. An active user runs about 10 exchanges per working day, 21 working days a month: roughly 210 inferences, or half a million tokens, per user per month.

Developer assumption: a developer using agentic coding tools (Claude Code-class agents that read repositories, run commands, iterate on failures, and verify their own work) consumes 5 to 15 million tokens per working day depending on intensity, about 210 million per month at the 10M/day midpoint.

Those assumptions translate headcount into infrastructure. All figures are two-year cumulative costs:

Organization	Monthly tokens	2-yr flagship API bill	2-yr owned cost	Verdict
10 agentic developers	2B	~\$327,000	~\$1.13M	Rent + optimize
40 agentic developers	8.4B	~\$1.37M	~\$1.35M	The 1:1 threshold
100 agentic developers	21B	~\$3.4–3.7M	~\$1.80M	Own wins, 1.9–2.0x
500 agentic developers	105B	~\$17.2–18.3M	~\$5.4M	Own wins, 3.2–3.4x

Table 1: Two-year cumulative cost comparison of renting versus owning AI.

API figures use Claude Opus 5 and GPT-5.6 Sol at list; owned figures are fully loaded, net of hardware resale, and derived in the model section below.

The 1:1 line itself includes a two-person platform team (\$900,000 for two years) carried by the deployment: 1:1 lands at roughly 7 to 8 billion tokens per month, between 35 and 40 agentic developers depending on which flagship you'd otherwise rent (node quantization puts the exact crossing at 37 developers against Sol and 40 against Opus 5). This is the verified number for an organization standing up its first deployment, and it is the headline of this paper.



Agentic development crosses the threshold with a single team: 40 developers is not a hyperscaler, it is an engineering organization at a mid-size company. Boards budgeting AI by headcount are measuring the wrong axis. The axis that matters is autonomy: how many tokens your systems consume when no human is watching.

One variable deserves flagging even though the model deliberately excludes it: the threshold above assumes consumption holds flat. Per-developer token consumption has in fact been rising as agentic tools mature, roughly 1 million tokens per developer-day in 2024 to 10 to 20 million now, and any continuation of that trend moves organizations toward the line faster than headcount alone would. The model does not assume it continues; readers planning capacity should at least ask whether it will.

03 The model you'd actually host

A rent-versus-own comparison is only as credible as its own side. "Self-hosting" in 2026 does not mean settling for a small model and a quality haircut. Two releases in recent months moved open weights to within arm's reach of the closed flagships, and they anchor the two ownership configurations this paper prices.

- **GLM-5.2** (Z.ai, released June 16, 2026, MIT license) is the deployable-today assumption. It is a 744B-parameter Mixture-of-Experts model with about 40B active parameters per token and a 1M-token context window, landing near the frontier on single-shot coding benchmarks. The MIT license permits commercial use, fine-tuning, and air-gapped deployment without usage clauses. The hardware floor is set by memory, not compute: all 744B parameters must reside in GPU memory even though only 40B fire per token, so the FP8 weights (~750 GB) require a single 8x H200 node (1,128 GB aggregate) as the vendor's reference deployment. That node is the ownership unit priced below.
- **Kimi K3** (Moonshot AI, released July 16, 2026; weights published July 26) is a second priced configuration. At 2.8 trillion parameters with 16 of 896 experts active per token, it is the largest open-weight model ever shipped, and Moonshot published the full checkpoint on Hugging Face under permissive, commercially usable terms.

This paper prices two different pieces of hardware. K3 was trained with MXFP4 quantization-aware training rather than quantized after the fact, so the shipped checkpoint is the same precision Moonshot serves in production, not a lossy community reduction. But quantization-aware training reduces bits per parameter, not the parameter count: at roughly 4.5 bits per weight including scales, the published repository is 1.56 TB across 96 shards.



The hardware floor here is memory, not compute, and 1.56 TB does not fit the 8x H200 node that GLM-5.2 runs on. The ownership unit for K3 is therefore a single 8x B300 node. At 288 GB of HBM3e per GPU, 2,304 GB aggregate holds the checkpoint with roughly 740 GB left for KV cache at long context, keeps the whole model inside one NVLink domain, and executes MXFP4 in hardware rather than under emulation. Hopper can run K3 across two 8x H200 nodes at tensor-parallel 16, and vLLM's default recipe does exactly that, but this paper does not price it: you are buying two chassis to hold a model that fits in one, paying the cross-node interconnect penalty on every token, and emulating the numeric format the model was trained in. Anyone standing up K3 deliberately in the second half of 2026 buys Blackwell Ultra.

On the independent leaderboards Opus 5 sits second overall at 82.81, ahead of Sol (81.39) and K3 (79.89). On Terminal-Bench 2.1 the three are within a point of each other: Opus 5 at 89.1%, Sol at 88.8%, K3 at 88.3%.

The strategic point is unchanged and bigger than either model: the quality gap between closed flagships and open weights has compressed to weeks, and at the current cadence the model you can own trails the model you must rent by less than one release cycle. Ownership no longer requires accepting last year's intelligence.

What about renting the open models instead?

Both models are also available as APIs, at \$1.40/\$4.40 per million tokens for GLM-5.2 and \$3/\$15 for Kimi K3, far below flagship rates, and K3's cached input bills at \$0.30 against a reported 90% cache hit rate on coding workloads. Routing volume to them is excellent Phase 2 (Optimization) work, and it pushes the economic threshold outward. But renting an open model resolves neither the compliance threshold (the data still leaves, and for K3's hosted API it leaves to a PRC-based provider, which is a harder review than a domestic one) nor the strategic one (the endpoint can still vanish), and at fleet volumes the per-token bill still grows linearly while owned infrastructure does not. The open-model APIs are a waypoint, not the destination.

04 The cost model

Why two years

Capital decisions are not made on a single fiscal year, and a one-year frame structurally misprices ownership: the entire hardware purchase lands in one budget cycle while the asset serves for three or more, and the residual value of the hardware, which is real and liquid, never appears at all.



The model therefore compares two-year total cost of ownership against two-year cumulative API spend. Ownership is priced on a cash basis: hardware purchased outright at the start, two years of power, colocation, and platform engineering, minus what the hardware resells for at month 24. H100-class systems have held 75% to 85% of acquisition value through their first 24 months; the model assumes a more conservative 60% residual for H200 given Blackwell and Rubin ramp pressure, with 50% to 70% as the stated range.

Consumption is held flat across both years; the model deliberately assumes no growth, so every figure here is what the economics look like if your usage today is your usage in month 24.

The workload

The unit of enterprise AI work isn't a prompt; it's a corpus. Contract review, claims processing, knowledge-base construction, and agentic workflows all share the same shape: read large volumes of input, produce comparatively small volumes of structured output.

The model assumes a 10:1 input-to-output token ratio, reflecting document-processing and agentic pipelines. It runs at four sustained volumes: Pilot (50M tokens/month), Production (2B/month), Platform (20B/month), and Agent-scale (100B/month), which the opening table already translated into humans.

The rent side

The rent side prices the two frontier flagships an enterprise standardizing on maximum available quality would choose between as of July 2026. **GPT-5.6 Sol**, announced June 26 as a limited preview, is priced at \$5 per million input tokens and \$30 per million output, with a 1.05M-token context window; prompts above 272K input tokens bill at 2x input and 1.5x output, and the model assumes competent chunking that avoids the surcharge.

Claude Opus 5, released July 24 and now Anthropic's generally available flagship, is priced at \$5/\$25. At the 10:1 ratio they blend to \$7.27 and \$6.82 per million total tokens respectively, a 6% spread that quality evaluations on your workload should dominate. Enterprise-agreement premiums and committed-volume discounts run in opposite directions and are treated as a wash. Cheaper tiers exist on both sides (Terra, Sonnet); routing to them is Phase 2 work that roughly doubles the ownership threshold, and the sensitivity discussion covers it.



The own side

The primary ownership unit is the 8x H200 node that GLM-5.2's reference deployment requires. Per Q2 2026 OEM pricing, an integrated HGX H200 system runs \$320,000 to \$420,000; the model uses \$370,000 purchased outright, a 13kW system draw at \$0.12/kWh, \$25,000 per year in colocation, and resale at 60% of acquisition after 24 months. Net two-year node cost: about \$225,000, or \$188,000 to \$262,000 across the resale range.

K3's unit is a single 8x B300 node. Integrated 8-GPU B300 systems are anchored at \$300,000 to \$350,000 by public pricing trackers, though board-level GPU pricing near \$53,000 each implies OEM configurations well above that; the model uses \$400,000 to stay conservative and to absorb the direct-liquid-cooling facility work that B300 makes mandatory. It assumes an 18kW node draw at \$0.12/kWh, \$40,000 per year in colocation to reflect liquid-cooled rates, and the same 60% resale. Net two-year node cost: about \$278,000, or \$238,000 to \$318,000 across the resale range.

Engineering is a two-person platform team at \$450,000 per year fully loaded, \$900,000 over the horizon, shared across the fleet; this line item creates the model's most important structural property, covered below.

Throughput is the softest assumption in the paper, so both configurations state it explicitly.

For GLM-5.2 on 8x H200: a blended effective rate of 3,000 tokens per second per node at 70% sustained utilization, aggregate under continuous batching for a 40B-active MoE on a prefill-heavy workload. Single-stream decode is far lower; aggregate batched throughput is what determines cost. That gives each node roughly 66 billion tokens per year and a net node-only cost of about \$1.70 per million tokens over two years.

For K3 on 8x B300: 5,000 tokens per second per node at the same 70% utilization. K3 activates roughly 50B parameters per token against GLM's 40B and routes across 896 experts, which adds scheduling overhead, but B300 answers both with native MXFP4 execution and roughly 1.7x the memory bandwidth per GPU, on a model that fits in a single NVLink domain. That gives each node about 110 billion tokens per year and a net node-only cost of about \$1.26 per million tokens.

The results

Two-year cumulative cost, rent versus own, with specific models:

Agentic developers	Monthly tokens	2-yr volume	Own - GLM (8x H200)	Own - Kimi (8x B300)	GPT, 2 years	Claude Opus, 2 years
10	2B	48B	\$1.13M (1)	\$1.18M (1 node)	\$349K (0.31x)	\$327K (0.29x)
40	8.4B	202B	\$1.35M (2)	\$1.18M (1 node)	\$1.47M (1.09x)	\$1.37M (1.02x)
100	21B	504B	\$1.80M (4)	\$1.73M (3 nodes)	\$3.67M (2.03x)	\$3.44M (1.91x)
500	105B	2.52T	\$5.41M (20)	\$4.23M (12 nodes)	\$18.33M (3.39x)	\$17.18M (3.18x)

Table 2: Two-year cumulative cost comparison of renting versus owning AI across specific models.

Multiples in the API columns are against the GLM-5.2 configuration; against K3 on B300 they are 1.99x to 2.12x at 100 developers and 4.06x to 4.33x at 500. Multiples above 1.0x mean renting costs more than owning. A single B300 node carries K3 through the 40-developer row on throughput, and K3 cannot go below one node regardless of volume, which is what sets the 10-developer figure.

The table cuts against both camps. At ten developers, renting either flagship is roughly 3.4x cheaper than standing up either owned configuration. The picture inverts past the 1:1 line: at 100 agentic developers, ownership wins by 1.9x to 2.1x. At 500 developers it wins by 3.2x to 4.3x, between \$11.8 and \$14.1 million retained over two years.

The 1:1 line sits in the 40-developer row on the GLM-5.2 configuration, at 37 developers against Sol and 40 against Opus 5, and that remains the headline: it is the deployable-today, most-verified, most-conservative case. Hardware moves the line, though, and it moves it inward. A single 8x B300 node serves roughly 44 developers' worth of K3 throughput for \$1.18 million fully loaded over two years, which crosses 1:1 at about 32 developers against Sol and 34 against Opus 5. GLM-5.2 on the same hardware, which this paper does not price, would move it inward too. Every version of the calculation that uses current-generation silicon puts the threshold below 40, not above it.



The choice between the two owned configurations is now more a hardware question than a model question. GLM-5.2 runs on Hopper you may already own, needs no liquid cooling, and is near-frontier on coding. K3 needs Blackwell Ultra and the facility work that comes with it, and in exchange gives you the strongest open-weight model in existence, within a point of both closed flagships on agentic coding, on fewer nodes than GLM requires at every volume above the threshold. If you are buying hardware for this workload in the second half of 2026, that is the decision. At fleet scale the rational architecture runs both: GLM for volume, K3 for the hard tasks, and a rented flagship for the exceptions.

The practical reading of the threshold: below 10 developers, rent without a second thought. Between 10 and 40, optimize aggressively and model your own break-even with your own telemetry, because you are approaching the line and current-generation hardware may already have moved it behind you. Past 100, ownership is winning by roughly 2x fully loaded.

05 The framework: Rent → Optimize → Own

This is not simply buy versus build. That binary is too crude to be useful. Every AI capability moves through three phases, and many capabilities will never reach the third. The mistake isn't renting. The mistake is owning too early, or renting for too long.

- **Phase 1:** Rent. Nearly every capability should begin here: low utilization, rapid experimentation, unstable workflows, fast-moving model generations. The goal is not cost efficiency. The goal is learning what token consumption actually looks like once real users touch the system. At pilot volumes the model below shows renting winning by two orders of magnitude; any infrastructure spend at this stage is a capital allocation error.
- **Phase 2:** Optimize. Usage rises, workflows stabilize, consumption becomes forecastable. Now the optimization toolkit matters: routing requests between model tiers by complexity, prompt caching, batch processing, distillation, retrieval optimization. Both flagships reward this work directly: cached input reads bill at roughly a tenth of the standard rate on both providers (Opus 5 cache reads run \$0.50 against \$5 standard input), and batch processing takes a flat 50% off asynchronous jobs. A cached and batched workload can run at roughly a quarter of list price. Most organizations should spend a surprisingly long time in this phase, and every dollar of optimization pushes the ownership threshold further out.
- **Phase 3:** Own. Ownership begins when AI stops being software and becomes infrastructure, when the question shifts from which model should we call to which capabilities should we permanently operate ourselves. It should happen only when at least one of three thresholds has been crossed.



06 The three thresholds

Ownership is rarely triggered by a single variable. Three distinct thresholds exist, crossed for different reasons, at different times, by different kinds of organizations. Only the first is about money.

Threshold 1: Economic

At some sustained volume, renting becomes more expensive than operating. This is the threshold quantified above and derived below, because it is the only one of the three that can be calculated rather than argued: over a two-year horizon it sits at roughly 40 agentic developers fully loaded.

Threshold 2: Compliance and data privacy

For a large class of organizations, the ownership decision is made long before the economics resolve it, because the data cannot leave. Healthcare organizations operating under HIPAA, financial institutions under data-residency and record-keeping mandates, firms bound by attorney-client privilege, and any enterprise processing regulated personal data all face the same structural problem: an API call is a data transfer. Every prompt that crosses the vendor boundary triggers the machinery of data-egress review, privacy assessment, vendor risk management, and business associate or data processing agreements, and some categories of data cannot cross at all.

The API route prices some of this in: providers now charge measurable premiums for regional data residency, and Anthropic's US-only inference option for Opus 5 bills at 1.1x input and output — a 10% uplift, stated on the price sheet. But a surcharge doesn't dissolve the underlying exposure. Prompts, retrieved documents, and agent trajectories are the most sensitive data an enterprise produces, a live feed of what the organization knows, decides, and worries about, and workflow leakage through AI tooling is the silent risk most governance programs haven't caught up to. A model running inside your own boundary, air-gapped if necessary, retires the entire category: nothing leaves, so nothing needs review. For regulated enterprises, this threshold typically triggers first, and the economics arrive later as confirmation rather than cause.

The K3 weight release sharpens this into a specific, common case. Organizations that want K3's capability but cannot route prompts to a Chinese-operated endpoint — a live constraint in finance, healthcare, defense, and legal — previously had no option. As of July 26 they have one, and it runs entirely inside their own jurisdiction on hardware they control.



Threshold 3: Strategic

The third threshold is control: guaranteed availability, decision auditability, the ability to modify the system without vendor permission, and continuity if a vendor changes pricing, policy, or existence. The past month alone made this concrete. GPT-5.6 launched as a limited preview restricted to government-vetted partners, and Anthropic's newest flagship spent most of June offline under a US export-control order before being restored on July 1. Frontier capability is now subject to policy decisions that no enterprise controls, on timelines no procurement cycle can absorb. Weights you hold cannot be paused, repriced, deprecated, or export-controlled out from under a production workflow.

That protection is real but it is not unconditional, and the K3 release illustrates the edge. Weights already downloaded under a permissive license cannot be recalled. But US policy toward Chinese open-weight models is unsettled — Commerce has reportedly considered Entity List additions and hosting restrictions, and the White House OSTP publicly accused Moonshot of training K3 on export-controlled silicon and distilling US models. An enterprise that has the weights on its own disks is insulated from all of that in a way that an enterprise calling a hosted endpoint is not. For organizations where a single vendor or policy change is an existential dependency, ownership is insurance, and insurance is allowed to cost money.

The rest of this paper models Threshold 1, because it's the one enterprises can put in a spreadsheet. But note the order of operations: if you've crossed Threshold 2 or 3, the economics are context, not the decision.

07 Why the gap widens: linear versus sublinear

The structural argument matters more than any current calculations, because list prices will change and the multiples will move.

Per-token pricing makes cumulative cost a strictly linear function of volume and of time. That is what per-token pricing means: year two costs exactly what year one did, forever. Ownership is front-loaded and then flattens: the capex lands once, resale value comes back at the end, and in this model per-node two-year cost including the engineering share falls from \$1.13 million at one node to \$450,000 at four to \$270,000 at twenty, because the platform team amortizes across the fleet while volume scales 50x. Two curves with those shapes always diverge. The only question is where they cross, and everything after the crossing compounds in ownership's favor. Extend the horizon to three years and the multiples grow again; the two-year frame is the conservative one.



There's a second-order effect that matters more than the bill. When the marginal token costs money, teams ration. They cap agent iterations, sample instead of processing the full corpus, summarize instead of reading everything. Every rationing decision degrades output quality. When the marginal token is free, the rational behavior flips: process everything, rerun whenever the pipeline improves, let agents iterate until the task is done rather than until the budget is. Fixed-cost infrastructure doesn't just lower the bill. It changes what the organization is willing to attempt.

And there's a budgeting argument CFOs appreciate more than engineers do. A per-server cost is a number you can put in next year's budget, twice. A per-token cost is a forecast, and token forecasts have a habit of being wrong by multiples once a project succeeds and every adjacent team wants in. Predictability has independent value, and over a two-year horizon it has a lot of it.

08 A practical decision framework

Before moving any capability from Phase 2 to Phase 3, ask five questions:

- 1. Is utilization predictable?** Ownership economics collapse below sustained utilization. If the workload is spiky or experimental, stay in Phase 2.
- 2. Can the data leave?** If prompts, documents, or agent trajectories carry regulated or privileged content, the compliance threshold may already have decided for you.
- 3. Does control create strategic value?** The past month's availability disruptions at both frontier labs are the argument in miniature: weights you hold cannot be paused or export-controlled.
- 4. How many agentic developers do you have?** Forty is the two-year 1:1 line; seven per node once a fleet exists; five hundred is decisive. Count honestly, and count where you'll be in twelve months, not where you are.
- 5. Would switching vendors materially disrupt the business?** Deep vendor dependence on a critical workflow is a risk that ownership retires.

Zero or one yes: rent, and revisit quarterly. Two or three: optimize aggressively and start modeling the break-even with your own numbers. Four or five: you are likely past the threshold already, and every month of delay is a transfer of margin to your vendors.



09 What this model does not claim

The model is about cost, and cost is the second question. Quality is the first: an owned deployment that can't clear the accuracy bar for its workload has no value at any price. The open-weight releases narrow this concern substantially for coding and agentic work, where K3 sits within a point of both closed flagships on Terminal-Bench and GLM-5.2 lands near the frontier, but benchmark parity is not workload parity, and the bar must be established by evaluation on your tasks, not assumed. Opus 5's launch is a reminder that the ordering is unstable: K3 led Anthropic's GA flagship for eight days and then didn't. For the highest-difficulty reasoning tasks the closed flagships currently remain the right answer regardless of volume, and the rational fleet architecture routes exceptions to them.

The model holds consumption flat across the two-year horizon. That is a deliberately conservative choice, not a forecast: per-developer consumption has been rising as agentic tools mature, and if that continues, the effective threshold arrives at lower headcounts than the static figures state. The model excludes engineering time from the per-node break-even on the grounds that both routes need pipeline work, and shows the fully-loaded figures alongside so you can disagree. It excludes the compliance costs that the API route triggers and local deployment largely avoids; counting those would move the threshold toward ownership. The model assumes a 10:1 input-output ratio; chat-heavy workloads shift the blend toward expensive output tokens, while cache-friendly workloads favor renting. The human-scale conversions are mid-range figures with stated ranges; your per-exchange and per-developer consumption will differ.

K3's self-hosting economics are no longer directional on availability, since the weights shipped, but they remain directional on throughput. The 1.56 TB checkpoint size is a published fact and the memory floor follows arithmetically from it, which is why K3 is priced on B300 rather than on the H200 node GLM-5.2 uses. The throughput figure does not follow from anything published. Five thousand tokens per second per node is an estimate for a 50B-active, 896-expert MoE running native MXFP4 on Blackwell Ultra; engine support for KDA and Stable LatentMoE was landing in vLLM and SGLang the same week as the weights, and there is not yet a body of independent production benchmarks to check it against. Treat the K3 column as a first estimate with wider error bars than the GLM column. The model also charges no cost of capital on the upfront purchase; at a 10% hurdle rate, add roughly \$40,000 per node per year on either platform, which moves the fully loaded 1:1 line by a handful of developers.

And it freezes prices at a point in time. API prices have been falling 30% to 50% per year; GPU prices fall too. Sol is in limited preview and its general-availability terms could shift. Anyone using this model a year from now should rerun it with current numbers. The assumptions are stated precisely so that you can.



One thing price cuts do not change: the geometry. A 50% API price cut at agent-scale turns \$17.2 million into \$8.6 million against \$5.4 million owned over two years, and the linear-versus-sublinear divergence resumes immediately. Only a structural change in API pricing, genuinely unmetered flat-rate enterprise tiers at these volumes, would break the argument. Nobody offers that today, and frontier model shops are actively moving away in the opposite direction.

10 What it means for planning an AI budget

The practical guidance falls out of the threshold. If you are running one pilot workload or a ten-developer agentic team, the highest-ROI work is Phase 2 optimization, not procurement. If you have 30 to 40 developers on agentic tooling, you are at the 1:1 line right now, and the decision should be made deliberately rather than discovered in an invoice. This paper also assumes agentic engineering as the predominant task - the equation will shift if you are focusing on other workflows. If you are regulated, the compliance threshold likely fires before the economic one, and the arrival of permissively licensed frontier-class weights — now including the largest open model ever released — means crossing it no longer costs a quality haircut. And if you have over 100 developers with agents, the math resolved some time ago: over the next two years, renting will cost you multiple times what the capability costs to own.

Most organizations will continue renting foundation models, and they should. Very few should ever train one, and after this month, none need to. But the winners of the next decade won't simply consume intelligence at metered rates. They will own the systems, the context, the evaluation infrastructure, and at sufficient scale the weights and the compute, that make intelligence uniquely valuable to their business.

The model becomes infrastructure. Ownership becomes strategy. And the point at which renting should give way to ownership is no longer an abstraction: it is a threshold you can calculate in developers, users, or tokens, over the same horizon used for capital planning.



11 FAQ

Why is this only coming to light now?

Prior waves of AI adoption centered around chatbot deployments. Simple chats required far fewer tokens (500-3,000) while agentic engineering tasks require far more (30,000-1,000,000). Agentic solutions first reached commercial viability in early 2026. Finally, open-weights models similarly reached frontier-level intelligence in mid-2026.

Which open-weight model does the model assume?

GLM-5.2 is the primary priced configuration: 744B-parameter MoE (~40B active), MIT license, 1M context, near-frontier coding quality, reference deployment of a single 8x H200 node holding ~750 GB of FP8 weights, about \$1.70 per million tokens node-only.

Kimi K3 is priced as a second configuration since its weights shipped July 26: 2.8T parameters, 16 of 896 experts active, and a 1.56 TB MXFP4 checkpoint. That is a hard memory floor, and it does not fit one 8x H200 node. The paper prices K3 on a single 8x B300 node (2,304 GB), where the checkpoint fits with roughly 740 GB of KV cache headroom and MXFP4 executes natively, at roughly \$1.26 per million tokens node-only. That figure has the widest error bars in the model.

Doesn't the quality gap between open and closed models break the comparison?

Less every month, but the gap is real and it moves in both directions. K3 led Claude Opus 4.8 at launch; Opus 5 shipped eight days later and now leads K3 (82.81 vs 79.89 on the composite leaderboard), as does Sol. On Terminal-Bench 2.1 all three sit inside a point: Opus 5 at 89.1%, Sol at 88.8%, K3 at 88.3%. GLM-5.2 lands near the frontier on coding. For the hardest reasoning tasks the closed flagships keep an edge, and the rational architecture routes those exceptions to a rented API while owned capacity serves the volume. Establish the bar by evaluation on your workload, not by benchmark tables, including this paper's.

Why a two-year horizon instead of one?

Because a one-year frame structurally misprices ownership: the full hardware purchase lands in a single budget cycle, and the hardware's resale value never appears. Over two years the capex amortizes against a cumulative rent bill and the residual comes back. Three years would widen it further; two is a conservative choice.



What would change the conclusion?

A structural change in API pricing, such as truly unmetered flat-rate enterprise tiers at these volumes. Price cuts alone don't: they move the crossover outward but leave the linear-versus-sublinear geometry, and therefore the eventual inversion, intact. On the other side, a decline in per-developer token consumption would move the threshold out; the model already takes the conservative position by assuming consumption merely holds flat.

Pricing sources: GPT-5.6 Sol rates from OpenAI's June 26, 2026 announcement and API documentation; Claude Opus 5 rates, release date (July 24, 2026), and US-only inference uplift from Anthropic's pricing documentation and system card; GLM-5.2 specifications and reference deployment (Z.ai/Hugging Face, June 2026); Kimi K3 specifications, weight release (moonshotai/Kimi-K3, July 26, 2026), checkpoint size, and deployment guidance from Moonshot's model card and independent deployment write-ups; leaderboard positions from BenchLM and Artificial Analysis as of July 27, 2026; Kimi K3 checkpoint size from the moonshotai/Kimi-K3 repository file manifest (1,560,936,091,448 bytes across 96 safetensor shards, July 2026); deployment topologies from SGLang's Kimi-K3 cookbook and vLLM's July 27, 2026 day-0 support post; B300 system pricing, 288 GB HBM3e per-GPU capacity, and liquid-cooling requirement from public 2026 pricing trackers and OEM quotes; H200 system pricing and H100-class residual-value data verified against public 2026 pricing trackers and OEM quotes (BenchLM, OpenRouter, Mercatus GPU Index, vLLM deployment guides), July 2026. All ownership figures are US-market estimates; recompute with your own power, colocation, throughput, resale, and per-developer consumption assumptions before making capital decisions.

About Datasaur

Datasaur is a private LLM provider and data labeling platform designed for companies to build their AI ecosystem with ease and efficiency. It assists organizations and universities in setting up custom LLMs and annotating data more efficiently and accurately through automation, quality control, and human-in-the-loop workflows. For more information, visit www.datasaur.ai.

Schedule a demo