

The AI Infrastructure & Semiconductor Investment Atlas

Global Systems, Supply Chains, and Security Analysis

Second edition — reader orientation

Version: 2.0

Evidence cutoff: 25 July 2026

Coverage: 544 public, private, subsidiary, and state-owned participants across the global AI stack

Companion outputs: analytical reader edition · company and access handbook · archival full edition

This atlas is a map of the global AI economy and an investment decision framework. It distinguishes three questions that are often blurred together:

1. **Is a company strategically important to the AI stack?**
2. **Does that importance materially affect the security an investor can own?**
3. **Does the current price offer an attractive expected return?**

The first question determines coverage. The second determines investability. The third determines action. A monopoly, national champion, or private frontier lab can be essential to the map without being an attractive or directly investable security.

Three ways to use the work

- **Analytical reader edition:** Parts I–IV, the AI-factory capstone, the investment decision layer, and the methodology. Read this for the argument, economics, scenarios, and actionable shortlist.
- **Company and access handbook:** all 544 approved company records, organized by primary layer and accompanied by evidence, access treatment, monitoring triggers, and source links.

- **Archival full edition:** the reader edition and the complete handbook in one deterministic manuscript.

Contents

1. Part I — The Argument
2. Part II — The Stack, Layer by Layer
3. Chapter 2.13 — From GPU to AI Factory
4. Part III — Systems in the Round
5. Part IV — Where It Is Going
6. Part V — Investment Decision Layer
7. Company and Access Handbook
8. Part VI — Reference and Apparatus

Ten conclusions

1. **AI demand is real, but infrastructure revenue still outruns disclosed end-customer economics.** The durability question is not whether models improve; it is whether production workloads produce enough cash to service the capital built for them.
2. **The bottleneck migrates.** Accelerator scarcity moved into HBM, packaging, networking, power, and construction. A durable investment process follows lead times and capacity response rather than yesterday's shortage.
3. **Strategic importance does not equal equity capture.** Parent dilution, regulated pricing, customer bargaining power, state ownership, and valuation determine how much supply-chain leverage reaches shareholders.
4. **Layer diversification can conceal one common factor.** Much of the Western stack depends on a small group of hyperscaler capital budgets; much of the Chinese stack depends on state procurement and domestic fabrication capacity.
5. **Inference expands the addressable market and intensifies the depreciation debate.** Falling unit costs support usage, while rapid hardware turnover can weaken reported returns even when demand remains strong.
6. **The allied chokepoint system is as important as either superpower.** Taiwan, Japan, Korea, and the Netherlands control fabrication, memory, lithography,

substrates, and engineered materials that neither the United States nor China can quickly replace.

7. **The third bloc is a demand system, not a neutral spectator.** Gulf sovereign capital, Indian and Southeast Asian capacity, European regulation, and non-aligned procurement influence which technical stack reaches global scale.
8. **Power is regional.** National electricity forecasts do not determine whether a specific campus can secure land, interconnection, equipment, water, permits, labor, and an economically acceptable tariff.
9. **Private-company marks and public-market prices are different objects.** The handbook preserves private strategic coverage; the actionable shortlist contains liquid securities and scenario-based valuation discipline.
10. **The atlas is a dated decision system, not a timeless ranking.** Prices, sanctions, private valuations, model leadership, and quarterly spending require scheduled refreshes and explicit stale-after dates.

Confidence language

The manuscript uses the following reader-facing states:

- **Observed:** reported historical or current data.
- **Company-reported:** an issuer or private company metric that is not independently audited for the purpose used.
- **Estimate:** a transparent third-party or modeled value.
- **Forecast:** a forward projection attributed to its source.
- **Analytical synthesis:** a conclusion or scenario constructed from multiple inputs.
- **Contested:** credible sources materially disagree or the available evidence is weak.

Market prices in the investment decision layer are observations at the latest available close before the cutoff. Return ranges and scenario weights are analytical synthesis, not price targets or individualized advice.

Part I — The Argument

The whole thesis in plain language. Most readers need only this part; the rest of the atlas is the evidence and the detail behind it.

1.0 How to read this atlas

This report is built to be read at whatever depth you need. **Part I**, which you are reading, is the argument in about twenty minutes. **Part II** is the core: twelve stack layers plus an AI-factory bill-of-materials capstone, each a self-contained briefing on role, economics, participants, monitoring indicators and risks. **Part III** compares the American, Chinese, allied-chokepoint and third-bloc systems. **Part IV** is the forward view and monitoring dashboard. **Part V** separates a 30-security investment decision layer from the 544-company evidence handbook. **Part VI** is reference: glossary, methodology, and sources.

A few conventions. The report's fixed evidence cutoff is 25 July 2026, and every quantitative figure also carries its own observation or forecast period; the fastest-moving numbers, private valuations and quarterly capex above all, will have drifted by the time you read this, so verify before acting. Numbers that are contested, estimated, or drawn from weaker sources are flagged with a warning mark and a note on the disagreement, because false precision is more dangerous in this field than honest uncertainty. Company names are introduced with their ticker, and the full tagging lives in Part V. This is research and structural analysis. It is not investment advice.

1.1 The one-page thesis

The artificial-intelligence build-out is the largest and fastest capital-formation event in the history of technology, and understanding it as an investor comes down to five shifts, each of which cuts against the obvious trade.

The first is that **the bottleneck has broadened beyond the chip**. Three years ago the scarce thing was the GPU. Today electricity, memory, advanced packaging, networking and grid interconnection can all constrain deployment, while more system value is migrating from the logic die into memory and packaging. Bottleneck owners may capture the build more efficiently than compute vendors, but only when durability and valuation support the strategic thesis.

The second is that **the frontier of intelligence moved from training bigger models to running them longer**. The old approach of simply enlarging models has hit diminishing returns; the new approach makes models reason at the moment of use, which consumes far more compute per query and does so on every use rather than once. That keeps aggregate compute demand climbing even as the models themselves commoditize.

The third follows from it: **cheaper intelligence expands spending rather than shrinking it**, at least for now. Efficiency improves several-fold a year but usage grows faster, so total compute and capital spending keep rising. The deflation is real but it lands on the model layer and on per-unit pricing, not on the suppliers of compute and power.

The fourth is that **the money has turned reflexive and increasingly leveraged**. Capital spending now exceeds the cash the spenders generate, so the build is funded by debt and by a web of circular deals in which the industry finances its own demand, a structure that both central banks have named a stability concern.

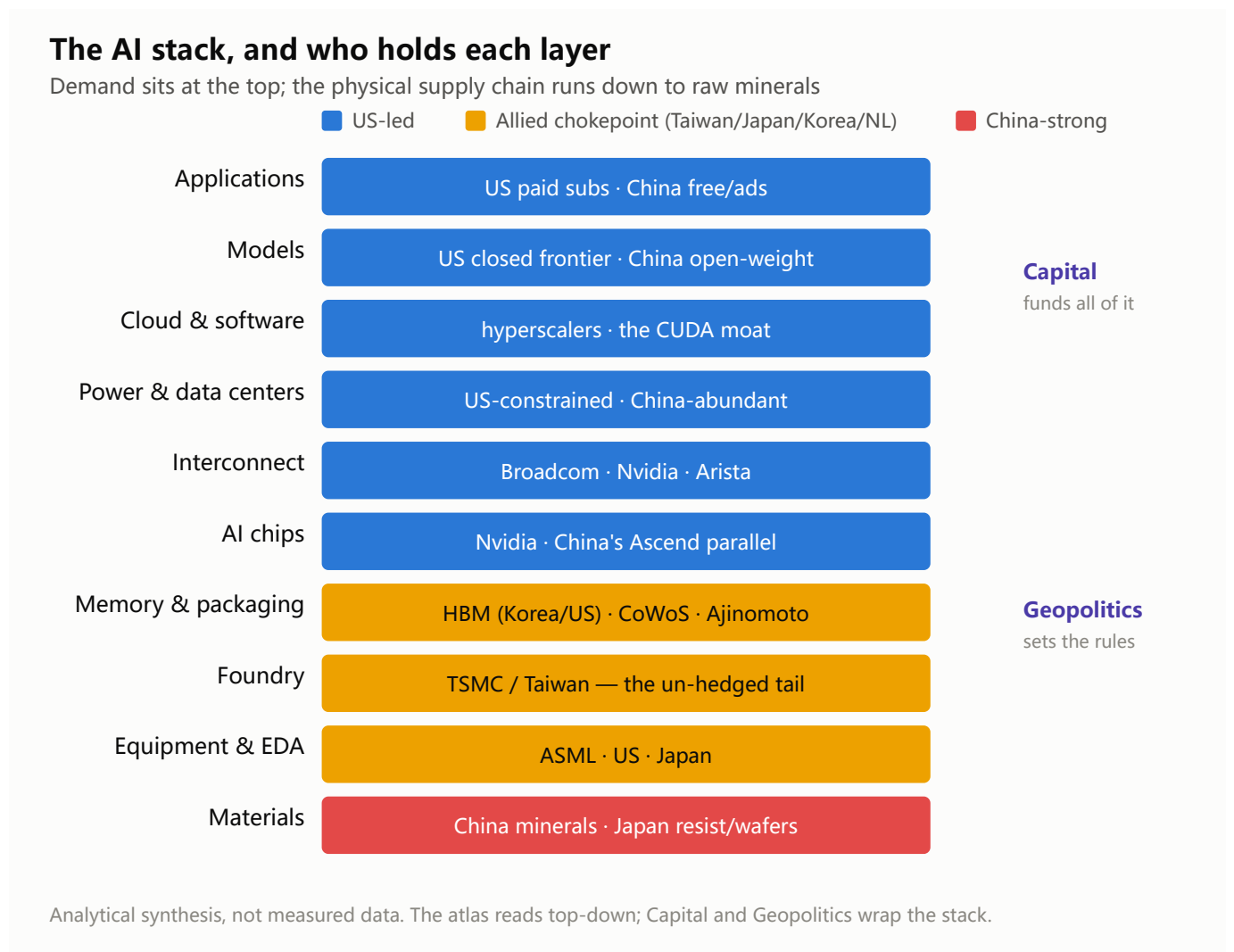
The fifth is that **the world is splitting into two technology stacks under a policy regime that has become transactional**, with Taiwan as the one risk that can be sized but not hedged. An American stack and a Chinese stack are diverging at every layer, the contest has moved to the uncommitted "third bloc," and beneath all of it the leading-edge manufacturing that the entire build depends on sits on one island.

Underneath these five sits the single question that decides whether any of it is durable: whether enterprises earn a real return on AI. The evidence so far is that they do where the workflow is deep, coding above all, and do not yet where it is shallow. The market

has noticed, and has begun to reward monetization over ambition.[^atlas-global-menlo-enterprise-ai-2025]

1.2 The stack, explained

To follow the argument you need a mental model of how AI is actually built, because every later chapter refers back to it. The industry is a stack of layers, each depending on the one below.



At the top are the **applications** people pay for and the **models** that power them, the demand that justifies everything beneath. Models run in the **cloud**, on clusters that require enormous **power** and the data centers to house it. The clusters are built from **AI**

chips, bound together by **interconnect** and fed by **memory**, all of it assembled through advanced **packaging**. The chips are manufactured by a **foundry**, using **equipment** and design **software**, out of raw and refined **materials**. Wrapping the whole structure are two forces that touch every layer: the **capital** that funds it and the **geopolitics** that sets its rules.

This atlas reads the stack demand-first, starting at the top with what drives spending before descending to what that spending buys. The diagram's central message is *where the advantage sits*: the United States leads much of the application, cloud and compute stack; allied chokepoints in Taiwan, Japan, Korea, and the Netherlands hold difficult-to-replace manufacturing inputs; and China is strong in selected raw materials, power deployment and a rapidly localizing parallel stack.

1.3 The five meta-theses

The five shifts above translate into five investable theses that recur throughout the atlas.

Own the bottleneck, not the compute. As transistor shrinks deliver less and demand outruns supply, value pools into the constraints: high-bandwidth memory, now more than half a GPU's cost and sold out; advanced packaging, the gating step on how many chips can be built; the electrical equipment and power that a data center cannot run without; and the lithography tools without which no advanced chip exists. These are the picks-and-shovels of the build, and they carry order-backed earnings and real pricing power.

The frontier is compute-hungry, so total demand keeps rising even as models commoditize. Reasoning, agents, and video all consume more compute, not less, which is bullish for silicon, power, and memory in aggregate. The offsetting caution is that the model layer itself is deflating, so owning the labs is a bet on distribution and product, not on model supremacy.

Jevons beats deflation in aggregate, but deflation bites the crowded trades.

Efficiency gains are met by faster usage growth, so capital spending climbs; the deflation shows up in frontier-lab margins and in the periodic efficiency shocks that re-rate the compute names on fear before demand backfills.

The money is reflexive and levered, so the risk is now financial. Circular financing, depreciation-flattered earnings, passive-flow concentration, and GPU-backed debt link the players together and shift the danger from "will the technology work" to "will the financing hold." Quality-of-earnings analysis has become a source of edge.

Two stacks, a transactional policy regime, and an un-hedged Taiwan tail. Bifurcation is the base case, the contest is for the uncommitted middle, and the concentration of manufacturing in Taiwan is the one exposure that can only be respected, not diversified away before roughly 2030.

1.4 The two framing lenses

Two ideas help make sense of the whole map.

The first is the **neutral chokepoints**. The hardest, least-substitutable points in the entire supply chain are held not by either superpower but by third parties both depend on: extreme-ultraviolet lithography by ASML in the Netherlands, leading-edge fabrication and advanced packaging by TSMC in Taiwan, high-bandwidth memory by the Koreans, and key engineered materials by Japan. These chokepoints are where the real leverage sits, and they are why a US-China framing alone is incomplete: some of the most important companies in AI are neither American nor Chinese.

The second is **bifurcation**. One interdependent global supply chain is splitting into two, an American-led stack built on Nvidia and CUDA and a Chinese-led stack built on Huawei's Ascend and CANN, each with its own hardware, software, and standards. This split now defines the geopolitics of the industry, and the decisive question has become which stack the uncommitted countries of the world will adopt. These two lenses run through every chapter that follows.

1.5 From strategic map to security decision

Strategic importance is not an investment rating. A monopoly supplier can be a poor security at an extreme price; a diversified parent can own an essential asset that is immaterial to consolidated earnings; and several apparently different stocks can be one correlated hyperscaler-capex trade. The **Investment Decision Layer** therefore narrows the atlas to 30 liquid US-listed securities and tests four separate axes: strategic criticality, equity capture, valuation and expected return. Each card records current price, exposure, catalysts, thesis-break conditions, factor overlap and analytical two-year bear/base/bull ranges. Those ranges are scenario disciplines, not price targets.

The resulting shortlist spans compute, cloud, interconnect, memory and packaging, equipment and materials, and power and electrical infrastructure. It deliberately includes securities with unattractive base cases: seeing that a strategically important company is already priced for success is part of the conclusion. The companion **Company Handbook** preserves the global 544-company evidence universe, including private firms and restricted securities, without pretending all are actionable ideas.

The one screen that matters most differs by side of the Pacific. For the American names, it is concentration and circular financing: before buying any AI-infrastructure name, ask whose money is really behind its revenue, external demand or the industry funding itself. For the Chinese names, it is access and reality: whether a foreign investor can even own the security, and whether the "domestic" champion is truly self-sufficient or still dependent on the foreign tools and memory it claims to have replaced. Hold those two screens, and the rest of this atlas is detail.

This part synthesizes the thirteen chapters of Part II and the forward view of Part IV. The evidence cutoff is 25 July 2026; quantitative observations and forecasts retain their own periods.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-global-[iea-energy-ai-2025](#)] [[^]atlas-global-menlo-enterprise-ai-2025]

[^atlas-global-iea-energy-ai-2025]: Energy and AI (<https://www.iea.org/reports/energy-and-ai>). International Energy Agency, 2025-04-10; accessed 2026-07-25.

[^atlas-global-menlo-enterprise-ai-2025]: 2025: The State of Generative AI in the Enterprise (<https://menlovc.com/perspective/2025-the-state-of-generative-ai-in-the-enterprise/>). Menlo Ventures, 2025-12-19; accessed 2026-07-25.

Chapter 2.1 — Models & the Intelligence Frontier

Two shifts define this layer in 2026, and both cut against the obvious trade. The frontier has moved from training bigger models to making them think longer at the moment of use, which keeps compute demand climbing. And the model itself is commoditizing fast: open and Chinese systems now match the closed frontier on many tasks at a fifth to a hundredth of the price, and by mid-2026 they carried the majority of the world's API tokens. The durable money is in the compute underneath and the distribution on top. "Who has the best model" is a title that changes almost monthly.

Almost everything else in this atlas exists to serve the models in this chapter. The power plants being built, the memory sold out two years forward, the hundreds of billions in capital spending: all of it is downstream of the demand these models create, which is why the report begins here. A foundation model is a large neural network trained on a vast body of text, code, or images. The money it consumes arrives in two very different forms. **Pre-training** is the single heavy run that reads the corpus once and bakes its patterns into billions of parameters. **Inference** is the compute spent every time someone uses the model, generating a reply one token, or word-piece, at a time. Pre-training is a capital cost paid once; inference is an operating cost paid on every query, for as long as the model is used. Which of the two grows faster is the whole investment story of this layer, and the answer flipped in 2025.

The wall, and the pivot that replaced it

For most of the past decade the business ran on a rule reliable enough to underwrite the entire boom: enlarge the model, add data, and it improves at a predictable rate. Through 2025 that rule bent, and the people best placed to see it said so plainly. Ilya Sutskever, who coined the phrase "data is the fossil fuel of AI" at NeurIPS in December 2024, described in a November 2025 interview the close of "the age of scaling" and the opening of an "age of research," where the bottleneck is ideas rather than compute. The products bore him out. OpenAI's GPT-5, released on 7 August 2025, was widely read as a

refined *packaging* of existing capability — a router unifying older models with reasoning modes — rather than a generational leap.

The frontier did not stall; it changed axis. The newest systems improve by *reasoning at the moment of use*, generating long internal chains of thought before they answer, rather than front-loading all their intelligence into a larger pre-trained model. Google's Gemini 3 Pro, released 18 November 2025, drew its reasoning edge explicitly from inference-time "Deep Think" rather than a bigger base model. xAI's Grok 4 was trained with reinforcement learning "at pre-training scale." The economic consequence is the point: a reasoning model can burn ten to a hundred times more tokens on a single query than a plain chatbot reply, and it does so on every use rather than once in training. Epoch AI put inference above 60% of total compute among the major providers by early 2025, and the share has climbed since; industry estimates put 80–90% of *compute dollars* on inference. When a later chapter tells you power has become the binding constraint, the cause traces back here.

There is a genuine debate underneath this. The bears — Sutskever, Gary Marcus — argue that each compute doubling now yields shrinking gains, and that knowledge benchmarks flatten past a certain scale. The bulls counter that a *new* scaling axis opened with its own reliable curve, so the frontier keeps moving; Gemini 3's strength is their exhibit. For an investor the resolution matters less than the direction: whichever camp is right, the compute intensity per query is rising, not falling.

When the benchmarks stopped discriminating

One visible symptom of the shift is that the old benchmarks stopped telling the models apart. MMLU-Pro sits saturated at 88–94% across every frontier model; AIME competition math runs 95–100%, with at least one model reported at a perfect score; GPQA-Diamond graduate science is near-saturated around 87–94%, above the ~70% human-expert baseline. When every leading model scores in the nineties, the benchmark no longer discriminates, so the industry has moved to a new generation of harder tests: autonomous-coding suites (SWE-bench Verified, Terminal-Bench), the deliberately brutal Humanity's Last Exam, the abstract-reasoning ARC-AGI-2, and the crowd-voted LMArena.

Even these are climbing fast. Humanity's Last Exam, designed to be near-impossible, has frontier models above 50%. ARC-AGI-2, on which every frontier model scored roughly 0% at its March 2025 launch, had leaders in the eighties within a year. The practical

consequence is that benchmark leadership is now both fast-moving and hard to verify — the newest model names change monthly and many quoted scores come from aggregator sites of uncertain reliability — so the honest way to read the leaderboard is as a snapshot of a close race rather than a settled order.

Model Providers and Competitive Positioning

Before ranking the models, it helps to know the roster, because these are the names that fill the headlines and earnings calls. The families worth recognizing, and who makes them:

Model family	Maker	Country	Known for	Access
GPT-5.x	OpenAI	US	the consumer default; strong coding	PVT (via Microsoft)
Claude (Opus/Sonnet/Haiku)	Anthropic	US	enterprise & coding leader	PVT (via Amazon, Google)
Gemini 3	Google DeepMind	US	multimodal + reasoning; Search/Workspace reach	GOOGL
Grok 4	xAI	US	real-time, X-integrated	PVT (in SpaceX)
Llama 4	Meta	US	open-weight pioneer (now retreating)	META
DeepSeek V4 / R-series	DeepSeek	China	the efficiency shock; open	PVT
Qwen 3	Alibaba	China	most-downloaded open family	BABA / 9988.HK
Kimi (K-series)	Moonshot AI	China	long-context; open	PVT
GLM	Zhipu AI	China	open, enterprise	PVT
Ernie	Baidu	China	free consumer assistant	BIDU / 9888.HK
Hunyuan / Yuanbao	Tencent	China	super-app distribution	0700.HK
Doubao	ByteDance	China	#1 China consumer app (382M MAU)	PVT
MiniMax / Hailuo	MiniMax	China	agents & video	PVT
Mistral / Magistral	Mistral AI	France	Europe's open-weight champion	PVT

Model family	Maker	Country	Known for	Access
Command	Cohere	Canada	enterprise retrieval	PVT
(no product yet)	Safe Superintelligence (SSI)	US	Sutskever's "age of research" bet	PVT
(no product yet)	Thinking Machines Lab	US	Mira Murati's lab	PVT

Three patterns fall out of the roster. The closed frontier is a handful of US labs. The open-weight world is increasingly Chinese, and it is a crowded field — Alibaba, DeepSeek, Moonshot, Zhipu, MiniMax, and the big platforms (Baidu, Tencent, ByteDance) all ship competitive models. And two of the best-funded new labs, Sutskever's SSI and Murati's Thinking Machines, have no product at all, betting that the next breakthrough comes from research rather than from scaling what exists.

Three scoreboards, three different winners

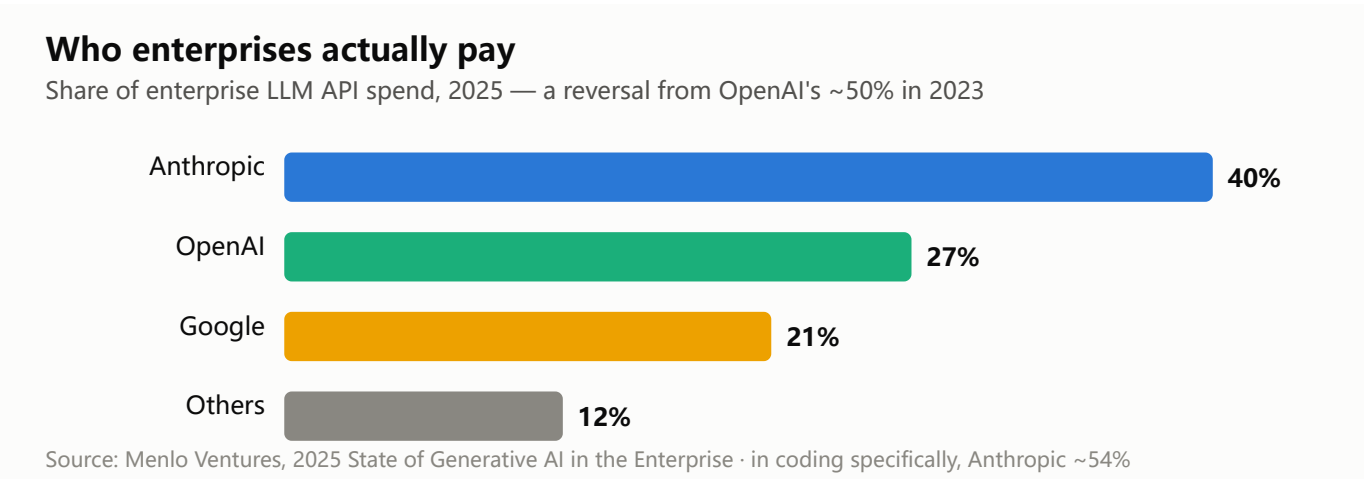
Ask "which model is best" and the honest answer is that it depends which scoreboard you read, and the three that matter give three different answers.

On **raw capability**, the race is close and reshuffles almost monthly. At the 25 July 2026 evidence cutoff, Anthropic's Claude family tops the general-ability and composite-intelligence leaderboards, while OpenAI's GPT-5-series leads autonomous coding.

Scoreboard (observed by 25 July 2026)	Leader	Level	Source
General ability — LMArena Elo	Anthropic Claude family	~1500–1507	arena.ai (live)
Composite — Artificial Analysis Intelligence Index	Claude (Opus tier)	~61 (of ~100)	Artificial Analysis
Hardest expert exam — Humanity's Last Exam	Claude family	~53%	Artificial Analysis
Autonomous coding — SWE-bench Verified ⚠️	OpenAI GPT-5.x	~96%	aggregator ⚠️
Coding agents — Terminal-Bench 2.0	GPT-5.5 harnesses	~85%	tbench.ai
Graduate science — GPQA-Diamond	near-saturated (Grok-4 87% verified)	~87–94%	Epoch AI

Leaderboards move week to week; treat the ranking as a close race, not a fixed order. Scores marked ⚠️ are aggregator-sourced.

The **second scoreboard is enterprise money**, and it tells a cleaner, more durable story. Menlo Ventures' 2025 enterprise survey put Anthropic at 40% of enterprise LLM API spend, OpenAI at 27%, and Google at 21% — the top three taking 88% of a market where foundation-model APIs captured \$12.5B of \$18B in infrastructure spend, and where total enterprise AI spend hit \$37B, up 3.2× year on year.[^models-menlo-2025]



The reversal is the headline: OpenAI held roughly 50% of this market in 2023 and Anthropic 12%. In coding specifically, Anthropic's lead is wider still, near 54% against OpenAI's 21%, up from 42% just six months earlier. Enterprises, in other words, increasingly pay Anthropic even as consumers overwhelmingly use OpenAI — a split that is central to how you value each lab.

Which is the **third scoreboard: consumer reach**, where OpenAI's lead is enormous.

Consumer product	Active users	As-of	Source
ChatGPT	900M weekly	27 Feb 2026	OpenAI[^models-openai-users]
Google Gemini app	750M monthly (+ AI Overviews ~2B)	Feb 2026	TechCrunch
Meta AI	~1B monthly	2026	company ⚠
Grok	~117M monthly (from 35M Dec-25)	Mar 2026	⚠
DeepSeek (China)	81.6M weekly	Feb 2026	Reuters

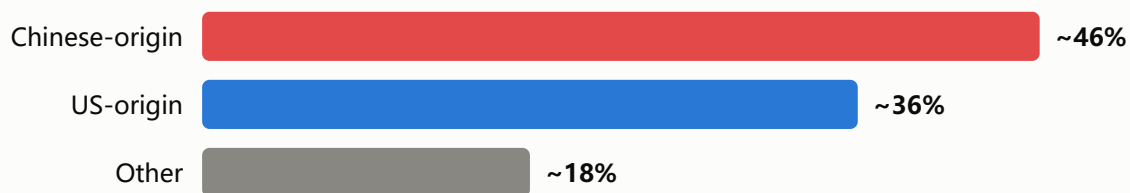
ChatGPT's more than 900M weekly actives, announced by OpenAI in February 2026, dwarf every rival. The three scoreboards resolve into a clean picture: OpenAI owns the consumer surface and the mindshare, Anthropic owns the enterprise wallet and the coding workflow, and Google owns distribution through Search and Workspace. None owns "the best model" for longer than a release cycle.

The developer layer is quietly going Chinese

The scoreboard most investors miss is the developer one, and it has moved against the United States. On OpenRouter, a large model-routing marketplace, models of Chinese origin carried roughly 46% of identified tokens by June 2026 against about 36% for US-origin models, and open-weight models overtook proprietary ones to roughly a 60/40 split. DeepSeek was the single largest provider on the platform.

The developer layer is quietly going Chinese

Share of OpenRouter tokens by model origin, June 2026 (🚩 aggregator)



Source: OpenRouter via aggregators (🚩). US-model token share fell from ~70% (Jun 2025) to ~30% (Jun 2026); DeepSeek is the sing

The trajectory is the striking part: US-model token share on the platform fell from around 70% in June 2025 to roughly 30% a year later. On Hugging Face, Alibaba's Qwen passed 700M downloads by January 2026, became the most-downloaded open family, and spawned more fine-tuned derivatives than Google and Meta combined; by one count Qwen derivatives were over 40% of all new language-model variants. These are aggregator figures and the exact percentages will move, but the direction is unmistakable and corroborated across sources.

This is strategy, not accident. Denied the best training silicon, China made openness an asymmetric weapon: a capable model anyone can download and run cheaply spreads Chinese technical influence, undercuts the pricing power of the American closed labs, and grows an ecosystem that needs no one's permission. DeepSeek's arrival in early 2025 briefly erased hundreds of billions in US AI market value precisely because it questioned the premise that compute scarcity was a durable moat. There is a wrinkle worth flagging: reports in 2026 that both Meta (pausing open Llama) and Alibaba (moving flagship Qwen toward closed) are partially reversing their open strategies, which, if it holds, would reshape this dynamic — treat it as developing.

The price of intelligence is collapsing

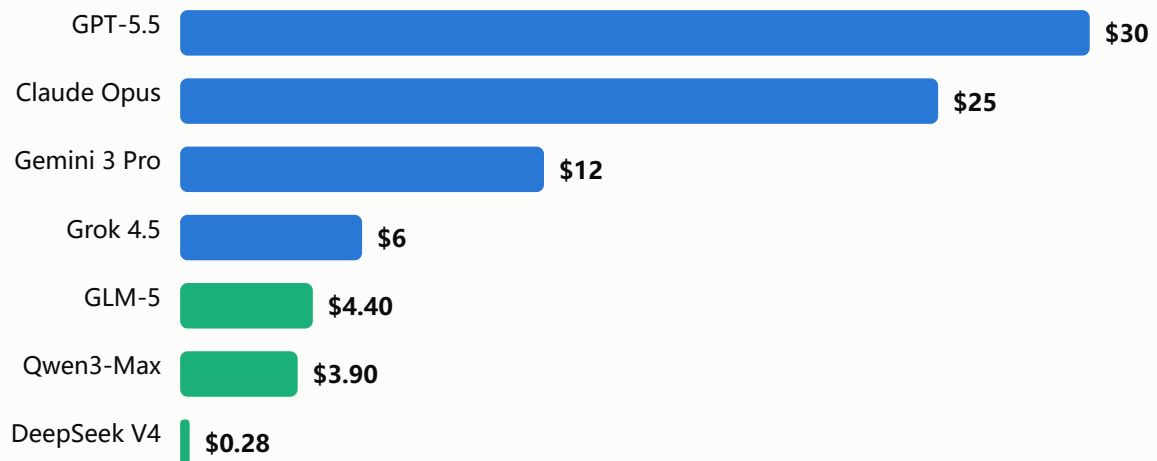
The reason the model layer is a hard place to make money is that the thing it sells is deflating faster than almost any product in history. Stanford's AI Index found the cost of GPT-3.5-level inference fell more than 280× between late 2022 and late 2024, from about \$20 to about \$0.07 per million tokens.^[^models-stanford-cost] Epoch AI's work on frontier pricing finds the cost to reach a *fixed* capability level fell between 9× and 900× per year across specific milestones, while warning that the fastest recent declines may not persist.^[^models-epoch-cost] Anthropic launched Claude Opus 4.5 in November

2025 at roughly a 67% price cut to its predecessor; Chinese labs cut prices around sixfold in the first half of 2026, declaring several cuts permanent.

The price of intelligence is collapsing

■ US closed-frontier ■ Open / Chinese

Output price per 1M tokens, mid-2026 — closed frontier vs open/Chinese models



Source: aggregator pricing (🚩, mid-2026) · frontier-closed runs ~5×–100× the price of leading open models for near-comparable q

The live spread makes the commoditization concrete: frontier-closed output runs roughly 5× to 100× the price of the leading Chinese-open models for capability that, on many tasks, is within a release cycle. Alongside the price collapse, the models have grown vastly more capacious: one-million-token context windows became generally available on the frontier models in 2026, with some pushing toward two million and Meta's Llama 4 Scout advertising ten million, though independent tests show *effective* context is closer to 50–65% of the advertised figure.

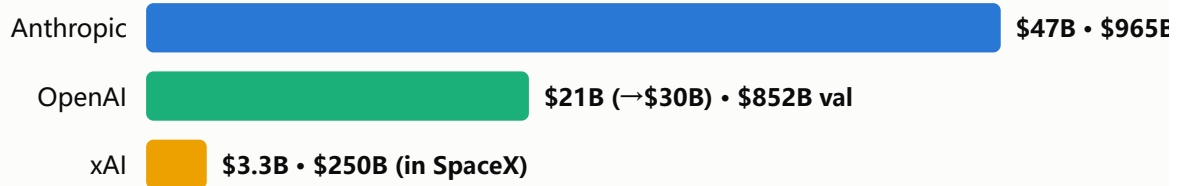
The central investment question in this layer is whether cheaper compute *shrinks* total spending or *expands* it. Current evidence favors expansion: efficiency improves rapidly, but usage has grown faster, so aggregate compute and capital spending continue to rise even as the cost of a query falls. That relationship is an observed regime, not a physical law. A large efficiency shock or slower usage growth would test whether deflation reaches infrastructure suppliers rather than remaining concentrated in model margins and token pricing.

The money behind the models

The valuations attached to this layer price a future the revenue has not yet reached, though the revenue is growing fast.

The money behind the models

Annualized revenue run-rate, mid-2026 (\$B); private valuation in the label



Source: Epoch AI dataset, CNBC/Bloomberg (2026). Valuations are illiquid marks, not prices. Anthropic's run-rate went ~\$9B→\$47B in

OpenAI's annualized revenue crossed roughly \$20B during 2025 and stood at about \$21.4B at year-end per Epoch AI's dataset, tracking toward a ~\$30B target for 2026, against reported losses near \$14B a year and a private valuation of \$852B set in a \$122B round in March 2026, with a confidential IPO filing reported for mid-2026. Anthropic's run-rate went from roughly \$9B at the end of 2025 to over \$47B by May 2026, on which it raised at a \$965B valuation, the first challenger to exceed OpenAI's mark. xAI reported about \$818M in the first quarter of 2026, roughly a \$3.3B run-rate, against a \$2.47B operating loss, and was valued at \$250B inside SpaceX after an all-stock merger.

The Anthropic figure is the one to sit with, because it is the clearest evidence in the whole atlas that AI monetizes durably where the workflow is deep and technical. Its Claude Code product went from a \$500M run-rate in September 2025 to over \$2.5B by February 2026, passed 2M weekly active users, and by April 2026 had over 500 customers spending more than \$1M a year. Coding is where the money is, and it is the empirical anchor for the demand-durability debate that runs through Chapter 2.2. The valuations, by contrast, should be read as illiquid marks with liquidation preferences, not as market caps — they price AGI optionality, and the discipline is to own the labs, where possible, through their listed backers rather than through premium-priced access vehicles.

The release treadmill

The pace of releases is itself a feature of this layer, and a strain on everything beneath it. The cadence over the past year:

Date	Release
Aug 2025	OpenAI GPT-5; gpt-oss-120B/20B (OpenAI's first open-weight in six years)
Sep 2025	Anthropic Claude Sonnet 4.5; DeepSeek V3.2
Oct 2025	Anthropic Claude Haiku 4.5
Nov 2025	OpenAI GPT-5.1; Google Gemini 3 Pro + Deep Think; Anthropic Claude Opus 4.5 (–67% price)
Dec 2025	OpenAI GPT-5.2
Apr 2026	GPT-5.4; Gemini 3.1 Pro; DeepSeek V4 ("near-frontier, a fraction of the price")
May 2026	OpenAI GPT-5.5; Anthropic Claude "Mythos" preview
Jun–Jul 2026	Claude Sonnet 5; further Claude, GPT-5.6, Grok 4.5, Kimi K3 releases ⚠️

Meta's Llama 4 (Scout and Maverick) shipped, but its flagship "Behemoth" was effectively shelved, never formally released or cancelled. The relentless roughly-quarterly cadence is why benchmark leadership is so fleeting, and why the strain on the supply chain beneath — the chips, packaging, and power of later chapters — never lets up.

Agents: the capability is racing, the reliability is not

The product story of 2026 is agents — models that take actions over long horizons rather than only answering — and the early revenue is genuine. But this is the layer's clearest example of capability outrunning dependability.

Capability is improving faster than almost anywhere in AI. METR, which measures the length of task an agent can complete at 50% reliability, estimates that its broad historical horizon doubled roughly every seven months over the six years through 2025. Newer domain-specific slices can imply faster rates, but the estimate is sensitive to the task mix; METR explicitly warns that measurements beyond sixteen hours exceed the reliable range of its current suite.^[^models-metr-horizon] Anthropic's Model Context Protocol, an open standard for connecting agents to tools, was donated to the Linux Foundation and adopted by OpenAI, Google, and Microsoft, a genuine interoperability milestone.

The catch is reliability, not capability. A step that is 95% reliable, compounded across a long task, fails too often for production. MIT's Project NANDA found roughly 95% of enterprise generative-AI pilots delivered no measurable profit-and-loss impact; by most estimates 86–88% of agent pilots never reach production; and Gartner projects that more than 40% of agentic-AI projects will be canceled by the end of 2027 on cost, unclear value, and inadequate controls. Amazon's own AGI director stated flatly that it is agent reliability, not capability, that keeps enterprises in pilot. The gap between a demo that works most of the time and a system a business can trust with unattended work is where the next wave of value, and a new tooling layer built to make agents dependable, will be won. This is the same reliability gate that Chapter 2.2 treats from the application side.

Beyond text: video, world models, and robots

The models layer is expanding sideways as fast as it is advancing, and a whole roster of companies sits in these adjacent frontiers. **Video generation** matured into a production tool over 2025–26: OpenAI's Sora, Google's Veo, Kuaishou's Kling and ByteDance's Seedance in China, and MiniMax's Hailuo now produce short clips with synced audio and plausible physics, and studios routinely route a scene across two or three of them. Around them sits a field of **creative-media** specialists worth recognizing: Midjourney and Black Forest Labs (the Flux models) for images, Runway, Pika, and Luma for video, Suno for music, ElevenLabs for voice, and Stability AI.

A newer frontier is **world models** — systems that simulate an interactive environment rather than describe one — which many researchers see as the path to physical intelligence. DeepMind's Genie 3 generates real-time, persistent 3D worlds; Fei-Fei Li's World Labs, which raised \$1B in February 2026 from backers including AMD and NVIDIA, is building the same for simulation; and NVIDIA's Cosmos supplies world-model foundations to robotics developers. That connects directly to **embodied AI**, where "vision-language-action" models drive humanoid robots. The model-makers here — Physical Intelligence, Skild AI, and NVIDIA's GR00T — feed machines from Figure, which put around forty humanoids into a BMW plant at a \$39B valuation, and Tesla's Optimus. Robotics startups raised roughly \$8.5B in 2025.

For an investor these adjacencies matter because they are the next compute driver: video and world-model training and inference are far heavier than text, so they extend the demand thesis of this whole layer. NVIDIA is the full-stack expression — its Cosmos,

GR00T, and Omniverse tools make it a robotics-and-world-model play as much as a chip company — while most of the pure-plays (World Labs, Physical Intelligence, Figure, Runway, Kling via Kuaishou) are private or embedded in larger listed parents.

The data wall and the scramble for new inputs

The quietest constraint in this layer is running out of things to learn from. The supply of high-quality human text is finite and largely exhausted — the "data wall" — which is pushing training toward two substitutes: synthetic data generated by models themselves, and reinforcement-learning "environments," interactive worlds in which a model practices a task against a verifiable reward. The post-training recipe has re-based from human-feedback tuning to reward-based methods plus synthetic self-play; essentially every recent frontier model uses a distinct variant.

This works where the reward is verifiable — code that passes its tests, math with a checkable answer, a game with a score — and is riskier for open-ended language, where recursively training on model output risks "model collapse," a degradation of the distribution. So the wall is being bypassed in *verifiable* domains, not universally, which is part of why capability gains have concentrated in coding, math, and tool use. A small industry has been funded to supply these training grounds: Mechanize (building coding environments with Anthropic), Prime Intellect (positioned as a "Hugging Face for RL environments"), and Mercor and Surge expanding from static labeling. The Information reported Anthropic leaders discussing more than \$1B of spending on RL environments over the coming year. For an investor this is a genuine new input market and a picks-and-shovels layer, and also a possible early bubble pocket if the valuations attached to it (some already in the tens of billions) prove ahead of the reality.

The trade: long the compute, rent the model

The evidence supports more confidence in aggregate infrastructure demand than in any standalone model business. Reasoning, agents, and video increase inference demand and the power, networking, and memory beneath it regardless of which lab leads a given month. That is a strategic demand conclusion, not a security recommendation; the decision layer separately tests valuation, equity capture, and downside.

The labs are the subtler call. OpenAI is the consumer franchise and Anthropic the enterprise and coding franchise, both reachable only through their backers — Microsoft, Amazon, Google, and Nvidia, detailed in Part V — and both worth owning for distribution and product rather than for model supremacy, because the middle of the model market is commoditizing underneath them at 5× to 100× lower prices. On the open-weight side, Alibaba is the strategically significant name, and the developer-token data suggests Chinese open models are winning a layer the US assumed it led, which makes Alibaba a real, if China-risk-laden, way to own the open frontier. The RL-environments vendors are an early, private, higher-risk expression of the data-wall trend. Throughout, the discipline is to treat the top benchmark score as noise and to treat enterprise API share, developer-token share, and token-efficiency as signal.

What would break this read

Three developments would break this read. A real *pre-training re-acceleration* — a model dramatically better because it is bigger rather than because it thinks longer — would re-concentrate the advantage around whoever owns the most compute and reward the model layer all over again, inverting the "rent the model" conclusion. A *second efficiency shock* larger than DeepSeek's could tip the near-term balance from expansion toward deflation and puncture the compute-scarcity case that underwrites much of this atlas. And *agents crossing the reliability threshold* into broad production would pull enterprise revenue forward and validate the entire build faster than the market expects — a boon to everything downstream, and to the application layer of Chapter 2.2 most of all. A quieter risk sits under the open-weight thesis: if the reported Meta and Alibaba retreats from open models hold, the developer-layer story could reverse as fast as it arrived.

Chapter 2.1 endnotes

[^models-menlo-2025]: 2025: The State of Generative AI in the Enterprise (<https://menlovc.com/perspective/2025-the-state-of-generative-ai-in-the-enterprise/>). Menlo Ventures, 19 December 2025. Survey and bottom-up market estimate; not audited issuer revenue.

[^models-openai-users]: Scaling AI for everyone (<https://openai.com/index/scaling-ai-for-everyone/>). OpenAI, 27 February 2026. Company-reported product metric.

[^models-standford-cost]: Artificial Intelligence Index Report 2025 (https://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf). Stanford HAI, 7 April 2025.

[^models-epoch-cost]: LLM inference prices have fallen rapidly but unequally across tasks (<https://epoch.ai/data-insights/llm-inference-price-trends>). Epoch AI, 12 March 2025.

[^models-metr-horizon]: Task-Completion Time Horizons of Frontier AI Models (<https://metr.org/time-horizons/>). METR, updated 8 May 2026.

Full data and sourcing: the models data appendix (research/2.1-models-data.md) and the master figures table (§B, §H). Items marked ⚠ — the newest model-name leaderboard positions, OpenRouter shares, and some run-rates — are aggregator-sourced or contested and flagged as such. Company tags and access: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [^atlas-chart-model-pricing-snapshot] [^atlas-chart-openrouter-ranking] [^atlas-chart-private-ai-revenue]

[^atlas-chart-model-pricing-snapshot]: Model pricing snapshot (https://openrouter.ai/models?fmt=cards&o=pricing-high-to-low&output_modalities=text). OpenRouter, undated; accessed 2026-07-25.

[^atlas-chart-openrouter-ranking]: Chinese AI models gain ground on U.S. rivals as low costs drive adoption (<https://www.cnbc.com/2026/07/07/chinese-ai-models-costs-us-openai-anthropic.html>). CNBC, 2026-07-07; accessed 2026-07-25.

[^atlas-chart-private-ai-revenue]: AI companies: revenue and valuation data (<https://epoch.ai/data/ai-companies>). Epoch AI, 2026-07-21; accessed 2026-07-25.

Chapter 2.2 — Applications & Monetization

This is the layer where the entire build must eventually pay for itself, and it holds the biggest open question in the atlas. The bearish headline is real: by one widely-cited study, about 95% of enterprise AI pilots produce no measurable profit. The bullish reality sits underneath it: a handful of application companies are compounding toward billion-dollar revenue, enterprise AI spending tripled in a year to \$37B, and the money concentrates precisely where the workflow is deep, coding above all. Monetization is not failing; it is bifurcating, and telling the two halves apart is the whole job here.

Everything below this chapter is cost. Applications are where AI becomes a product someone pays for, either a consumer subscription or an enterprise contract, and the durability of that revenue is what ultimately justifies the hundreds of billions being spent on chips and power. The question that hangs over the whole industry is whether enterprises are getting a return, because if they are not, the capex of Chapter 2.11 has no foundation. The honest answer is that it depends entirely on what the AI is doing.

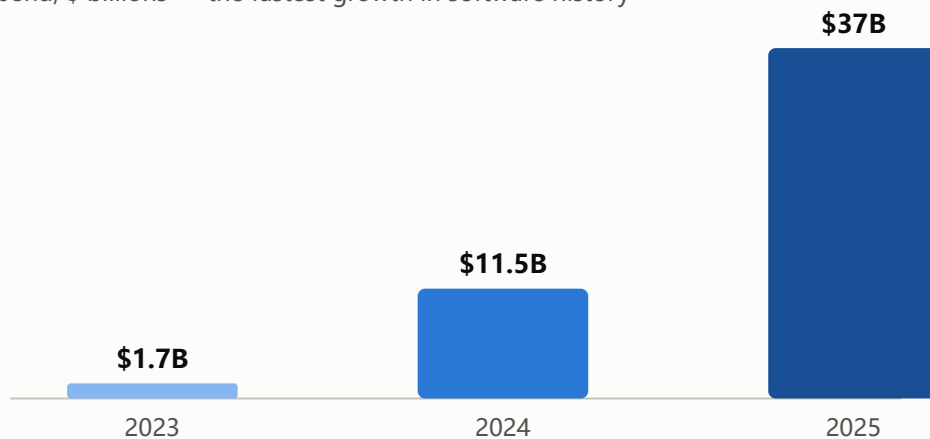
The 95% headline, and what it hides

The bearish case has a number attached. MIT's Project NANDA, in an August 2025 study, found that roughly 95% of organizations saw no measurable profit-and-loss return on their generative-AI pilots despite \$30–40B of spending, with only about 5% of integrated projects extracting real value. The study's own explanation matters: the failures came from tools that did not fit workflows or retain context, not from weak models. That distinction is the key to the whole layer.

Set against that failure rate is an equally real explosion in spending. Enterprise AI spend reached \$37B in 2025, up 3.2 times in a single year from \$11.5B, the fastest growth in the history of enterprise software.

Enterprise AI spending tripled in a year

Total enterprise AI spend, \$ billions — the fastest growth in software history



Source: Menlo Ventures, 2025 State of Generative AI in the Enterprise. Applications took \$19B (51%); coding alone ~\$4B.

The two facts reconcile once you look at where the money goes. Of that \$37B, applications took \$19B, slightly more than half, and within applications the spend is heavily concentrated: coding alone accounts for about \$4B of the departmental total, healthcare \$1.5B, legal \$650M. AI deals now convert at roughly 47% versus 25% for traditional software. So the picture is not "AI does not pay." It is "AI pays handsomely in a few deep workflows and disappoints in shallow, horizontal ones," and the market is already sorting winners from the rest. A further tell: startups now capture about 63% of application-layer revenue, up from 36% a year earlier, earning roughly two dollars for every one the incumbents earn.

Application-Layer Market Structure

The winners cluster by vertical, and the roster is worth recognizing because these are the names that dominate AI business headlines.

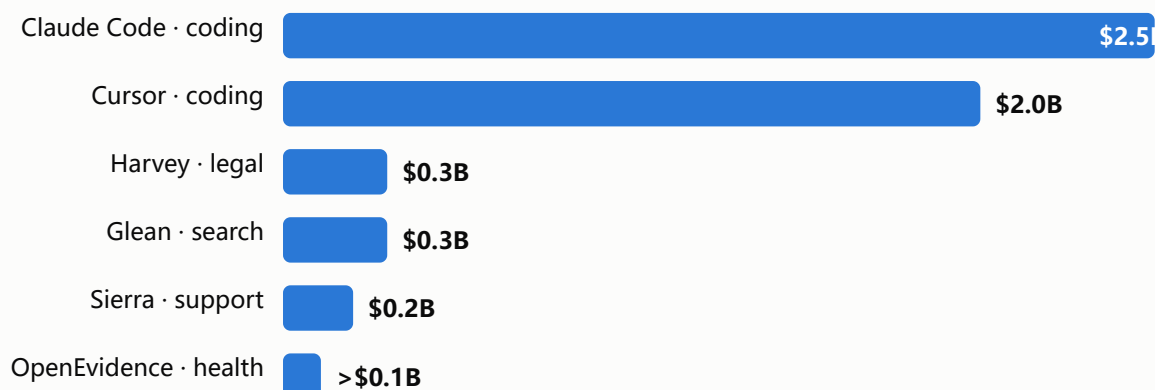
Company	Vertical	Scale (approx.)	Access
GitHub Copilot	coding	~20M users, 4.7M paid	MSFT
Cursor (Anysphere)	coding	~\$1B ARR; \$29B val	PVT
Claude Code (Anthropic)	coding	~\$2.5B run-rate	PVT
Cognition (Devin) / Windsurf	coding	~\$490M run-rate	PVT
Sierra	customer support	~\$200M ARR; > \$15B val	PVT
Decagon / Intercom Fin	customer support	Decagon \$4.5B val	PVT
Harvey / Legora	legal	Harvey ~\$300M ARR / \$11B val	PVT
Abridge	healthcare (scribing)	\$117M contracted ARR	PVT
OpenEvidence	healthcare (clinical Q&A)	~40% of US physicians; \$12B val	PVT
Glean	enterprise search	~\$300M ARR; \$7.2B val	PVT
Clay	sales / GTM	~\$100M ARR	PVT
Salesforce Agentforce	enterprise agents	~\$1.2B ARR	CRM
ServiceNow Now Assist	enterprise agents	~\$750M ACV	NOW
Palantir	applied-AI platform	AI-product revenue not separately disclosed	PLTR
Doubao / DeepSeek / Ernie / Yuanbao	China consumer	Doubao 382M MAU	PVT / BIDU / 0700.HK

Who is actually making money

The application layer has produced a cohort of companies with real, fast-growing revenue, most of them still private, and the pattern is that the winners own a specific, high-value workflow end to end.

The app layer is real — annualized revenue of the leaders

Approx. run-rate, mid-2026; mostly private companies owning one deep workflow



Source: company disclosures and reputable reporting through 25 July 2026. Private-company run-rates are unaudited estimates.

In **coding**, the standouts are Anthropic's Claude Code, at roughly a \$2.5B annual run-rate by early 2026, and Cursor (Anysphere), which crossed \$1B in annualized revenue and raised at a \$29.3B valuation, with GitHub Copilot's 20M-plus users and Cognition's Devin (which absorbed Windsurf) rounding out the field. In **customer support**, Sierra reached about \$200M in revenue at a valuation above \$15B, and Decagon tripled to a \$4.5B valuation, both on outcome-based pricing. In **legal**, Harvey grew to about \$300M of run-rate at an \$11B valuation, with Legora close behind. In **healthcare**, Abridge's ambient clinical documentation reached \$117M of contracted revenue across Kaiser, Mayo, and Johns Hopkins, and OpenEvidence, used by roughly 40% of US physicians, doubled to a \$12B valuation. In **enterprise search and sales**, Glean tripled to over \$300M and Clay reached around \$100M with net revenue retention above 200%. The common thread is defensibility through workflow depth and proprietary data, not model access.[^atlas-chart-anthropic-claude-code][^atlas-chart-cursor-arr][^atlas-chart-glean-arr][^atlas-chart-private-app-arr]

The pricing revolution: from seats to outcomes

A quiet but important shift underpins the whole layer: when an agent does the work a person used to do, charging per seat stops making sense. Seat-based pricing fell from 21% to 15% of vendors in a year, while hybrid and usage-based models rose to 41%. The vanguard is outcome-based pricing, where the vendor is paid only when the AI delivers a result: Intercom's Fin support agent charges \$0.99 per resolved ticket, and Decagon prices per resolution. This transition reprices the entire software industry. It is a tailwind for the AI-native winners, who can capture a share of the labor budget rather than the software budget, and a headwind for any incumbent whose value was the number of seats it sold, which is why the enterprise-software incumbents are racing to bundle agents into every tier.

Agents in production, and the reliability gate

The next phase of this layer is agents, software that takes actions rather than just answering, and the early revenue is genuine. Salesforce's Agentforce reached a \$1.2B run-rate, up more than 200% year on year, and ServiceNow's Now Assist passed \$750M in annual contract value and restructured its entire product line around AI tiers. Real deployments show the potential: Klarna's OpenAI-powered assistant handled 2.3M conversations, the equivalent of 700 human agents, and cut resolution time from eleven minutes to under two, an estimated \$40M profit improvement, though Klarna later rebalanced toward a human-plus-AI hybrid after over-automating.[^atlas-global-menlo-enterprise-ai-2025]

The constraint, again, is reliability rather than capability, and it is the same gate described in Chapter 2.1 from the model side. Gartner projects that more than 40% of agentic-AI projects will be canceled by the end of 2027 on cost, unclear value, and inadequate controls, even as it forecasts that a third of enterprise applications will embed agents and 15% of day-to-day work decisions will be autonomous by 2028. Only about 16% of enterprises run "true" plan-and-adapt agents today; the rest are fixed workflows. The signal to watch is the conversion rate from pilot to production, which remains in the single digits for genuine autonomous agents.

America sells subscriptions, China gives it away

The national contrast in this layer is about business model as much as capability. The United States monetizes applications directly, through enterprise API contracts and consumer subscriptions; ChatGPT alone has around 50M paying subscribers, and the \$19B of enterprise application spend flows to a dense field of \$100M-to-\$1B revenue startups. China largely does not charge. Its leading assistant, ByteDance's Doubao, reached 382M monthly users, and DeepSeek's app around 130M, but the dominant model is free access funded by advertising and the broader super-app ecosystem. Baidu made its Ernie assistant fully free in 2025 after subscriptions failed, and ByteDance only began piloting paid Doubao tiers in mid-2026, priced from around \$10 a month.

The consequence is that China's application layer generates enormous usage and little direct revenue, while its models spread through free distribution, the open-weight strategy of Chapter 2.1 applied to consumer products. For an investor this means the Chinese app opportunity is monetized indirectly, through the platform companies' ecosystems and advertising, not through the standalone application revenue that defines the US winners.

The thin-wrapper question, and where value accrues

A recurring skepticism is that most AI applications are "thin wrappers" around someone else's model, with no defensible moat. The concern has teeth: application gross margins are estimated at 50–60% against 70–90% for traditional software, compressed every time the underlying model gets cheaper, and the time for a capability to commoditize has fallen from about eighteen months to twelve. The counter-evidence is the \$19B enterprises actually spent, the startups' 63% share of app-layer revenue, and the ten-plus applications already above \$1B in ARR. The resolution is that value accrues not to model access, which is universal, but to distribution, proprietary data, and workflow lock-in. As one framing put it, "every company is now an AI wrapper, so go-to-market is the new moat." The durable winners own a specific workflow so deeply that switching is painful; the vulnerable ones sell a feature a model provider can absorb.

The adoption tests that matter through 2028

The load-bearing question for the whole atlas is whether enterprise ROI broadens from the deep workflows to the shallow ones. Watch the pilot-to-production conversion rate, and whether a follow-up to the MIT study shows the 5% success figure rising. Watch whether the coding-and-vertical winners sustain their growth and margins as the model layer beneath them commoditizes, which would confirm that workflow depth is a real moat. And watch the agent-reliability threshold, because agents crossing into dependable production is what would turn the demand-durability question from open to settled, and pull forward the revenue that underwrites everything downstream.

The trade: own the workflow, not the wrapper

The uncomfortable truth is that most of the clearest winners in this layer are private, reachable only through their venture backers or eventual IPOs. The public expressions are the incumbents embedding AI into existing distribution: Salesforce (CRM) and ServiceNow (NOW) in enterprise agents, Microsoft (MSFT) through Copilot, and the platform companies of other chapters, with Palantir (PLTR) as the high-multiple, high-volatility applied-AI pure-play. Anthropic's coding-driven profitability, reachable through Amazon and Google, is the single best public-market evidence that the application layer monetizes durably where the workflow is deep. The discipline is to favor owners of a specific high-value workflow with proprietary data over thin wrappers whose only asset is model access, because the latter's margins compress every time the underlying model gets cheaper.

What would settle the ROI question

The bullish read on this layer, and by extension on the whole capex thesis, breaks if enterprise ROI stays stuck, if the MIT 95% figure fails to improve and the winners stay confined to coding and a couple of verticals rather than broadening. It also breaks if the agent-reliability gate does not open, leaving Gartner's 40% cancellation forecast to play out and stranding the spend. In the other direction, the thesis strengthens materially if agents cross into broad production and outcome-based pricing lets AI capture a share of labor budgets rather than software budgets, which would expand the addressable market by an order of magnitude and validate the build faster than the market expects.

Data and sourcing: the applications data appendix (research/2.2-applications-data.md) and the master figures table (§H). Tier-1 anchors: Menlo Ventures enterprise survey, MIT Project NANDA, Gartner, CNBC/Bloomberg valuations, Salesforce and ServiceNow filings. Private valuations and ARR figures move monthly and several are flagged ⚠ in the appendix. Company tags and access: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [^{^atlas-chart-anthropic-claude-code}] [^{^atlas-chart-cursor-arr}] [^{^atlas-chart-glean-arr}] [^{^atlas-chart-private-app-arr}] [^{^atlas-global-menlo-enterprise-ai-2025}]

[^{^atlas-chart-anthropic-claude-code}]: Anthropic raises \$30 billion in Series G funding at \$380 billion post-money valuation (<https://www.anthropic.com/news/anthropic-raises-30-billion-series-g-funding-380-billion-post-money-valuation>). Anthropic, 2026-02-12; accessed 2026-07-25.

[^{^atlas-chart-cursor-arr}]: Cursor Recurring Revenue Doubles in Three Months to \$2 Billion (<https://www.bloomberg.com/news/articles/2026-03-02/cursor-recurring-revenue-doubles-in-three-months-to-2-billion>). Bloomberg, 2026-03-02; accessed 2026-07-25.

[^{^atlas-chart-glean-arr}]: Glean Surpasses \$300M ARR (<https://www.glean.com/press/glean-surpasses-300m-arr-unrivaled-enterprise-context-fuels-ai-adoption>). Glean, 2026-05-28; accessed 2026-07-25.

[^{^atlas-chart-private-app-arr}]: Big Ideas 2026 (https://research.ark-invest.com/hubfs/1_Download_Files_ARK-Invest/Big_Ideas/ARKInvest%20BigIdeas2026.pdf). ARK Invest, 2026-02-01; accessed 2026-07-25.

[^{^atlas-global-menlo-enterprise-ai-2025}]: 2025: The State of Generative AI in the Enterprise (<https://menlovc.com/perspective/2025-the-state-of-generative-ai-in-the-enterprise/>). Menlo Ventures, 2025-12-19; accessed 2026-07-25.

Chapter 2.3 — Power & Data Centers

Three years ago the scarce input in AI was the GPU. By mid-2026 it is the electron, and the unglamorous gear that moves it. Global data-center electricity is on track to roughly double by 2030, US grid-connection queues stretch past eight years, and the transformers and gas turbines needed to energize a site now carry multi-year waits. The trade has migrated from silicon to the grid, and the companies that can deliver a megawatt on a schedule now capture the rent. This is also where the two superpowers diverge most sharply: America is power-constrained, and China is not.

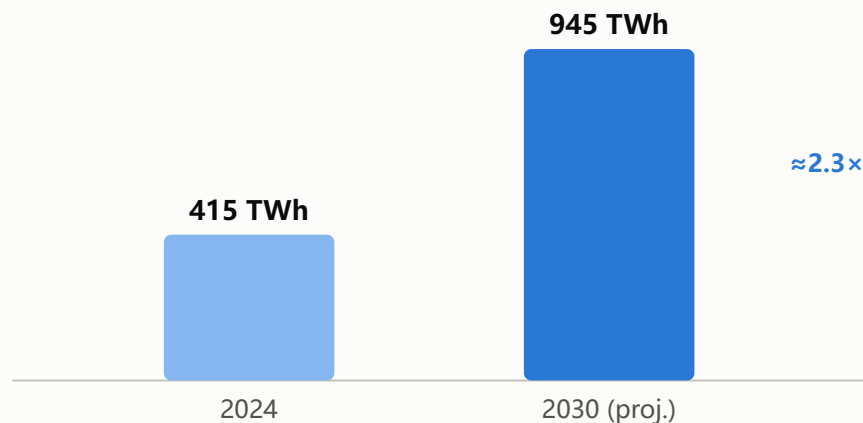
Every chapter so far has described things that consume electricity. This one is about whether the electricity exists. A modern AI data center is less a building full of computers than a substation with servers attached: a single frontier campus now draws as much power as a mid-sized city, and the constraint on building more of them has shifted decisively from the chips inside to the power outside. Jensen Huang and Sam Altman both say it plainly now, the binding limit on AI is electricity, not compute. Understanding this layer means understanding two things in sequence: how big the demand has become, and why the supply cannot keep up.

The demand wall

The numbers are stark and they come from sober sources. The International Energy Agency puts global data-center electricity use at roughly 415 TWh in 2024, rising to about 945 TWh by 2030, close to 3% of all electricity on earth and roughly a 2.3-fold increase in six years.^[^power-iea-energy-ai] In the United States, the Department of Energy's Berkeley Lab found data centers consumed about 4.4% of national electricity in 2023 and projected 6.7% to 12% by 2028.^[^power-lbnl-us] Nearly half of all US electricity-demand growth this decade traces to data centers.

The demand wall — global data-center electricity

Terawatt-hours per year; $\sim 2.3\times$ in six years, toward $\sim 3\%$ of all electricity on earth



Source: IEA, Energy and AI (2025). US data centers alone add ~ 240 TWh; nearly half of US demand growth this decade.

Demand of that size runs straight into a grid that cannot connect it fast enough. Berkeley Lab's 2025 survey found roughly 2,060 GW of generation and storage waiting in US interconnection queues, about twice the entire existing US generating fleet, with median waits around five years nationally and more than eight in PJM, the mid-Atlantic grid. The price signal is unmistakable: PJM's capacity auction cleared at \$329.17 per megawatt-day for 2026/27, against \$28.92 two years earlier, roughly a tenfold jump, and data centers accounted for about 40% of the resulting \$6.3B cost. Morgan Stanley estimates US data-center demand reaches about 74 GW by 2028 against a power shortfall near 49 GW, and Gartner projects that around 40% of AI data centers will be power-constrained by 2027. There is a skeptics' case, worth holding: interconnection queues are inflated by duplicate and speculative "phantom" requests, and more than 750 GW of requests were withdrawn in 2025, so some of the demand is not real. But the price and the withdrawals point the same way: the scarcity is real, and it is priced.

Power Infrastructure and Key Suppliers

The build divides into those who *generate* the power, those who *make the gear* that moves it, those betting on *new nuclear*, and those who *build* the sites. The roster worth recognizing:

Player	Ticker	Role
Constellation	CEG	nuclear IPP; restarting Three Mile Island
Vistra	VST	IPP; nuclear + gas fleet
Talen Energy	TLN	IPP; data-center PPAs
GE Vernova	GEV	gas turbines + grid equipment
Siemens Energy	ENR.DE	gas turbines + grid
Mitsubishi Power	7011.T	gas turbines (sold out to 2028)
Eaton	ETN	electrical distribution gear
Vertiv	VRT	power + liquid cooling
Schneider Electric	SU.PA	electrical / data-center systems
Quanta Services	PWR	electrical construction / EPC
nVent, Powell	NVT, POWL	connection/protection, switchgear
Cummins, Caterpillar	CMI, CAT	interim gensets / engines
Oklo, NuScale	OKLO, SMR	small modular reactors (pre-revenue)
X-energy, Kairos	PVT (Amazon, Google)	SMRs backed by hyperscalers
Cameco	CCJ	uranium / nuclear fuel

The bridge: gas and nuclear restarts

Solving a power shortage takes years, so the industry has split its response into a near-term bridge and a long-term endgame; distinguishing the two is essential for underwriting this layer.

The bridge is natural gas and nuclear restarts. Behind-the-meter gas, generation built on-site to bypass the grid queue, can be deployed in roughly eighteen months, and around 101 GW of it has been announced, though only about 2 GW is actually operating so far. The catch has moved to the turbine. GE Vernova's gas-turbine backlog reached

116 GW in mid-2026, targeting 125 GW by year-end; Siemens Energy is carrying a record order book above €130B; and Mitsubishi Power is sold out through 2028, so new reservations now book four to five years out. The interim gap is being filled with reciprocating engines from Cummins and trucked mobile turbines. Nuclear restarts are the other real near-term lever: Constellation is restarting the former Three Mile Island Unit 1 (its Crane Clean Energy Center) on a twenty-year Microsoft contract, fast-tracked to 2027 and backed by a \$1B federal loan, and Holtec brought the Palisades plant in Michigan back online at the end of 2025, the first-ever US reactor restart.

The endgame: small modular reactors, and the timeline gap

The endgame is small modular reactors, and here timing discipline matters because the market can price a 2030s technology as if it solves the 2027 crunch. The marquee deals are real but distant: Amazon with X-energy targeting 5 GW by 2039, Google with Kairos aiming at 500 MW by 2035, Meta committing to 1.2 GW from Oklo. Oklo itself broke ground at Idaho National Laboratory in 2025 and guides first power to late 2027 at the earliest, has no revenue and no reactor-design approval, and its market value has swung between roughly \$5B and \$13B on sentiment alone. When Truist initiated coverage of Oklo, NuScale, and Nano Nuclear in 2026, it rated all three Hold and told investors they "want proof." The timing distinction is straightforward: restarts and gas can add megawatts this decade; most new nuclear belongs to the next.

Cooling: the wall inside the rack

Between the power source and the chips sits cooling, which has quietly become mandatory rather than optional. Nvidia's GB200 and GB300 rack-scale systems use direct-to-chip liquid cooling, with public rack-power estimates clustering around roughly 120–150 kW depending on configuration and measurement boundary. Nvidia had not published a comparable Vera Rubin rack-power specification at the cutoff, so precise next-generation figures remain estimates rather than product facts. Vertiv is the scaled incumbent, and its reference architecture claims large energy, footprint, and space savings; a field of specialists sits around it, including CoolIT, Boyd, LiquidStack, and Submer. Direct-to-chip has won the current generation while immersion cooling remains a niche waiting for the density that may eventually force it, and the coolant-distribution

units that feed these systems are themselves a near-term supply constraint. The normalized package-to-rack-to-100-MW comparison is in \$2.13.

The gigawatt campus

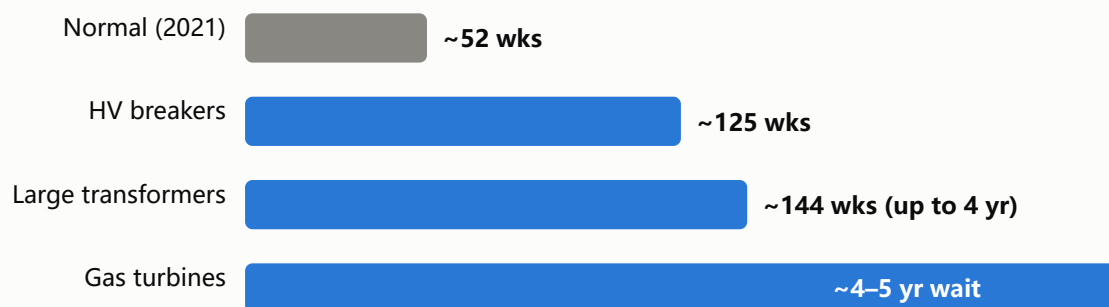
The demand shows up physically as a new class of building: the multi-gigawatt campus. OpenAI and Oracle's Stargate flagship in Abilene, Texas is targeting roughly 450,000 Nvidia GB200 GPUs and about 1.2 GW, part of a program now running toward 7 GW and more than \$400B on the way to a \$500B, 10 GW headline. Meta's Hyperion campus in Louisiana scaled to 5 GW and more than \$50B, drawing scrutiny over its water use, while its Prometheus site targets going online before the end of 2026. xAI's Colossus 2 in Memphis is heading toward roughly 1 GW, about 40% of the city's daily load, powered in part by dozens of portable gas turbines that ran without permits, the poster child for the permitting and environmental-justice risk these campuses create. The bull case, from UBS and others, is that queues and permitting physically cap overbuild, so what looks like excess behaves like a rolling upgrade; the bear case, from Man Group, is that power built for 2024–25 demand could strand by 2027–28. Both are testable against site cancellations, of which there have already been a few.

The bottleneck is the boring gear

If demand is the story's villain, the electrical supply chain is its choke point, and it is the least glamorous and most investable part of the whole atlas. A gigawatt campus needs transformers, switchgear, and high-voltage equipment, and all of it is now on allocation.

The bottleneck is the boring gear

Equipment lead-times, order to delivery — mid-2026 vs a ~2021 baseline

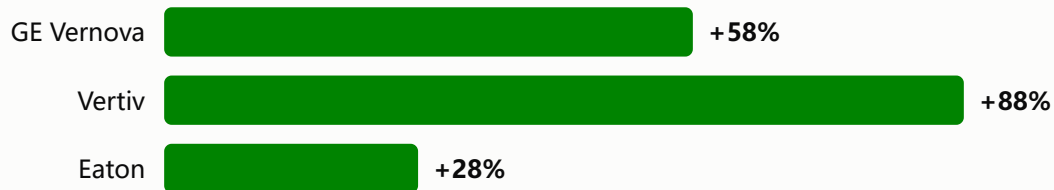


Source: PowerMag, Utility Dive (2026). GE Vernova turbine backlog 116 GW; Siemens Energy €130B+; Mitsubishi sold out to 2028.

Large power transformers run 128 to 144 weeks, with the biggest high-voltage units quoted up to four years, against roughly a year before the boom; high-voltage breakers run about 125 weeks. The order books show the demand in real time: GE Vernova's electrification backlog reached about \$76B against \$38B of prior-year sales, a book-to-bill near 2.5; Eaton's data-center orders rose about 240% year on year; and nVent lifted its 2026 organic-growth guidance to 21–23%. The market has noticed, which is the risk.

The picks-and-shovels have already re-rated

Total return, 1 January–23 July 2026 — the fundamentals are real, and so is the crowding



Source: public market total-return data through 23 July 2026, rounded; dividends reinvested where reported.

Through 23–24 July 2026, approximate total returns were 58% for GE Vernova, 88% for Vertiv, and 28% for Eaton. The fundamentals are genuine and order-backed, but the crowding is genuine too, so the live edge is now in monitoring lead-times and book-to-bill for the inflection rather than in the thesis itself. One signal to respect: in May 2026 the North American reliability regulator issued a rare Level 3 alert after more than 1,000 MW of data-center load tripped off the grid in seconds, a reminder that the physics of connecting these loads is not yet solved.

America is power-constrained; China is not

This is the layer where the two systems diverge most, and it is a durable Chinese advantage rather than a temporary one. Everything above describes the American problem: a grid that struggles to add capacity, a permitting process measured in years, and a gas-and-nuclear bridge running into its own supply limits. China has close to the opposite situation. It carries well over twice the installed generating capacity of the United States and adds more new capacity in a single year than many countries operate in total, across coal, hydro, nuclear, wind, and solar. When a Chinese province needs a new substation or transmission line for a compute cluster, the state builds it on a timescale American permitting cannot match.

China has also turned this into policy. Its "east-data, west-compute" program (东数西算) deliberately routes data centers to the western interior, where hydro, wind, and coal power are cheap and abundant, and runs the fiber back to the eastern population centers. The result is that power, the thing throttling the American build, is close to a non-issue for the Chinese build. China's binding constraints are the ones in other chapters, the advanced chips and the memory it cannot yet make at scale; power is not among them. For an investor this means the power trade is an American trade, and the Chinese expression of the AI build is in compute and localization, not in megawatts.

The megawatts that can actually arrive

The central question is whether power scarcity caps the AI build or merely reshapes it. Watch data-center site cancellations and renegotiations, the clearest real-time tell. Watch whether the roughly 750 GW of interconnection requests withdrawn in 2025 becomes a trend, which would signal that some of the demand was speculative. Watch the next PJM capacity auction for another record or a mean-reversion. Watch the transformer and gas-turbine lead-times, the single best monthly gauge of the whole cluster. And watch for regulatory intervention on cost allocation and large-load ride-through, which the reliability alert makes a question of when, not if.

Owning the bottleneck

The power and electrical complex offers order-backed exposure to a physical deployment constraint, but the securities are not substitutes for one another. Constellation (CEG), Vistra (VST), and Talen (TLN) combine generation with contract, commodity and regulatory risk. Eaton (ETN), Vertiv (VRT), GE Vernova (GEV), Quanta Services (PWR), nVent (NVT), Powell (POWL), and Siemens Energy (ENR.DE) supply electrical and grid equipment; Cummins (CMI) and gas midstream names add fuel-cycle exposure. Many already carry premium valuations, so lead times, backlog quality, book-to-bill and customer concentration matter more than the narrative alone. Pre-revenue small-modular-reactor names such as Oklo (OKLO) and NuScale (SMR) are long-dated technology and regulatory options, not solutions to the current power shortage.

The paradox of being long the problem

The trade breaks on a single event above all others: a hyperscaler cutting or pausing its 2027 capital-spending guidance, which would leave peak-priced turbine and transformer orders stranded and de-rate the entire electrical complex hard and fast. Short of that, an acceleration in interconnection-queue withdrawals would signal that the demand was partly speculative, and a genuine step-change in energy efficiency per token, of the kind a second efficiency shock in the models layer could produce, would soften the demand curve. And here is the paradox at the heart of owning this layer: faster grid and permitting reform, or a breakthrough in transformer and turbine manufacturing capacity, would relieve the scarcity that is the entire investment case. To be long the power bottleneck is to be long the problem staying unsolved.

Chapter 2.3 endnotes

[^power-iea-energy-ai]: Energy and AI (<https://www.iea.org/reports/energy-and-ai>). International Energy Agency, 10 April 2025. The 2030 value is a projection.

[^power-lbnl-us]: 2024 United States Data Center Energy Usage Report (https://eta-publications.lbl.gov/sites/default/files/2024-12/us_data_center_energy_usage_report_lbnl-2001637_0.pdf).

Lawrence Berkeley National Laboratory, December 2024. The 2028 value is a scenario range.

Data and sourcing: the master figures table (§E). Tier-1 anchors include IEA Energy and AI (2025), DOE/Berkeley Lab, PJM, EIA/FERC, and company backlogs and filings. Items marked ⚠ are explicitly treated as estimates or contested claims. Company tags and access: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [^atlas-chart-market-etn-2026] [^atlas-chart-market-prices-mid-2026] [^atlas-chart-power-equipment-leadtimes] [^atlas-chart-utility-dive-gas-turbines]

[^atlas-chart-market-etn-2026]: ETN YTD Return (<https://www.ytdreturn.com/etn/>). YTDReturn.com, undated; accessed 2026-07-25.

[^atlas-chart-market-prices-mid-2026]: VRT vs GEV Stock Comparison (<https://portfolioslab.com/tools/stock-comparison/VRT/GEV>). PortfoliosLab, undated; accessed 2026-07-25.

[^atlas-chart-power-equipment-leadtimes]: Transformers in 2026: Shortage, Scramble, or Self-Inflicted Crisis? (<https://www.powermag.com/transformers-in-2026-shortage-scramble-or-self-inflicted-crisis/>). POWER, 2026-01-01; accessed 2026-07-25.

[^atlas-chart-utility-dive-gas-turbines]: 5-year waits and rising costs: How demand is redefining the gas turbine market (<https://www.utilitydive.com/news/5-year-waits-and-rising-costs-how-demand-is-redefining-the-gas-turbine-mar/813385/>). Utility Dive, 2026-03-23; accessed 2026-07-25.

Chapter 2.4 — Cloud & the Software Moat

The cloud is where models actually run, and it is where the AI build turns into revenue. Three companies still own most of it, but the more important fact is that renting GPUs has spawned a second tier of "neoclouds" whose backlogs dwarf their revenue and whose balance sheets are stacked with GPU-backed debt. The durable moat in this layer is not the data center; it is the software, above all Nvidia's CUDA, and the question that decides the next few years is whether that moat holds at the inference layer or cracks.

Running an AI model at scale requires renting, or building, an enormous cluster of accelerators, the networking to bind them, and the power to feed them. The cloud is the business of providing that on demand. It splits into two tiers. The **hyperscalers** — Amazon Web Services, Microsoft Azure, and Google Cloud — are the diversified giants that sell everything from storage to AI. The **neoclouds** — led by CoreWeave and Nebius — are specialists that do essentially one thing, rent out Nvidia GPUs, and have grown from nothing to a meaningful share of the market in two years. Sitting above both is the layer that actually locks customers in: the software stack, and in particular Nvidia's CUDA, the programming environment that fifteen years of developer investment have made the default way to run AI.

Cloud Infrastructure Provider Landscape

The roster spans the diversified giants, the GPU-rental specialists, the inference-optimized upstarts, and the Chinese incumbents building on a separate stack.

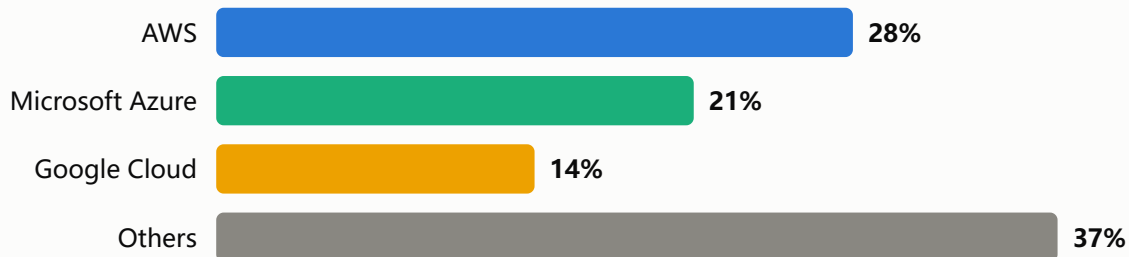
Provider	Ticker	Tier	What to know
AWS	AMZN	hyperscaler	#1, ~28% share; Trainium silicon
Microsoft Azure	MSFT	hyperscaler	AI run-rate ~\$37B, +123%
Google Cloud	GOOGL	hyperscaler	fastest-growing (+82%); TPUs
Oracle	ORCL	hyperscaler	Stargate; debt-funded
CoreWeave	CRWV	neocloud	~\$99B backlog; GPU-backed debt
Nebius	NBIS	neocloud	"sold out"; \$7–9B ARR guide
Crusoe	PVT	neocloud	~\$30B valuation; ~4.9 GW contracted
Lambda	PVT	neocloud	\$9B val; IPO planned H2 2026
Together AI	PVT	neocloud	~\$1B ARR; inference serving
Groq, Cerebras, SambaNova	PVT / CBRS	inference specialist	custom inference silicon-as-a-service
Alibaba Cloud	BABA / 9988.HK	China	~36% China share; AI run-rate ~\$5.2B
Huawei Cloud	PVT	China	Ascend/CANN "CloudMatrix" stack
Tencent Cloud	0700.HK	China	~9% share; GPU-supply-constrained

The giants, and how fast AI is growing inside them

The cloud market is large and accelerating. Total enterprise cloud infrastructure spending reached \$128.6B in the first quarter of 2026, up 35% year on year, the fastest growth in more than four years, on an annualized run-rate above half a trillion dollars. The three leaders hold roughly two-thirds of it.

The cloud giants — worldwide infrastructure share

Q1 2026; the Big Three hold ~two-thirds of a market growing 35% a year



Source: Synergy Research Group, Q1 2026 (worldwide cloud infrastructure). Google Cloud grew fastest, +82% YoY.

The headline market share understates how quickly AI is reshaping these businesses. Microsoft's AI operation reached a roughly \$37B annual run-rate, up 123% year on year, and its commercial book of committed future revenue hit \$627B, nearly doubling. Google Cloud grew 82% in the second quarter of 2026 to \$24.8B, and its backlog roughly doubled in a single quarter to about \$514B. Amazon's AWS, the largest at 28% share and \$37.6B in quarterly revenue, grew 28%, its fastest in fifteen quarters, though its operating margin slipped from 39.1% to 37.7% as the depreciation on all those AI servers began to bite. That margin detail is a preview of a theme in Chapter 2.11: the capital intensity of AI is starting to show up in the income statement.^[^atlas-chart-synergy-cloud-q1-2026]

The neoclouds, and the backlog that dwarfs the revenue

The more novel part of this layer is the neocloud, and it is best understood through one company's numbers. CoreWeave reported first-quarter 2026 revenue of about \$2.08B, up 112% year on year, which annualizes to roughly \$8B. Against that, it disclosed a revenue backlog of \$99.4B, including a \$21B commitment from Meta and multi-year deals with Anthropic and others. The backlog is more than ten times the revenue run-rate.^[^atlas-chart-coreweave-q1-2026]

The neocloud bet — backlog dwarfs revenue

CoreWeave, Q1 2026: contracted future revenue is $\sim 12\times$ the current run-rate

Revenue (annualized)  **~\$8B**

Backlog (RPO)  **\$99.4B**

Source: CoreWeave IR, Q1 2026. Funded by $\sim \$25B$ of (largely GPU-backed) debt; quarterly capex $\sim 3.7\times$ revenue.

That gap is the whole model, and the whole risk. CoreWeave spent about \$7.7B on capital in the quarter, close to four times its revenue, funded by roughly \$25B of debt, much of it collateralized by the Nvidia GPUs it buys, and it carries over 1 GW of active power against more than 3.5 GW contracted. The rest of the field shows the same shape at different scales: Nebius guides to \$7–9B of committed annual revenue on a base that was under \$2B and tells investors it is "sold out"; Crusoe is raising at around a \$30B valuation with roughly 4.9 GW contracted; Lambda, at a \$9B valuation, plans an IPO in the second half of 2026 on the back of a large Microsoft deployment; and Together AI has reached roughly \$1B of annualized revenue serving inference. The economics work if the GPUs stay highly utilized — a debt-financed cluster roughly breaks even around 70% utilization ⚠️ — and if the take-or-pay contracts hold; they break if either fails. GPU rental prices are themselves falling, with H100 rates down sharply as Blackwell supply floods in, which squeezes the margin on older fleets. Nvidia sits at the center of it all, having invested around \$2B each in CoreWeave and Nebius, whose purchases of Nvidia chips its own money helps fund, the circular arrangement Chapter 2.11 dissects.

The CUDA moat, and where it is cracking

If the data center is a commodity, the software is not. Nvidia's CUDA has roughly six million developers and first-class support in every major machine-learning framework, and the switching cost is real: years of CUDA-optimized training pipelines and custom kernels would have to be rewritten to move off it. This is the true source of Nvidia's pricing power, more durable than any single chip.

The moat is not uniform, though, and the interesting question is where it thins. For **training** at the frontier, it holds firmly. For **inference**, it is eroding. AMD's ROCm software has closed enough of the gap that AMD claims parity by the end of 2026, and its data-center revenue grew 57%. Cross-hardware abstraction layers — OpenAI's Triton and

PyTorch 2.0 among them — increasingly let a model run on AMD or a custom chip with little rewriting, which shifts the buying decision toward raw cost per token, exactly where the custom ASICs of Chapter 2.5 compete. Huawei's CANN, the Chinese alternative, claims that around 80% of standard PyTorch inference workloads run with only minor adjustments. The consensus of the evidence is that the moat is intact for training and genuinely eroding at inference, which matters because inference is now the larger share of compute.

The inference specialists

A distinct group of companies is betting that inference specifically will run best on purpose-built silicon rather than general GPUs. **Groq** (valued around \$6.9B) builds deterministic low-latency inference chips; **Cerebras** (which staged a large 2026 IPO) makes wafer-scale processors and signed a multi-hundred-megawatt OpenAI deal; and **SambaNova** (around \$11B) offers full-stack inference systems. Their pitch is inference several times cheaper per token than a general GPU ⚠️, and their relevance rises as the compute mix shifts toward inference. They are mostly private or newly public, higher-risk, and dependent on a few anchor contracts, but they are the names to recognize when the conversation turns to inference-optimized compute.

America's clouds, China's parallel

The national split repeats the pattern of the layers above. The United States has both tiers, the hyperscalers and the neoclouds, running on Nvidia and CUDA. China has a domestic cloud market roughly a tenth the size, concentrated among three players: Alibaba Cloud at about 36% share, Huawei Cloud at 16%, and Tencent Cloud at 9%. Alibaba is the standout, its external cloud revenue growing 40% with AI products now about 30% of it and an AI run-rate near \$5.2B, eleven consecutive quarters of triple-digit AI growth, which makes it the investable expression of both the Chinese cloud and the open-weight model layer.

The constraint on the Chinese cloud is the one that runs through this whole atlas: chips. Huawei Cloud's external revenue actually fell 3.5% in 2025 despite a growing market, held back by US export controls on the compute it needs, and Tencent's cloud growth is capped by GPU supply. China's answer is the same vertical stack it is building everywhere else, Huawei's Ascend chips with the CANN software, deployed through Huawei's

CloudMatrix systems that bind hundreds of Ascend chips into a single cluster. It is a genuine second ecosystem, years behind on software but improving, and it means a Chinese enterprise increasingly rents Ascend-powered compute running CANN rather than Nvidia running CUDA.

Utilization, backlog and cash conversion

Three things. The first is whether CUDA's inference moat holds or gives way, which will show up in AMD's and the custom-ASIC vendors' share of inference workloads and in whether the abstraction layers make switching genuinely painless. The second is the financial health of the neoclouds, where the signals to watch are utilization rates, any customer that fails to take contracted capacity, GPU-rental pricing, and stress in the GPU-backed debt that Chapter 2.11 treats as a systemic question. The third is margin: whether hyperscaler cloud profitability holds as AI-server depreciation mounts, the early sign of which is already visible in AWS's slipping margin.

Where to own it

Microsoft (MSFT), Alphabet (GOOGL), and Amazon (AMZN) combine cloud exposure with diversified cash generation; Oracle (ORCL) is more concentrated in large AI contracts and carries greater financing sensitivity. Nvidia (NVDA) owns a software and systems moat across cloud providers, while Broadcom (AVGO) and Arista (ANET) supply custom silicon and networking. CoreWeave (CRWV) and Nebius (NBIS) provide more direct GPU-cloud exposure with materially greater balance-sheet, customer and utilization risk. Cerebras, Groq and SambaNova are higher-risk bets on specialized inference. Alibaba (BABA, 9988.HK) is a listed expression of China's domestic cloud and open-weight model ecosystem. These securities share demand factors but have very different equity-capture and valuation profiles.

What would break the cloud thesis

The bullish case for the hyperscalers weakens if cloud margins compress faster than AI revenue grows, turning the capex into a drag rather than an engine, the early warning of which is already in AWS's margin line. The neocloud thesis breaks on a utilization or

funding shock, a single large customer failing to take capacity, or a downgrade of the GPU-backed debt, either of which would expose the leverage quickly, especially with rental prices already falling. And Nvidia's software moat, the most valuable asset in this layer, would be revealed as thinner than assumed if a major lab or hyperscaler moved a frontier training run onto non-CUDA hardware without a painful rewrite, which has not yet happened but is the thing to watch.

Data and sourcing: the cloud data appendix (research/2.4-cloud-data.md) and the master figures table. Tier-1 anchors: Synergy Research (market share), company earnings (AWS/Azure/Google Cloud, CoreWeave, Nebius, Alibaba), Canalys (China). Items marked ⚠️ (GPU rental pricing, some AI run-rates, neocloud break-even, inference cost claims) are flagged in the appendix. Company tags and access: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [^{^atlas-chart-coreweave-q1-2026}] [^{^atlas-chart-synergy-cloud-q1-2026}]

[^{^atlas-chart-coreweave-q1-2026}]: CoreWeave Reports Strong First Quarter 2026 Results (<https://investors.coreweave.com/news/news-details/2026/CoreWeave-Reports-Strong-First-Quarter-2026-Results/>). CoreWeave, 2026-05-07; accessed 2026-07-25.

[^{^atlas-chart-synergy-cloud-q1-2026}]: Cloud Market Annual Revenue Run Rate Topped Half a Trillion Dollars in Q1 (<https://www.srgresearch.com/articles/cloud-market-annual-revenue-run-rate-topped-half-a-trillion-dollars-in-q1-as-growth-surge-continues>). Synergy Research Group, undated; accessed 2026-07-25.

Chapter 2.5 — AI Chips & Accelerators

This is the layer everyone means when they say "AI chips," and it holds the central paradox of the whole build. By unit count, custom chips designed by the hyperscalers are about to overtake merchant GPUs. By revenue and profit, Nvidia still takes roughly four dollars in five and shows no sign of letting go. Both are true because units, revenue, and computing power are three different scoreboards. The durable question for an investor is not "who ships the most chips" but "who keeps the margin," and the answer runs through Nvidia's software and system integration, not the silicon alone.

An AI accelerator is a chip built to do the one kind of arithmetic that neural networks need, the multiplication of large matrices, thousands of times in parallel. A general-purpose CPU does a few things in sequence very fast; an accelerator does one thing across thousands of cores at once, which is exactly what training and running a model requires. The performance that matters is measured in floating-point operations per second (FLOPS), increasingly at low precision (the FP4 and FP8 formats that trade a little accuracy for a lot of throughput), and in how much high-bandwidth memory sits next to the compute. The most important shift in this layer is that the unit of sale is no longer the chip. Nvidia now sells the rack, an integrated "AI factory" of dozens of GPUs, CPUs, switches, and liquid cooling, and prices it accordingly.

AI Accelerator Market Structure

The roster of accelerators divides into three groups: the merchant GPUs anyone can buy, the custom ASICs the hyperscalers design for themselves, and the Chinese chips built behind the export-control wall.

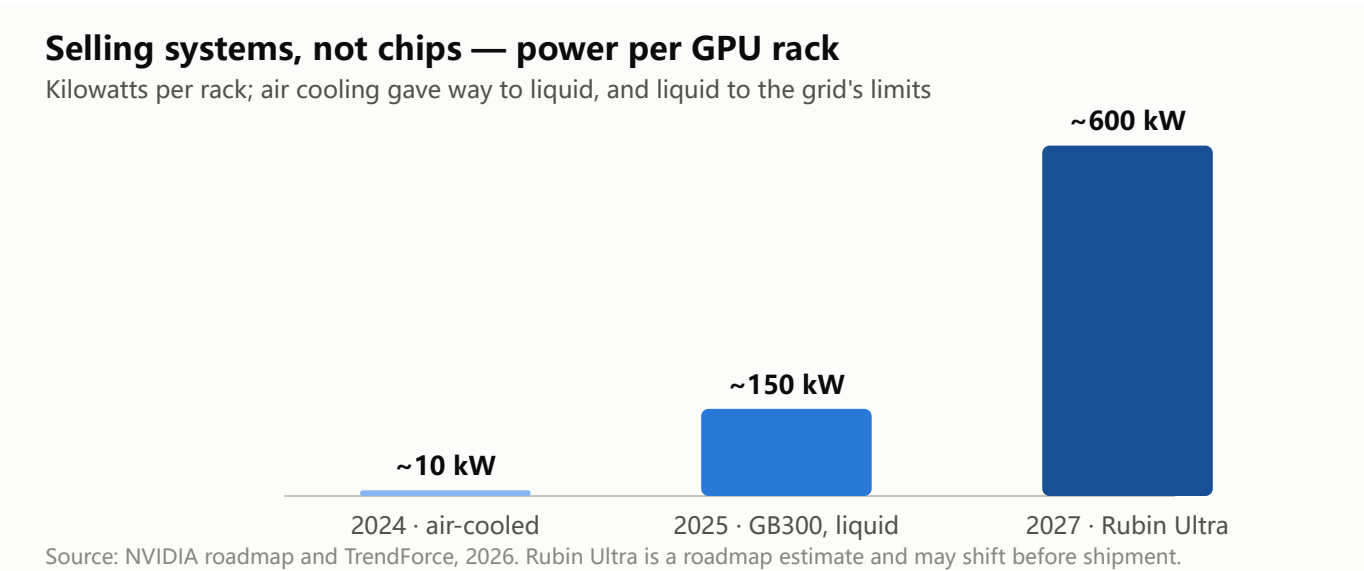
Chip / family	Maker	Type	What to know
Blackwell (GB300) → Rubin	Nvidia	merchant GPU	the standard; CUDA lock-in; ~75% of 2026 accelerator revenue (estimate)
Instinct MI300 → MI400	AMD	merchant GPU	the credible #2; OpenAI & Meta anchors
Gaudi → Jaguar Shores	Intel	merchant GPU	the also-ran; relevance pushed to 2027
TPU (Ironwood, v7)	Google	custom ASIC	the most mature hyperscaler chip
Trainium / Inferentia	AWS	custom ASIC	largest deployed ASIC fleet (500k+)
Maia	Microsoft	custom ASIC	in-house Azure silicon
MTIA	Meta	custom ASIC	in-house inference silicon
XPU (10 GW program)	OpenAI + Broadcom	custom ASIC	first volume ~2027
Ascend 910C / 950	Huawei	China champion	~600k units planned 2026; self-HBM
MLU (Siyuan)	Cambricon	China (listed)	own architecture; ~500k target
BR-series	Biren	China	training-focused
MXC / C-series	MetaX	China	domestic general-purpose GPU / accelerator program
MTT	Moore Threads	China	GPU-style
DCU	Hygon	China (listed)	x86-linked; catalogue-cleared

Chip / family	Maker	Type	What to know
(ASIC design partners)	Broadcom, Marvell	enablers	design ~95% of the hyperscaler ASICs

The names to carry forward: **Nvidia** and **AMD** are the merchant GPUs; **Broadcom** and **Marvell** are the arms dealers who design everyone else's custom chips; **Google**, **Amazon**, **Microsoft**, **Meta**, and **OpenAI** are all building their own; and **Huawei** and **Cambricon** are China's champions, trailed by **Biren**, **MetaX**, **Moore Threads**, and **Hygon**.

The one-year cadence, and why it is a moat

Nvidia has put itself on a one-year architecture rhythm that the rest of the industry struggles to match. Its Vera Rubin NVL72 rack, ramping into production in 2026, combines 72 Rubin GPUs and 36 Vera CPUs and delivers 3.6 exaFLOPS of NVFP4 inference with 20.7 TB of HBM4 across the rack. Nvidia had not published a rack-power specification at the cutoff, so third-party power estimates should not be presented as product facts. Rubin Ultra follows on the roadmap in 2027, with rack power discussed near 600 kW ⚠️. The economics of selling systems rather than chips are enormous: a GB200 or GB300 NVL72 rack runs around \$3M, a Vera Rubin rack an estimated \$5–7M, against bare GPUs at \$30–40k. The full physical BOM, evidence grades, and normalized 100 MW implications are in §2.13.[^atlas-chart-nvidia-rack-roadmap][^atlas-chart-trendforce-rack-power]



The cadence is itself the moat. Shipping a materially better rack every twelve months, co-designed across compute, networking, and cooling, raises the switching cost for any customer and strains every rival to keep pace. The risk is the mirror image: a one-year cadence stresses the entire supply chain beneath it, from the advanced packaging in Chapter 2.7 to the power in Chapter 2.3, and any slip would ripple outward. Reports of a possible Rubin Ultra "Kyber" delay circulated and were denied in 2026, the kind of rumor that will recur precisely because the cadence leaves no slack.

AMD, Intel, and the merchant challengers

The merchant-GPU challenge to Nvidia is real but narrow. **AMD** is the credible number two, with its MI400 family on TSMC's 2nm process and its double-wide Helios rack, due in the third quarter of 2026, giving it a genuine rack-scale product for the first time; AMD claims the MI455X delivers many times the token throughput of its predecessor. The demand base finally arrived with the OpenAI agreement to deploy 6 GW of AMD chips, sweetened with warrants for up to 160M AMD shares, plus a multi-gigawatt Meta commitment. AMD's data-center revenue reached \$5.8B in the first quarter of 2026, up 57%, though its share of AI accelerators remains in the mid-single digits.

Intel is the also-ran of this layer. Its Gaudi line missed even a modest \$500M target, and its go-forward story is Crescent Island, an inference GPU in customer testing in the second half of 2026, and Jaguar Shores, a 2027 rack-scale design built around silicon photonics. Intel's relevance in accelerators is now a 2027 question, and the roughly 10% US government equity stake in it (Chapter 2.8) is a bet on its foundry, not its GPUs.

The custom-ASIC surge, and the arms dealers who enable it

The most important competitive shift in this layer is that the hyperscalers are increasingly designing their own chips. Custom application-specific chips (ASICs) are growing at roughly a 45% annual rate against about 16% for merchant GPUs. Google's TPU is the most mature, now on its seventh generation ("Ironwood"); Amazon has deployed over 500,000 Trainium chips, the largest ASIC fleet by unit count; Microsoft has Maia and Meta has MTIA; and OpenAI signed a 10 GW custom-silicon program with Broadcom in October 2025, first volume expected in 2027.

Two scoreboards, two winners

Custom ASICs are overtaking GPUs by unit count — but Nvidia keeps the revenue



Source: JPMorgan research (2026 estimates / 2027 forecast). Units, revenue and computing power are different measures.

None of the hyperscalers designs these chips alone. **Broadcom** and **Marvell** are important enablers of custom silicon and can participate across several customer programs. That diversification reduces dependence on a single accelerator architecture but does not eliminate customer concentration, program timing or valuation risk. Industry forecasts that ASIC unit shipments overtake merchant GPUs should be treated as scenario inputs rather than a guarantee that design-service economics accrue equally to both suppliers.[^atlas-chart-accelerator-share]

Units, revenue, and computing power — three scoreboards

The competitive map only makes sense once you separate the scoreboards, because unit share is not revenue share and neither is compute share. By **units**, ASICs are about to lead. By **revenue and profit**, Nvidia dominates, with J.P. Morgan estimating about 75% of data-center accelerator revenue in 2026, protected by the CUDA software ecosystem of Chapter 2.4. By **computing power at the frontier**, Nvidia is more dominant still, because the custom chips mostly serve captive internal inference rather than frontier training.

This is the sharpest bull-bear debate in the layer. The most aggressive bears see Nvidia's *inference* share falling from above 90% toward 20–30% by 2028 as inference shifts to cheaper custom chips; Morgan Stanley's blunt counter is that custom silicon does not threaten Nvidia's dominance at all. The more cautious reading, and the one this atlas holds, is that ASICs take the low-margin, internal-inference slice while Nvidia keeps the high-margin merchant market, the training frontier, and the software tax, so the "crossover" is a unit-count event that need not become a revenue event. It remains the single biggest swing factor for Nvidia sentiment regardless.

America's chips, China's substitutes

The national split in this layer is the sharpest in the atlas. The United States, through Nvidia, AMD, and Broadcom, owns the frontier of AI compute and the software that runs on it. China has been cut off from that frontier and is building a parallel one under duress, with a full roster of its own. **Huawei's Ascend** line is the national champion, with something like 600,000 Ascend 910C accelerators planned for 2026 and a roadmap (the Ascend 950 series and beyond) that adds in-house high-bandwidth memory. **Cambricon** (寒武纪, 688256.SH) is the listed pure-play, tripling output toward roughly 500,000 accelerators, and behind them sit **Biren**, **MetaX**, **Moore Threads**, and **Hygon** (海光, 688041.SH), whose x86-linked DCU cleared the domestic security review. Nvidia's share of the Chinese AI-chip market has fallen from about 95% in 2022 to effectively zero, not because Chinese chips are better but because Beijing has mandated domestic silicon and Washington has restricted the good American parts.

The gap is real and worth stating precisely. Chinese accelerators trail on a per-chip basis, and Huawei's claim that an Ascend cluster beats Nvidia's is a system-level argument that uses several times as many chips and far more power to get there. Its software stack, Huawei's CANN, lags CUDA's fifteen-year head start badly. And the true ceilings on the Chinese ramp are not design but manufacturing: SMIC's yields and the domestic supply of high-bandwidth memory, both covered in later chapters. This is the layer where the bifurcation of Chapter 2.12 is most concrete, two separate compute ecosystems that no longer interoperate.

The roadmap evidence to watch

Three questions matter. The first is whether the ASIC unit-crossover becomes a revenue-and-margin crossover or stays confined to captive internal inference; watch whether Google externalizes its TPU beyond its own cloud and whether OpenAI's Broadcom silicon ships at volume in 2027. The second is the durability of Nvidia's inference share, the single biggest swing factor for its sentiment, where the bear case is contested and probably too aggressive but is the number the market fixates on. The third is the Chinese ramp, gated not by ambition but by SMIC capacity and domestic memory, where any breakout past those ceilings would reshape the Asian compute map. Underneath all three runs the cadence-and-supply-chain question: whether Nvidia can hold its one-year rhythm without a slip that would ripple through packaging, memory, and power.

The trade: own the integrator and the arms dealers

Nvidia (NVDA) remains the system integrator to beat, protected by CUDA, networking and an annual platform cadence; its risks include inference-share pressure, customer internalization and policy exposure. Broadcom (AVGO) and Marvell (MRVL) provide custom-silicon exposure across hyperscaler programs while adding networking and optics exposure; program concentration and valuation differ materially between them. AMD (AMD) is a merchant second source with real rack-scale demand but a smaller installed software base; Intel's listed thesis increasingly depends on foundry execution. On the Chinese side, Cambricon and Hygon are listed expressions of mandated domestic demand, but access, policy and localization expectations can dominate near-term fundamentals.

What would break this

The Nvidia thesis breaks if inference-share loss becomes real in the revenue line rather than the rumor mill, which a genuine, externally-adopted hyperscaler ASIC operating at training scale would signal. It also breaks, from the other direction, if a second efficiency shock in the models layer cuts the compute intensity of inference faster than volume grows, or if the one-year cadence slips and the supply chain seizes. The Chinese thesis changes if SMIC and domestic high-bandwidth memory break their current ceilings, letting Ascend and Cambricon ship at true frontier scale, which would turn a mandated, subsidized market into a competitive one and pressure the American incumbents in the third-country markets that Chapter 2.12 calls the real prize.

Data and sourcing: the chart sidecars (charts/sidecars/2.5-rackpower.json and charts/sidecars/2.5-share.json), the chart-source registry (audit/chart-sources.json), and the master figures table (§D). Tier-1 anchors: Nvidia and AMD roadmaps, J.P. Morgan and TrendForce accelerator forecasts, Broadcom and Marvell filings. Estimated rack specifications, market shares, the 20–30% inference-share bear case, and Ascend volumes remain explicitly qualified. Company access and evidence: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-chart-accelerator-share] [[^]atlas-chart-nvidia-rack-roadmap] [[^]atlas-chart-trendforce-rack-power]

[[^]atlas-chart-accelerator-share]: Eye on the Market: Semiquincententacles

(<https://am.jpmorgan.com/content/dam/jpm-am-aem/global/en/insights/eye-on-the-market/semiquincententacles-amv.pdf>). J.P. Morgan Asset Management, 2026-06-26; accessed 2026-07-25.

[[^]atlas-chart-nvidia-rack-roadmap]: NVIDIA Vera Rubin Pod: Seven Chips, Five Rack-Scale Systems, One AI Supercomputer

(<https://developer.nvidia.com/blog/nvidia-vera-rubin-pod-seven-chips-five-rack-scale-systems-one-ai-supercomputer/>). NVIDIA, undated; accessed 2026-07-25.

[[^]atlas-chart-trendforce-rack-power]: AI server rack power roadmap

(<https://www.trendforce.cn/presscenter/news/20260625-13122.html>). TrendForce, 2026-06-25; accessed 2026-07-25.

Chapter 2.6 — Interconnect & Networking

A modern AI cluster is not one computer but tens of thousands of chips that have to behave as one, and the wiring that binds them together has become both a bottleneck and a booming market in its own right. 650 Group forecast data-center AI-networking vendor revenue above \$25B in 2028; that is a narrower and more defensible measure than the former \$200B “TAM” used in this chapter. Nvidia owns the proprietary version of this plumbing; an open-standards coalition is trying to break it; and Broadcom sits in the rare position of winning almost regardless of who prevails. The next frontier is moving the optics that carry the signals onto the chip package itself.

Training a frontier model spreads the work across a cluster that can number more than a hundred thousand accelerators, and a cluster is only as fast as the slowest link between its chips. **Interconnect** is that link, and it comes in two layers. **Scale-up** is the very fast, short-range fabric that binds the dozens of GPUs inside a single rack into one giant processor; Nvidia's version is called NVLink. **Scale-out** is the broader network that connects thousands of racks across a data center; here the contest is between Nvidia's InfiniBand and ordinary Ethernet. As clusters grow, networking claims a rising share of the total bill, which is why this once-obscure layer is now a market that the biggest names in silicon are fighting over.

A >\$25B vendor-revenue market, and a proprietary-versus-open war

The prize is large and growing fast.

The market the interconnect wars are fighting over

> \$25B

in 2028 — forecast data-center
AI-networking vendor revenue.

Broadcom wins switching **and** optics; Arista wins scale-out Ethernet; Nvidia owns the integrated stack.

Source: 650 Group forecast published January 2024. Forecast, not current revenue.

The competitive structure is a familiar one: an incumbent with a proprietary, high-performance stack, and a coalition trying to pry it open. Nvidia's NVLink for scale-up and InfiniBand for scale-out are fast, mature, and locked to Nvidia's ecosystem, and they are a meaningful part of why buying Nvidia means buying a whole system. Against them, an industry coalition has built two open standards: **UALink** for scale-up and **Ultra Ethernet** for scale-out, both designed to let buyers mix hardware from different vendors. The twist is that Broadcom, rather than waiting for the UALink standard, shipped its own Ethernet-based scale-up technology first, so the near-term "open" winner has been Ethernet-flavored Broadcom silicon rather than the consortium's standard. Arista, meanwhile, is betting the scale-out network on Ethernet and calling it the eventual winner.

Interconnect Technologies and Suppliers

The layer runs from the switch silicon through the connectivity components to the optics, and Broadcom appears at almost every level.

Player	Ticker	Role
Nvidia	NVDA	NVLink (scale-up) + InfiniBand / Spectrum-X (scale-out)
Broadcom	AVGO	Tomahawk switches; Scale-Up Ethernet; co-packaged optics
Arista Networks	ANET	scale-out Ethernet switches
Marvell	MRVL	custom interconnect, DSPs
Astera Labs	ALAB	retimers / connectivity fabric
Credo	CRDO	active electrical cables
Coherent	COHR	optical components / lasers
Lumentum	LITE	lasers / optics
Fabrinet	FN	optical-module manufacturing

Scale-up versus scale-out: the two battles

The standards war is really two battles with different dynamics. In **scale-up**, Nvidia's NVLink moves roughly 1.8 terabytes a second per GPU, and the open answer, UALink, arrived late; because UALink silicon lagged, hyperscalers wanting a 2026 product leaned on Broadcom's Scale-Up Ethernet instead, delivered through its Tomahawk switches. So the near-term winner in the "open" camp has been Broadcom, not the consortium, and UALink risks being late to its own party. In **scale-out**, the contest is between Nvidia's InfiniBand and Ethernet, and here the momentum is clearer: Arista's leadership is betting publicly that Ethernet is the eventual winner, and its 1.6-terabit launch in 2026 is framed as an inflection point.

This is why **Broadcom** is the structurally advantaged name in the whole layer. It wins on the custom ASICs of Chapter 2.5, on both scale-up and scale-out Ethernet switching here, and on the co-packaged optics below, which means it collects revenue almost regardless of which hyperscaler, which chip, or which networking standard prevails. Arista owns the scale-out Ethernet opportunity, and a cluster of connectivity specialists, Astera Labs and Credo among them, supply the retimers and cables that hold these enormous fabrics together as they scale.

Optics moves onto the package

The physical limit the whole layer is running into is that pushing electrical signals fast enough over copper burns too much power and reaches only so far. The answer is optics, and 2026 is the year it moves from pluggable modules onto the switch package itself, so-called **co-packaged optics**. Nvidia's Spectrum-X Photonics switches, built on optical engines made with TSMC, and Broadcom's Tomahawk-based optical switch line are both arriving this year, and the same two companies that lead switching also lead optics. The transition is gradual, and moving optics onto the accelerator package rather than the switch is still years away, with the reliability of on-package lasers the technical crux, but it opens a durable new market for the laser and optical-component makers, Coherent, Lumentum, and Fabrinet.

Economics: content growth is not the same as supplier returns

Interconnect revenue can grow faster than accelerator units because larger clusters require more switch stages, retimers, cables and optical links. That attractive content-per-system direction still does not make every supplier equally attractive. Switch silicon and network operating systems capture architecture value; optical modules and cables generally face more manufacturing competition, qualification churn and price erosion. Co-packaged optics can raise technical barriers, but it can also transfer value from a replaceable module into the switch vendor's integrated platform.

The layer also carries unusually high customer and program concentration. A merchant supplier may win a large hyperscaler design and report several years of rapid growth, yet the economics can reverse when that customer internalizes a component, changes topology or shifts the next platform to a rival. Investors should therefore track gross margin, customer concentration and content per deployed accelerator alongside headline networking revenue. The \$25B forecast establishes the market opportunity, not the margin pool available to every participant.[^atlas-chart-ai-networking-tam]

America's network, China's mesh

The interconnect layer is another American stronghold, owned by Nvidia, Broadcom, and Arista. China's answer is instructive and ties back to Chapter 2.5. Because its individual

Ascend chips trail Nvidia's, Huawei compensates at the level of the network, building enormous clusters: its Atlas SuperPods bind thousands of Ascend chips, over eight thousand in one system, into a single machine using Huawei's own interconnect fabric. The strategy is to reach frontier-scale computing power through system-level density and networking rather than through per-chip superiority, which makes the interconnect the crucial enabler of China's whole compute effort. It is also years behind on the software and optics that make the Western fabrics efficient.

When open fabrics become deployed systems

Watch whether open Ethernet genuinely displaces Nvidia's proprietary stack or whether Nvidia's vertical integration keeps it locked at the frontier, and in particular whether the UALink standard becomes real in deployments or whether Broadcom's head-start Ethernet keeps winning the sockets. Watch Ultra Ethernet's share against InfiniBand in newly-built training clusters, the clearest measure of the open coalition's progress. Watch the volume ramp of co-packaged optics in the second half of 2026, and the reliability of the on-package lasers. And watch whether Nvidia opens NVLink, a move that would concede the standards war while trying to keep the customers.

Who captures the networking value

Broadcom (AVGO) has the broadest public-market exposure in this layer because it captures value across custom silicon, switching, and optics under several standards outcomes. Arista (ANET) is a more concentrated expression of the shift toward Ethernet in scale-out networking. Nvidia (NVDA) still owns the integrated, proprietary stack and benefits as long as buying its systems means buying its network. The connectivity specialists, Astera Labs (ALAB) and Credo (CRDO), offer higher-beta exposure to cluster growth, while Coherent (COHR), Lumentum (LITE), and Fabrinet (FN) are exposed to the co-packaged-optics transition. These are different security profiles despite sharing the same demand driver; valuation and customer concentration determine whether strategic growth reaches shareholders.

What would loosen Nvidia's grip

The bullish case for the open-networking names weakens if Nvidia's proprietary stack proves stickier than expected and NVLink and InfiniBand hold their ground at the frontier, keeping the value inside Nvidia's system rather than letting it flow to merchant switching and Ethernet. Broadcom's position is the most robust in the layer, so the main risk to it is a broad AI-capex slowdown rather than a competitive loss. And the co-packaged-optics thesis slips if the reliability of on-package optics disappoints and the industry stays on pluggable modules longer than the 2026 timeline implies, delaying the opportunity for the optical-component suppliers.

For the platform-specific translation from switch and optical content into complete U.S. and Chinese systems, see the four supply-chain teardowns in §2.13.

Data and sourcing: the quantitative sidecar (charts/sidecars/2.6-tam.json), the chart-source registry (audit/chart-sources.json), and the master figures table (§D). Tier-1 anchors: Broadcom filings, Arista disclosures, 650 Group (data-center AI-networking vendor revenue), and specialist engineering research on scale-up Ethernet. Company access and evidence: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [^{^atlas-chart-ai-networking-tam}]

[^{^atlas-chart-ai-networking-tam}]: Data Center AI Networking to Surge to Over \$25B in 2028 (<https://650group.com/press-releases/data-center-ai-networking-to-surge-to-over-25b-in-2028-according-to-650-group/>). 650 Group, 2024-01-24; accessed 2026-07-25.

Chapter 2.7 — Memory & Packaging

Two of the least visible layers in the stack have quietly become the most valuable and the most constrained. High-bandwidth memory now accounts for more of a top GPU's cost than the logic chip itself, and it is sold out. Advanced packaging, the step that bonds the memory and the compute together, is the single gating bottleneck on how many AI chips the world can build. As the shrinking of transistors delivers less each generation, the value of an accelerator is migrating off the logic die and into the memory and packaging around it. Own the bottleneck, not just the compute.

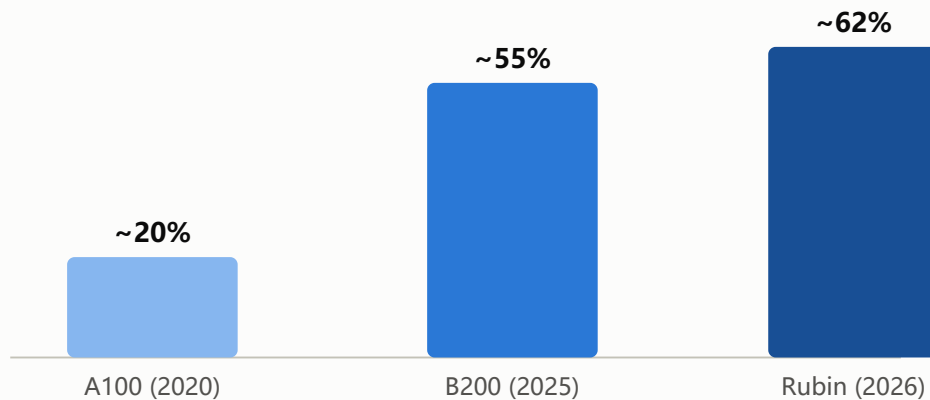
An AI accelerator is not one chip but an assembly. At its center is the logic die that does the computing, and stacked right beside it is **high-bandwidth memory (HBM)**, several dozen layers of DRAM bonded into towers that feed the processor. The reason this matters is that modern models are limited less by how fast a chip can calculate than by how fast it can be fed data, the "memory wall." More and faster memory, placed as close as physically possible to the compute, is what breaks that wall. **Advanced packaging** is the manufacturing art of putting all of it, logic and memory, onto one tightly-integrated module; the dominant method, TSMC's CoWoS, is what physically makes a modern GPU possible. Neither layer designs the chip, but without them there is no chip.

The value has moved into the memory

The clearest way to see this layer's importance is in the bill of materials. On Nvidia's A100 a few years ago, HBM was roughly a fifth of the component cost. On the current Blackwell generation it is more than half, and on the coming Rubin generation it is tracking toward nearly two-thirds.[^atlas-chart-hbm-bom-estimate]

The value moved into the memory

High-bandwidth memory as a share of a GPU's bill of materials, by generation



Source: Silicon Analysts estimate (▲). HBM now costs more than the logic die; its base layer is fabbed at TSMC on a logic process.

That shift is the whole investment thesis for this layer. As Chapter 2.8 explains, shrinking transistors now costs more and delivers less, so the extra performance each GPU generation needs comes increasingly from more memory and better packaging rather than from the logic die alone. Memory is no longer a commodity component bolted on at the end; it is becoming the most expensive and most differentiated part of the chip, and with HBM4 its base layer is being built on a logic process at TSMC, tying the memory makers and the foundry together and opening the door to semi-custom HBM designed for a specific customer's chip.

Memory and Advanced Packaging Ecosystem

The layer runs from the three memory makers through the packaging houses to a set of Japanese materials suppliers that most investors have never heard of but that gate the whole thing.

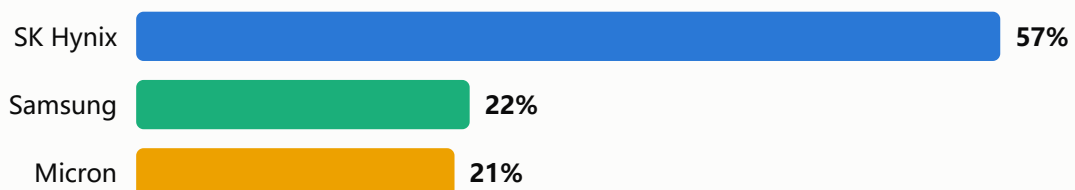
Player	Ticker	Role
SK Hynix	000660.KS	HBM leader (~57% bit share)
Micron	MU	HBM #3; the accessible US pure-play
Samsung	005930.KS	HBM #2 (turnaround) + foundry
CXMT	(A-share, new)	China's DRAM champion; HBM laggard
TSMC	TSM	CoWoS packaging monopoly
Amkor, ASE/SPIL	AMKR, ASX	OSAT packaging overflow
Besi	BESI	hybrid bonding (next-gen)
Ibiden, Shinko	(Japan)	high-end substrate (>70%)
Ajinomoto	(Japan)	ABF insulating film (>95%)
Resonac, Namics	(Japan)	molding compounds, underfills
SK Absolics, Corning	SKC, GLW	glass substrates (the coming shift)

The three-way memory race

The HBM market is a tight oligopoly of three. SK Hynix is the clear leader, with about 57% of HBM bit shipments, having won the Nvidia relationship early and held it. Samsung is fighting to recover lost ground at around 22%, and Micron, the sole American maker, has climbed to about 21%.^[^atlas-chart-counterpoint-hbm-q3-2025]

The HBM oligopoly — three makers, one leader

Share of HBM bit shipments, Q3 2025; HBM4 (2026) seen ~50 / 30 / 20



Source: Counterpoint Research, Q3 2025. HBM is sold out — pricing power sits with the sellers. Micron is the US pure-play.

The coming HBM4 generation is the 2026 event, and it changes the competitive stakes. All three vendors qualified for Nvidia's Rubin platform, with the split expected around 50% SK Hynix, 30% Samsung, and 20% Micron, and because HBM is sold out the pricing power sits with the sellers. HBM4 sells for around \$500 per stack, more than \$10 per gigabyte against \$7–8 for the prior generation. The deeper change is that HBM4's base die moves to a logic process fabricated at TSMC, which both couples the memory makers to the foundry and turns HBM from a commodity into a semi-custom product, potentially eroding the brutal price cyclicality that has always defined memory. The two open questions are whether Samsung's yields recover enough to hold its 30% slot after missing the prior generation's qualification, and whether TSMC's role in the base die hands the foundry leverage over the memory makers. Micron is the most accessible pure-play for a US investor; SK Hynix and Samsung trade in Korea.

The packaging monopoly and its hidden chokepoints

Packaging is even more concentrated than memory, and it hides some of the most extreme single-supplier dependencies in the entire supply chain. TSMC's CoWoS is the gating resource for the whole AI-chip industry, and its capacity, around 70,000–80,000 wafers a month at the end of 2025, is being pushed toward 120,000–130,000 by the end of 2026, with the overflow subcontracted to Amkor and Taiwan's SPIL. Nvidia alone is expected to consume roughly 700,000 CoWoS wafers in 2026. Beneath the package sits the substrate, and here a Japanese duopoly of Ibiden and Shinko holds more than 70% of the high-end, with Ibiden the dominant supplier for Nvidia's top GPUs and spending ¥500B to expand.

The most extreme chokepoint of all is nearly invisible: **Ajinomoto**, a Japanese company better known for seasoning, makes more than 95% of the insulating film (ABF) that every high-end chip package is built on, and it raised prices about 30% for the second half of 2026. Around it sit other Japanese materials names that gate the layer: Resonac supplies the epoxy molding compounds and Namics the underfills that hold these packages together. The next shift to watch is **glass substrates**, which promise large speed and power gains over today's organic substrates: SK Absolics (a unit of SKC) is preparing mass production in Georgia with samples to AMD, Intel has begun licensing its glass technology, Samsung targets glass interposers by 2028, and Corning and Japan's AGC supply the glass itself. These are the sort of single-supplier dependencies that do not show up in headlines until they break.

America's memory, Asia's packaging, China's gap

The national map of this layer is distinctive. Memory is a Korean and American strength: SK Hynix and Samsung in Korea, Micron in the US, with Japan and Taiwan supplying the materials and the packaging. China is conspicuously behind. Its DRAM champion, CXMT, is scaling fast and just completed the largest chip IPO of 2026, but it trails badly in HBM specifically, and Huawei's push to build its own HBM for the Ascend accelerators of Chapter 2.5 is an attempt to close a gap that is currently one of the two hardest constraints on China's entire AI-chip effort. The other, leading-edge fabrication, is Chapter 2.8.

This is the layer, more than almost any other, where export controls bite China hardest. High-bandwidth memory has been a specific target of restrictions, because it is the component China can least easily substitute, and its scarcity is why Chinese accelerators lean on stockpiled or domestically-produced memory that lags the frontier. Packaging is similar: the CoWoS capacity that makes Western AI chips is in Taiwan, and China is building domestic advanced-packaging capacity but from well behind. For an investor the read is that memory and packaging are where the American and allied position is strongest and the Chinese position weakest, the mirror image of the materials layer in Chapter 2.10.

Qualification, capacity and the memory cycle

Three things. The first is the HBM4 ramp and whether it stays sold out, which would preserve the pricing power that makes memory such an attractive part of the chain right now; the risk is that memory is historically cyclical, and every past boom has ended in oversupply and a price crash. The second is whether Samsung's yields recover enough to hold its 30% HBM4 share or whether SK Hynix extends its lead. The third is packaging capacity: whether TSMC's CoWoS expansion actually relieves the bottleneck or whether each new GPU generation, consuming ever more package area, keeps demand ahead of supply. And a slower structural question runs underneath: whether glass substrates begin to displace organic ones and reshuffle the substrate hierarchy later in the decade.

Owning the rising share of the chip

Micron (MU) is the most accessible US-listed pure-play in the memory group, while SK Hynix (000660.KS) is the HBM leader and Samsung (005930.KS) a turnaround case for investors with Korean-market access. All remain exposed to memory-cycle supply response. The packaging constraint accrues partly to TSMC (TSM), with Amkor (AMKR) and ASE (ASX) as outsourced-assembly alternatives and Besi (BESI) exposed to hybrid bonding. Ibiden and Ajinomoto are difficult-to-replace substrate inputs but may be less accessible; SK Absolics and Corning (GLW) are longer-dated glass-substrate options. Memory and packaging capture a rising share of system value, but capacity additions, qualification and entry valuation determine whether that trend produces excess returns.

The cycle that always turns

The biggest risk to this layer is its own history: memory is cyclical, and the current sold-out, high-price condition has always, eventually, given way to oversupply and a brutal price decline. A wave of new HBM capacity arriving faster than demand, or any stumble in AI-chip build rates, would hit the memory makers hard and fast, and the same efficiency shock that threatens the compute layer would flow through here too. On the packaging side, the thesis weakens if CoWoS capacity finally overshoots demand, turning a scarce, high-margin service into a commodity. And a genuine Chinese breakthrough in domestic HBM, which is not close today, would over time erode the allied advantage that makes this layer such a clean expression of the Western position in AI.

For the package-to-system mapping of HBM3E, HBM4, host memory and disclosed supplier relationships, see §2.13.

Data and sourcing: the quantitative sidecars (charts/sidecars/2.7-hbmbom.json and charts/sidecars/2.7-hbmshare.json), the materials data appendix (research/2.10-materials-data.md), the chart-source registry (audit/chart-sources.json), and the master figures table (§D). Tier-1 anchors: Counterpoint, company disclosures, and specialist packaging research. HBM bill-of-materials shares, HBM4 allocation, and market-size figures remain explicitly labeled as estimates. Company access and evidence: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-chart-counterpoint-hbm-q3-2025] [[^]atlas-chart-hbm-bom-estimate]

[[^]atlas-chart-counterpoint-hbm-q3-2025]: Global DRAM and HBM Market Share: Quarterly (<https://korea.counterpointresearch.com/global-dram-and-hbm-market-share-quarterly/>). Counterpoint Research, undated; accessed 2026-07-25.

[[^]atlas-chart-hbm-bom-estimate]: The Memory Triopoly (<https://www.wing.vc/content/the-memory-triopoly>). Wing Venture Capital, undated; accessed 2026-07-25.

Chapter 2.8 — Foundry & Fabrication

Every chip in this atlas is designed by one company and built by another, and the building is harder than the designing. Turning a blueprint into working silicon at the leading edge is the most difficult manufacturing humanity does, and one company on one island does the overwhelming majority of it. TSMC makes roughly nine in ten of the world's most advanced logic chips, its 2nm capacity is sold out, and its dominance is simultaneously the foundation of the AI build and its single largest point of failure. This is the chapter where the geopolitics of Chapter 2.12 become physical.

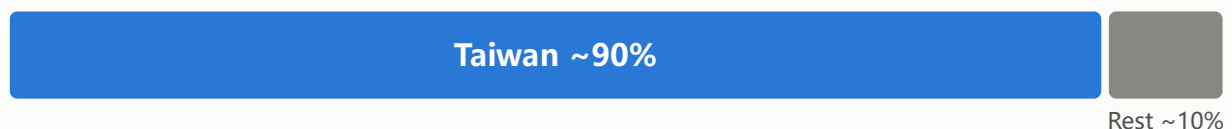
A **foundry** is a contract chip factory. Nvidia, AMD, Apple, and the hyperscalers design their chips but own no fabs; they hand the designs to a foundry, overwhelmingly TSMC, which manufactures them. The difficulty is almost impossible to overstate. A leading-edge fab costs upward of \$20B, runs for years to reach acceptable **yield** (the fraction of chips on a wafer that actually work), and etches features a few nanometers wide using the extreme-ultraviolet lithography of Chapter 2.9. Progress is measured in **process nodes**, the "2nm" and "3nm" generations, each denser than the last. The catch, and a theme of this chapter, is that each new node now costs more and delivers less than the one before, which is why value is migrating to the memory and packaging of Chapter 2.7.

One island makes almost all of it

The concentration at the leading edge is extraordinary and worth seeing plainly.

Where the world's most advanced logic is made

Estimated share of leading-edge (sub-5nm) logic capacity, 2026

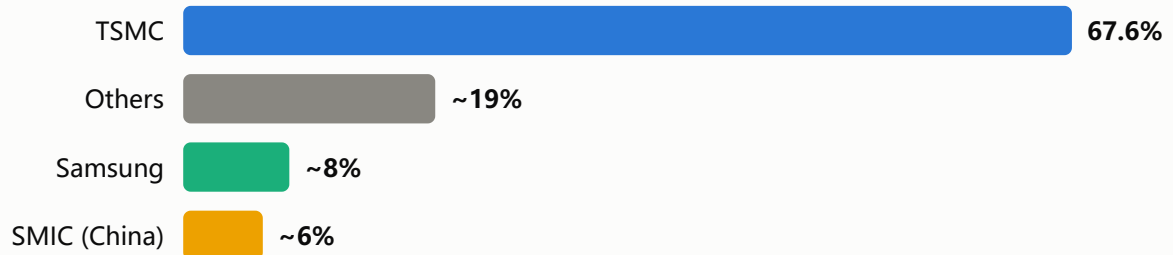


Source: Infosys semiconductor outlook (industry estimate). Capacity share, not corporate revenue or wafer starts.

TSMC generated about \$40.2B in revenue in the second quarter of 2026, its 2nm node entered volume production and is sold out through the year, and it plans a roughly 70% compound growth in advanced-node capacity through 2028. Its share of the pure-play foundry market dwarfs everyone else's.[^atlas-chart-taiwan-leading-edge-share][^atlas-chart-trendforce-foundry-share]

TSMC's foundry dominance

Share of the pure-play foundry market by revenue, Q1 2025



Source: TrendForce, Q1 2025. Rounded display values except TSMC.

Foundry Capacity and Process Leadership

Only a handful of companies can make a modern chip at all, and only three are in the 2nm-class race.

Foundry	Ticker	Leading node	Status
TSMC	TSM	N2 (2nm) → A16	~90% of leading-edge; sold out
Samsung Foundry	005930.KS	SF2 (2nm)	credible #2; Tesla AI6 win
Intel Foundry	INTC	18A → 14A	tech leads (PowerVia, High-NA); few external customers
SMIC	688981.SH / 0981.HK	7nm (via DUV)	China champion; ~20–30% yield
GlobalFoundries	GFS	mature nodes	US-based, not leading-edge
UMC	UMC	mature nodes	trailing-edge specialist

The node race: 2nm, backside power, and the High-NA bet

The leading edge is a three-way contest with diverging strategies. **TSMC's** N2 node entered volume production in late 2025 and ramps through 2026, its first gate-all-around transistor, with backside power delivery arriving in the A16 node in the second half of 2026; TSMC is deliberately *deferring* the ultra-expensive High-NA lithography tool, betting that mature machines keep its cost per transistor lower for longer. **Intel** took the opposite bet: its 18A node reached volume in mid-2025 *with* backside power (PowerVia), beating TSMC to it, and Intel was the first to put a production High-NA machine to work, genuine technology leads. But Intel Foundry still lacks a TSMC-scale roster of external customers, which is the one thing that would make its comeback real, and its 14A node will be more expensive precisely because of High-NA. **Samsung's** SF2 process reached above 60% yield and landed a \$16.5B Tesla contract as validation, though the AI6 chip reportedly slipped as 2nm production lagged.

The core disagreement, the timing of High-NA lithography, decides the economics of the leading edge for years: Intel is adopting it early to build an advantage, TSMC is deferring

it on cost, and ASML (Chapter 2.9) sells the machine to whoever buys. Underneath the contest runs the theme of the whole atlas's hardware layers: as node shrinks deliver less and cost more, value migrates into the packaging and memory of Chapter 2.7, which is why "who wins 2nm" matters somewhat less than it used to.

The Taiwan question, and China's ceiling

This layer is where the two-systems story becomes most consequential. The United States leads chip *design* and owns none of the leading-edge *manufacturing* on its own soil; that sits in Taiwan, ninety miles from China. TSMC is diversifying, with an Arizona program now committed to around \$265B and 2nm fabs breaking ground there, plus sites in Japan and Germany. But the diversification is slow and back-loaded: Arizona runs a generation behind Taiwan, the advanced packaging and the research and development stay in Taiwan, and roughly a third of TSMC's most-advanced capacity is only *targeted* for Arizona years out. The uncomfortable conclusion, developed in Chapter 2.12, is that a disruption in the Taiwan Strait would halt the entire AI build at once, and no amount of Arizona construction removes that risk before roughly 2030.

China sits on the other side of the same wall. Its national foundry, SMIC, produces 7nm-class chips using older deep-ultraviolet tools and multi-patterning, the workaround for being denied EUV, but at yields estimated in the 20–30% range that make the economics dependent on state support. SMIC is the manufacturing ceiling on China's entire AI-chip ambition: the Ascend and Cambricon accelerators of Chapter 2.5 can only ship in the volumes SMIC can yield, and pushing below 7nm on purely domestic equipment is judged unlikely before the end of the decade, though a SiCarrier-affiliated domestic immersion tool is in testing at SMIC toward a 28nm domestic flow in 2027. Fabrication and high-bandwidth memory are the two hard limits on Chinese AI, and this is one of them.

Foundry economics: utilization, yield and customer mix

A foundry's moat is expressed through economics as much as transistor density. New fabs absorb capital years before they generate revenue; low initial yields consume wafers without producing saleable dies; and fixed depreciation makes utilization a powerful driver of margins. A node can be technically competitive yet economically weak if it lacks enough anchor customers to fill the fab and spread process-development cost.

Conversely, a dense portfolio of customers supplies learning cycles, purchasing scale and cash to finance the next node.

That feedback loop explains the difference between process announcements and a durable foundry franchise. For TSMC, the key questions are pricing, leading-edge utilization, customer concentration and returns on overseas capacity. For Intel and Samsung, the decisive evidence is not a demonstration wafer but repeat external volume from customers that are not affiliated with the foundry owner. Capacity announcements should therefore be read with a lagged-normalization risk: synchronized subsidies can create excess capacity in mature nodes even while the leading edge remains scarce.

Yield, external customers and geographic diversification

The central technical contest is the timing of High-NA lithography, and ASML's order flow is a clean read on who is betting what. Watch TSMC's 2nm and A16 yields and its Arizona timeline; watch whether Intel Foundry lands a marquee external customer at the leading edge, which would introduce real competition for the first time in years; watch Samsung's yield stability and whether its Tesla win holds; and watch SMIC's yields as the real gauge of how fast China can scale domestic AI chips. Above all, the value question runs underneath all of it: as node shrinks deliver less, more of each chip's worth moves into the packaging and memory of Chapter 2.7.

Strategic necessity, geopolitical discount

TSMC (TSM) is presently difficult to substitute at the leading edge and sits upstream of Nvidia, Apple, AMD, and Broadcom. That strategic position supports pricing power, but security attractiveness still depends on valuation and the uncompensated Taiwan tail risk. Intel (INTC) is a high-variance US alternative, with genuine process capability and a government backstop but unproven external foundry demand; Samsung (005930.KS) is the credible second source; GlobalFoundries (GFS) is a US-based mature-node exposure. The equipment makers in Chapter 2.9 diversify foundry-customer risk but remain cyclical. On the Chinese side, SMIC (688981.SH, 0981.HK) is the mandated national champion, while low leading-edge yield makes it primarily a policy and localization thesis rather than a conventional margin thesis.

The one event that breaks everything

The dominant risk here is not competitive but geopolitical: a Taiwan Strait crisis is the single event that would invalidate the entire AI capex thesis in this atlas at once, and it can be sized but not hedged away. Short of that, the TSMC thesis would weaken if Intel Foundry genuinely won external leading-edge customers at scale, introducing real competition for the first time in years, or if Arizona and the other overseas fabs de-risked Taiwan faster than expected. The China thesis changes if SMIC breaks its yield ceiling or acquires a domestic path below 7nm, which would loosen the manufacturing limit on Chinese AI chips and pressure the allied advantage that this layer, like memory, currently represents.

For the platform-level distinction between confirmed fabrication, reported relationships and undisclosed next-generation allocation, see §2.13.

Data and sourcing: the quantitative sidecars (charts/sidecars/2.8-taiwan.json and charts/sidecars/2.8-foundryshare.json), the equipment & EDA appendix (research/2.9-equipment-eda-data.md), the chart-source registry (audit/chart-sources.json), and the master figures table (§D, §F). Tier-1 anchors: TSMC and Intel filings and TrendForce. The approximately 90% concentration of sub-5nm production in Taiwan remains an industry estimate. Company access and evidence: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-chart-taiwan-leading-edge-share] [[^]atlas-chart-trendforce-foundry-share]

[[^]atlas-chart-taiwan-leading-edge-share]: Semiconductor Industry Outlook 2026 (<https://www.infosys.com/iki/documents/semiconductor-industry-outlook2026.pdf>). Infosys Knowledge Institute, undated; accessed 2026-07-25.

[[^]atlas-chart-trendforce-foundry-share]: Global top ten foundries' revenue and market share in 1Q25 (<https://www.trendforce.cn/presscenter/news/20250609-12611.html>). TrendForce, 2025-06-09; accessed 2026-07-25.

Chapter 2.9 — Equipment & EDA

Beneath the foundries sit the companies that sell them their machines and their design software, and these are among the best businesses in the entire chain, because they get paid by every chipmaker regardless of which one wins. Five equipment makers dominate the tools, one of them, ASML, holds an outright monopoly on the single machine without which no advanced chip can be made, and two companies control the software used to design almost every chip on earth. This is also the layer where export controls are wielded most precisely, because denying a country these tools freezes its progress in place.

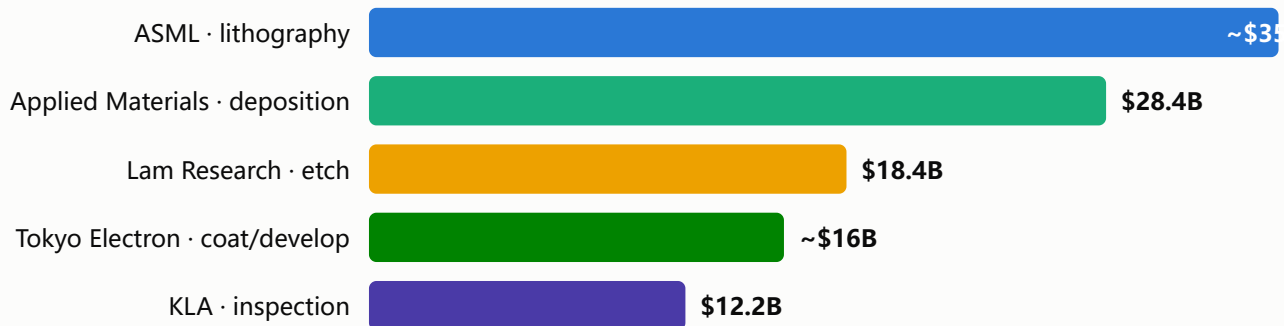
To build a chip you need two things the foundry itself does not make: the physical machines that etch, coat, and inspect the silicon, known collectively as **wafer-fab equipment (WFE)**, and the software used to design the chip in the first place, **electronic design automation (EDA)**. Both are quiet, unglamorous, and extraordinarily profitable, because they are picks-and-shovels sold to everyone in the gold rush. The WFE market runs around \$116B a year and is growing on the back of AI-driven memory and logic demand. It is also the most concentrated set of chokepoints in the supply chain, which is precisely why it has become the sharpest instrument of technology policy.

Five machines, and one monopoly

The equipment market divides cleanly, with each of the five giants owning a step of the process.

The five machines that make every chip

Wafer-fab-equipment makers, revenue by latest fiscal year (\$B); each owns one process step



Source: company filings, latest FY (SEMI for market). ASML has a 100% monopoly on EUV lithography — no EUV, no advanced chip.

The most important of them is **ASML**. It is the sole maker of extreme-ultraviolet (EUV) lithography machines, the systems that print the finest features on a leading-edge chip, and its monopoly is total: no EUV, no advanced logic or memory. Its next-generation High-NA machine costs roughly \$350–400M each, and the first units have gone to Intel, SK Hynix, and Samsung. Applied Materials and Lam Research dominate the deposition and etch steps, Tokyo Electron holds about 91% of the coater-developer niche, and KLA controls process-control and inspection with more than 85% of the optical-inspection market. Because each owns its step, they do not so much compete with one another as ride the same wave of fab construction together.

Semiconductor Equipment and EDA Landscape

The roster spans lithography, the process-step leaders, the design-software duopoly, the advanced-packaging machine makers, and the Chinese localization champions.

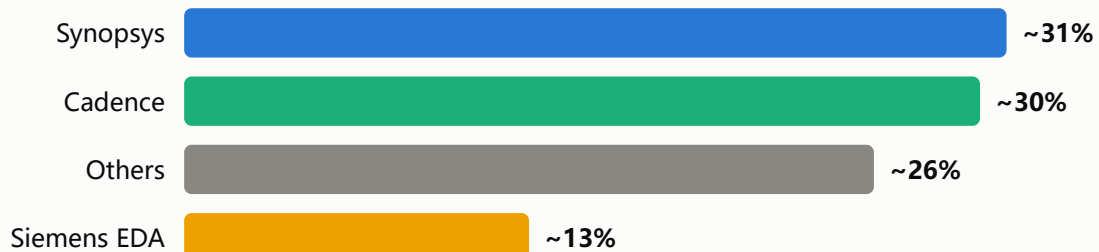
Player	Ticker	Role
ASML	ASML	lithography — 100% of EUV
Applied Materials	AMAT	deposition, ion implant, CMP
Lam Research	LRCX	etch + memory deposition
Tokyo Electron	8035.T	coater/developer, etch, clean
KLA	KLAC	metrology & inspection
Synopsys, Cadence	SNPS, CDNS	EDA design software (the duopoly)
Besi, ASMPT, Kulicke & Soffa	BESI, 0522.HK, KLIC	advanced-packaging / bonding equipment
Naura, AMEC, Hwatsing	002371.SZ, 688012.SH, 688120.SH	China WFE (etch, deposition, CMP)
SMEE, SiCarrier	(China, private/listed)	China lithography (DUV; the gap)
Empyrean, Primarius	301269.SZ, 688206.SH	China EDA

The EDA duopoly

The design-software layer is a story of concentration as extreme as the equipment.

The EDA duopoly — the software behind every chip

Approx. share of the chip-design-software market, 2025



Source: SemiAnalysis / company filings (approx., ⚠️). Synopsys + Cadence >60%; switching is near-unthinkable. China's Empyrean ~

Synopsys and Cadence, the two American EDA leaders, together hold well over half the market, with Siemens EDA a distant third, and their software is so embedded in how chips are designed that switching is nearly unthinkable; both are increasingly weaving AI

into the design process itself, with tools like Synopsys's DSO.ai and Cadence's Cerebrus. It is a duopoly with the pricing power that implies, and, like the lithography layer, a weapon: when the US briefly required export licenses for chip-design software to China in 2025, Synopsys suspended its guidance and halted China sales overnight, before the rule was rescinded weeks later in a trade truce. China's own EDA champions, Empyrean and Primarius, hold only about a tenth of their home market and lack a full-flow tool suite for the most advanced nodes.

The tools are the sharpest weapon against China

This layer is where the US and its allies hold their most decisive leverage over China, and where they have used it. China has pushed its equipment self-sufficiency to roughly 35%, and its national champions have grown fast: Naura has climbed to about the fifth-largest equipment maker in the world, its revenue heading toward \$7B, and AMEC's etch tools approach the cutting edge. But the gap that matters is lithography. China has no domestic EUV; its most advanced homegrown machine, from SMEE, does older deep-ultraviolet work, and while a SiCarrier-affiliated immersion-DUV tool is being tested at SMIC toward a 28nm domestic flow in 2027, integration into a production line at scale is years out. Denying China EUV is what caps SMIC's fabrication ceiling in Chapter 2.8, and denying it the servicing and spares for the tools it already owns is a newer, quieter escalation that the US has pressed the Netherlands and Japan to join. The exposure runs both ways, though: China is a large customer, historically around a fifth of Cadence's revenue and a meaningful slice of ASML's (guided down toward 20% for 2026), so the same controls that constrain China also cost the toolmakers.

The packaging-equipment boom

A newer growth vector inside this layer is the equipment for advanced packaging, driven by the CoWoS bottleneck of Chapter 2.7. As TSMC quadruples its packaging capacity, demand flows to the makers of the bonding and assembly machines: Besi, whose hybrid-bonding tools are central to the next generation and which runs a joint venture with Applied Materials; ASMPT and Kulicke & Soffa; and Applied Materials itself. High-bandwidth memory intensifies this, because a bit of HBM consumes roughly three times the wafer-fab capacity of a bit of ordinary memory, so the AI build drives

disproportionate equipment demand at exactly the packaging and memory steps where value is concentrating.

The installed base, qualification cycle and normalization risk

Equipment economics extend beyond the initial tool sale. A qualified process tool generates spare-parts, service, software and upgrade demand over a long installed life, making the installed base a stabilizer when new-fab spending slows. Process control and lithography can be especially sticky because changing a tool alters the process recipe and may force costly requalification. Technical service intensity, installed-base growth and recurring revenue therefore deserve separate attention from annual wafer-fab-equipment shipments.

The same stickiness does not remove cyclicity. Customers can pull equipment forward ahead of an export-control deadline, double-order during a shortage or pause once a fab shell has been equipped. A strong shipment quarter can therefore borrow from the future. The underwriting sequence should be: end-market wafer demand, customer utilization, fab construction, tool orders, revenue recognition and finally service attachment. SEMI's equipment forecast is an industry spending measure, not a promise of identical growth or margins for every vendor.[^atlas-chart-semi-wfe-2025]

Customer concentration and geographic mix are equally important. China restrictions can remove addressable products while localization spending temporarily raises demand for permitted tools. Memory suppliers' synchronized expansions can create a sharp up-cycle followed by digestion. The most defensible suppliers are those whose process step becomes more intensive at each node, whose installed bases compound, and whose tools remain difficult to replace without yield loss.

Orders that distinguish expansion from pull-forward

The defining technical question is the timing of High-NA lithography, the same debate as in Chapter 2.8, and ASML's order flow is a clean read on who is betting what. Watch China's self-sufficiency percentage and whether it can finally close the lithography gap; watch the export-control and servicing rules for further escalation or relaxation; watch the packaging-equipment makers as CoWoS capacity keeps expanding; and watch the

WFE cycle itself, because this layer is cyclical and a downturn would hit all five giants together regardless of the AI narrative.

Tools, software and the cycle

ASML (ASML) is the supply chain's most concentrated lithography chokepoint, with China exposure and the timing of leading-edge capacity as major swing factors. Applied Materials (AMAT), Lam Research (LRCX), KLA (KLAC), and Tokyo Electron (8035.T) diversify across process steps and customers, but remain exposed to wafer-fab-equipment normalization. Synopsys (SNPS) and Cadence (CDNS) form the EDA duopoly and carry software-like recurrence alongside semiconductor-cycle exposure. Besi (BESI) is a more concentrated packaging-equipment expression. Naura (002371.SZ) and AMEC (688012.SH) are Chinese localization beneficiaries, with Entity List exposure and a persistent domestic-lithography ceiling. None of these strategic positions eliminates entry-valuation risk.

The cycle, and the China swing

The equipment thesis is cyclical at its core, so the main risk is a WFE downturn, which would affect all five giants despite the AI narrative. ASML's specific exposure is China: tighter export controls could remove addressable revenue, while relaxation would add to it. A credible Chinese breakthrough in domestic lithography would gradually weaken a central allied point of leverage over advanced Chinese chipmaking. The EDA duopoly would also face a new question if AI-driven design tools lowered the barrier to a credible competitor, although current evidence does not establish such a challenger.

For how these upstream production constraints surface in the four anchor accelerator and AI-factory systems, see §2.13.

Data and sourcing: the equipment & EDA data appendix (research/2.9-equipment-eda-data.md) and the master figures table (§F). Tier-1 anchors: SEMI (WFE market), company filings (ASML, AMAT, LRCX, TEL, KLA, Synopsys, Cadence), TechInsights (share),

Caixin/CNBC (ASML China, EDA controls). Per-vendor WFE share % are estimates, flagged ⚠ in the appendix. Company tags and access: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-chart-eda-market-share] [[^]atlas-chart-semi-wfe-2025] [[^]atlas-prof-applied-q2-2026] [[^]atlas-prof-asml-results-2026] [[^]atlas-prof-c03-kla-q3fy26] [[^]atlas-prof-c03-lam-mar2026] [[^]atlas-prof-tel-fy2026]

[[^]atlas-chart-eda-market-share]: Peking University builds 3D chip design tool tailored to Huawei's LogicFolding architecture (<https://www.tomshardware.com/tech-industry/semiconductors/peking-university-builds-3d-chip-design-tool-tailored-to-huaweis-logicfolding-architecture>). Tom's Hardware, undated; accessed 2026-07-25.

[[^]atlas-chart-semi-wfe-2025]: SEMI Reports Global Total Semiconductor Equipment Sales Forecast to Reach \$125.5 Billion in 2025 (<https://www.semi.org/en/semi-press-release/semi-reports-global-total-semiconductor-equipment-sales-forecast-to-reach-125.5-billion-dollars-in-2025>). SEMI, 2025-07-22; accessed 2026-07-25.

[[^]atlas-prof-applied-q2-2026]: Applied Materials Announces Second Quarter 2026 Results (<https://investors.appliedmaterials.com/news-releases/news-release-details/applied-materials-announces-second-quarter-2026-results>). Applied Materials, 2026-05-14; accessed 2026-07-25.

[[^]atlas-prof-asml-results-2026]: Financial results (<https://www.asml.com/en/investors/financial-results>). ASML, 2026-01-28; accessed 2026-07-25.

[[^]atlas-prof-c03-kla-q3fy26]: Quarterly Reports: quarter ended March 31, 2026 (<https://ir.kla.com/sec-filings/quarterly-reports>). KLA, 2026-04-30; accessed 2026-07-25.

[[^]atlas-prof-c03-lam-mar2026]: Lam Research Reports Financial Results for the Quarter Ended March 29, 2026 (<https://investor.lamresearch.com/2026-04-22-Lam-Research-Corporation-Reports-Financial-Results-for-the-Quarter-Ended-March-29%2C-2026?asPDF=>). Lam Research, 2026-04-22; accessed 2026-07-25.

[[^]atlas-prof-tel-fy2026]: FY2026 Annual Results (<https://www.tel.com/ir/library/report/index.html/>). Tokyo Electron, 2026-04-30; accessed 2026-07-25.

Chapter 2.10 — Materials & Inputs

At the very bottom of the stack sit the raw and refined materials that everything else is built from, and this is the quiet chokepoint that most investors underweight. A ten-dollar consumable can gate a forty-thousand-dollar GPU. The layer's defining feature is a mutual hostage situation: China controls the raw minerals the West cannot easily replace, and Japan and the West control the ultrapure engineered materials China cannot yet make. Each side can hurt the other, which is why this layer is less an investment theme than a map of where the real leverage in the whole US-China contest actually lies.

Before a chip is designed, fabricated, or packaged, it begins as materials: the silicon wafer it is printed on, the light-sensitive photoresist that patterns it, the ultrapure gases that etch it, and, further upstream, the rare-earth elements and specialty minerals woven through the magnets, lasers, and components of the whole system. These inputs are cheap relative to the finished chip, which is exactly why their leverage is so disproportionate: a shortage of a single obscure material can halt a production line worth billions. The layer splits into two worlds with opposite balances of power, and understanding that split is the point of this chapter.

China's half: the minerals

China's dominance of the raw-mineral base is close to total in several categories, and it has shown, repeatedly, that it will use it.

China's grip on the raw minerals

China's share of global supply, 2025 — the leverage it weaponized in 2024–25



Source: USGS Mineral Commodity Summaries 2025. Antimony's export ban ran the price +2,600% before the late-2025 truce.

China controls roughly 90% of the world's rare-earth *processing*, about 98% of gallium and 77% of germanium refining, more than 95% of the graphite anode material for batteries, and around 90% of high-performance magnet production. Over 2024 and 2025 it turned these into weapons in sequence, banning gallium, germanium, and antimony exports to the US, then adding heavy rare earths and magnets, then a rule extending control to any foreign product containing even trace Chinese-origin rare earths. The effect was violent where it landed: antimony, used in everything from munitions to semiconductors, ran from about \$1,400 a tonne to nearly \$60,000, a rise of more than 2,600%, before a truce eased it. That truce, agreed after the Trump-Xi meeting in late 2025, suspends the harshest controls only until late 2026, so the leverage is paused, not removed.[^atlas-chart-usgs-china-minerals-2025]

Critical Materials and Component Suppliers

The materials layer runs across categories, and the balance of power flips from China to Japan and the West as you move from raw minerals to engineered inputs.

Category	Leaders	Where power sits
Silicon wafers	Shin-Etsu, SUMCO, GlobalWafers, Siltronic, SK Siltron	Japan-anchored (>50% of 300mm)
Photoresist	JSR, Tokyo Ohka, Shin-Etsu, Fujifilm	Japan ~90%, EUV resist >95%
Specialty gases	Air Liquide, Linde, Air Products	US/EU; neon exposure to Ukraine
CMP slurry / ultrapure chemicals	Entegris (CMC), Fujimi, Resonac, DuPont	US + Japan
ABF substrate & film	Ibiden, Shinko, Ajinomoto (>95% film)	Japan (see 2.7)
Rare earths (ex-China)	MP Materials, Lynas	US/Australia; policy-backed
Rare earths / gallium / graphite	Chinese state producers	China (the leverage)

The West and Japan's half: the engineered materials

Run the same analysis on the ultrapure, engineered materials, and the dominance flips to the other side of the Pacific.

Japan and the West's grip on the engineered materials

Allied share of global supply, 2025 — the ultrapure inputs China cannot yet make



Source: industry / company disclosures (🚩). The mirror image of the minerals: China holds the raw base, the allies the engineered ti

Japan makes roughly 90% of the world's semiconductor photoresist and more than 95% of the high-end EUV resist, through Shin-Etsu, JSR, and Tokyo Ohka; JSR's ownership of

Inpria, the leading metal-oxide EUV-resist maker, tightens that grip further. Shin-Etsu and SUMCO together supply more than half of all 300mm silicon wafers, with GlobalWafers, Siltronic, and SK Siltron rounding out an oligopoly that controls over 80% of the market. The ultrapure gases that etch and deposit run through Air Liquide, Linde, and Air Products, with a lingering vulnerability in neon (much of it historically from Ukraine, which spiked prices ninefold after the 2022 invasion). The CMP slurries and wet chemicals that polish and clean are led by America's Entegris and Japan's Fujimi and Resonac. And, as Chapter 2.7 detailed, Japan's Ibiden and Shinko hold more than 70% of high-end substrates while a single company, Ajinomoto, makes more than 95% of the insulating film. The materials that require the most process sophistication are held almost entirely by the US, Japan, and Europe, and China cannot yet substitute them at scale.^[^atlas-chart-japan-materials-share]^[^atlas-mat-ajinomoto]

This is the mirror image of the minerals story, and together they define the layer. China holds the upstream raw base; the allied bloc holds the downstream engineered tier. Each can inflict real damage on the other, which produces a fragile, mutually-deterred stability and a scramble on both sides to build redundant supply, the "de-Americanization" of China's bill of materials matched by the West's "de-risking" from China. Both efforts are expensive, slow, and incomplete.

Qualification economics: small cost, large switching friction

The investment logic in engineered materials differs from ordinary commodities. The material may be a small fraction of finished-chip cost, but a purity deviation can destroy yield across an expensive production line. Customers therefore qualify a supplier and a specific formulation through extended process testing; once qualified, they are reluctant to change merely to save a small amount on input cost. That creates retention and pricing power for proven suppliers, while making new capacity slow to convert into commercial share.

Scarcity can still be temporary. Inventory buffers, dual sourcing, recycling and customer redesign gradually weaken a shortage, and commodity prices can fall before a new mine or refinery reaches steady production. The relevant dashboard is consequently broader than spot price: qualification wins, long-term offtake, customer inventories, production yield, cash cost, government support and the time remaining on export controls.

Ajinomoto's position in build-up film illustrates the difference between a qualified process dependency and a generic raw material.[^atlas-mat-ajinomoto]

For rare-earth projects, strategic value and minority-shareholder returns can diverge. A government may rationally support redundant domestic supply even when the project's through-cycle economics are poor. Public investors must distinguish policy support that protects capacity from contracts, price floors or capital structures that actually protect equity returns.

Truce expiry, qualification wins and substitution

The principal near-term policy variable is the minerals truce, which expires in late 2026; whether China re-tightens rare-earth and gallium controls, or lets the truce hold, is a genuine swing factor for defense, magnet, and semiconductor supply chains. Watch whether Western rare-earth capacity, MP Materials and Lynas above all, scales fast enough to matter or stays a rounding error against Chinese output. Watch the engineered-materials chokepoints for their own squeezes, the Ajinomoto price increase for the second half of 2026 being a small example of pricing power at a true single-supplier bottleneck. And watch the slow migration to glass substrates, which could reshuffle the substrate hierarchy over the second half of the decade.

From chokepoint to shareholder return

The most direct public expressions of this layer are the ex-China rare-earth builders, MP Materials (MP) and Lynas (LYC.AX), which are as much policy bets as commodity plays, backed by government price floors and offtake because their strategic value exceeds their current economics. The engineered-materials champions, Shin-Etsu, SUMCO, JSR, Tokyo Ohka, Ibiden, and Ajinomoto, are concentrated process dependencies, though most trade in Japan and may be less accessible to a US investor; Entegris (ENTG) is a listed American materials supplier, and Air Liquide and Linde (LIN) provide diversified industrial-gas exposure. The layer is most useful as a map of bargaining power; the decision layer separately tests whether that strategic leverage is material to each security and already reflected in valuation.

When the leverage gets used

The mutual-deterrence framing breaks if one side decides the leverage is worth using despite the cost. A Chinese decision to let the truce lapse and re-impose hard mineral controls would spike prices and disrupt supply chains well beyond semiconductors, a tail risk that is paused rather than gone. In the other direction, the thesis for the Western rare-earth names weakens if Chinese supply floods back and collapses prices, which has ended every previous attempt to build a non-Chinese rare-earth industry; the government price floors under MP Materials exist precisely because this has happened before. And the allied advantage in engineered materials, seemingly secure, would erode over years if China's heavy investment in domestic wafers, photoresist, and substrates finally began to close the quality gap, which it has not yet.

For the downstream package, rack and 100 MW facility boundaries that consume these inputs, see §2.13.

Data and sourcing: the materials data appendix (research/2.10-materials-data.md) and the master figures table. Tier-1 anchors: USGS (minerals), SEMI (wafers), Reuters and SEC filings (MP Materials, rare earths), company disclosures (Ajinomoto, Shin-Etsu). Market-size and share figures beyond USGS/SEMI are aggregator estimates, flagged ⚠️ in the appendix. Company tags and access: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-chart-japan-materials-share] [[^]atlas-chart-usgs-china-minerals-2025] [[^]atlas-mat-ajinomoto]

[[^]atlas-chart-japan-materials-share]: Japan Seeks to Revitalize Its Semiconductor Industry (<https://www.csis.org/analysis/japan-seeks-revitalize-its-semiconductor-industry>). Center for Strategic and International Studies, undated; accessed 2026-07-25.

[[^]atlas-chart-usgs-china-minerals-2025]: Mineral Commodity Summaries 2025 (<https://www.usgs.gov/publications/mineral-commodity-summaries-2025>). U.S. Geological Survey, 2025-01-31; accessed 2026-07-25.

[^atlas-mat-ajinomoto]: Ajinomoto Build-up Film

(<https://www.ajinomoto.com/innovation/action/buildupfilm>). Ajinomoto, undated; accessed 2026-07-25.

Chapter 2.11 — Capital & Market Structure

This layer is not a technology but the money, and it is where the entire atlas is stress-tested. The AI build is being financed at a scale and in a manner that has begun to worry central banks: capital spending that now exceeds the cash the spenders generate, a web of circular deals in which the industry funds its own demand, a stock market more concentrated than at any point in a generation, and a rising tide of debt collateralized by chips that lose value fast. None of this means the boom is fake. It means the risk has shifted from "will the technology work" to "will the financing hold," and the two questions have different answers.

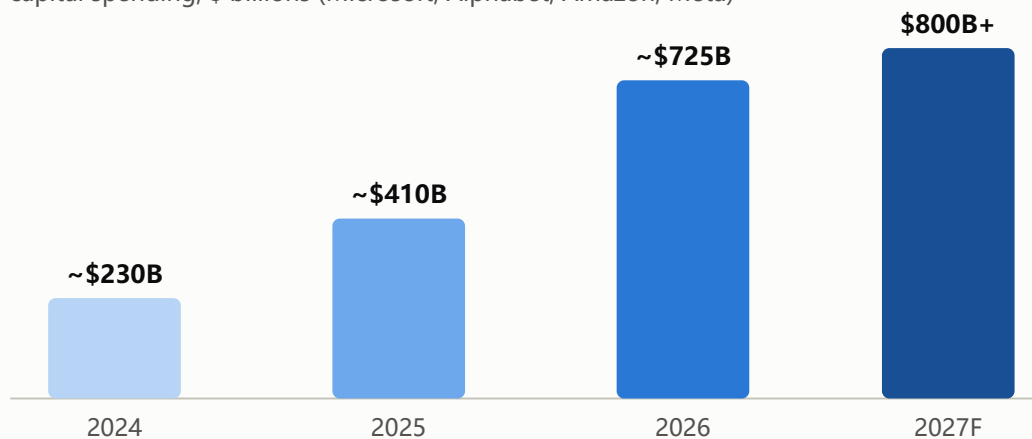
Every prior chapter described something being built. This one describes how it is paid for and how markets are pricing it, because the durability of the money is ultimately what determines whether the whole edifice stands. Several features define the current moment: capital spending on a scale without real precedent, reflexive financial loops that link the major players together, a stock market whose gains have concentrated into a handful of names, a rising reliance on debt, and a private market pricing artificial general intelligence as if it had already arrived. Each is a source of both the boom's momentum and its fragility.

The capex supercycle

The spending is the foundation, and its trajectory is vertical.

The capex supercycle

Big Four hyperscaler capital spending, \$ billions (Microsoft, Alphabet, Amazon, Meta)



Source: company guidance; Goldman Sachs (2027F 🚩). Capex now exceeds operating cash flow — the rest is debt-funded.

The four largest hyperscalers are guiding to roughly \$725B of capital spending in 2026, up about 77% from the prior year, with the five-company total including Oracle running higher still. Company by company, the guidance is staggering: Amazon around \$200B (roughly double 2025), Alphabet \$180–205B (raised twice, up from about \$85B in 2025), Microsoft above \$120B for the fiscal year, and Meta \$125–145B. Goldman Sachs projects cumulative hyperscaler capex across 2025–27 above \$1.15T, more than double the prior three years. The critical, and newer, fact is that this spending now exceeds the operating cash flow of the companies doing it; Amazon's trailing free cash flow has compressed to near \$1B. When capex outruns cash, the balance has to come from debt, and the market's mood has shifted from rewarding capital spending to interrogating the return on it: when Alphabet beat earnings but raised its capex guidance in July 2026, the stock sold off.[^atlas-chart-hyperscaler-capex][^atlas-chart-goldman-capex-2027]

Capital Providers and Financing Channels

The money flows through a recognizable cast: the hyperscalers who spend, the labs and neoclouds who consume, the vendor who finances its own customers, and the private-credit managers who increasingly fund it all.

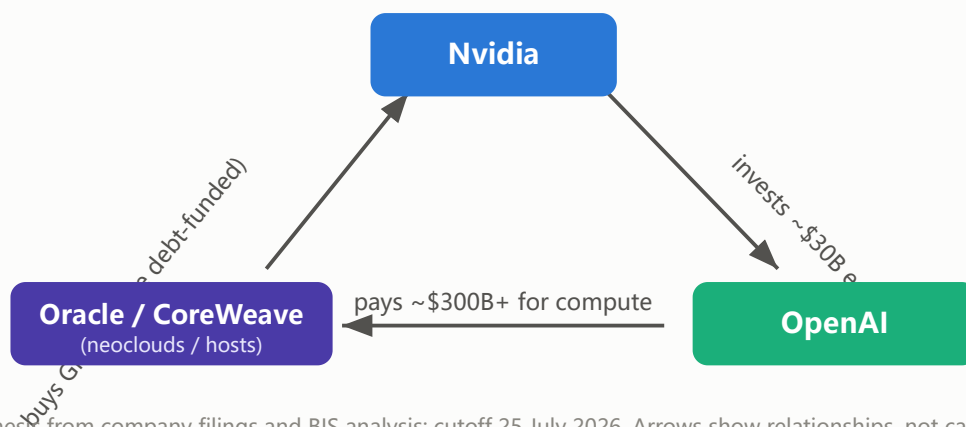
Player	Ticker	Role in the AI money
Microsoft, Alphabet, Amazon, Meta	MSFT, GOOGL, AMZN, META	the ~\$725B/yr capex spenders
Oracle	ORCL	Stargate; a debt-funded RPO of ~\$500B+
OpenAI, Anthropic, xAI	PVT	the labs raising record private rounds
CoreWeave, Nebius	CRWV, NBIS	neoclouds funded by GPU-backed debt
Nvidia	NVDA	vendor-financier: equity stakes in OpenAI, CoreWeave, Nebius
Blue Owl, Apollo, Blackstone, Ares, Brookfield	OWL, APO, BX, ARES, BAM	private-credit originators of the data-center debt
PIMCO, BlackRock	(PIMCO PVT), BLK	anchor buyers of the debt
SoftBank, MGX	PVT	Stargate and mega-round backers

The circular loop

The most-discussed feature of this cycle is the way the money moves in circles.

The circular loop — the industry funds its own demand

A simplified view of the money flow the BIS flagged as a stability concern



The pattern, in its simplest form, is that Nvidia invests in OpenAI, OpenAI commits hundreds of billions to Oracle and CoreWeave for computing capacity, and those

providers spend the money buying Nvidia chips, some of it collateralized by Nvidia's own equity stakes in them. The specific links: Nvidia's headline "up to \$100B" commitment to OpenAI (walked back in early 2026 to a roughly \$30B actual stake); the OpenAI-AMD deal in which the customer received warrants for up to 160M AMD shares; Oracle's roughly \$300B, five-year Stargate commitment sitting inside a remaining-performance-obligation book that ballooned past \$500B; and Nvidia's roughly \$2B equity investments in each of CoreWeave and Nebius, whose purchases of Nvidia chips its own money helps fund. Estimates of the aggregate "circular" web run from about \$800B to \$1.4T ⚠️. The concern the Bank for International Settlements has flagged is that these arrangements can inflate the appearance of demand and link the players' fortunes so tightly that a stumble at one node transmits faster than in a normal supply chain. The market-level point is simple: Nvidia's revenue quality and OpenAI's funding are now connected, and a problem at one shows up quickly at the other. The full node-by-node ledger is the subject of a dedicated deep-dive.[^atlas-chart-bis-ai-finance] [^atlas-chart-company-circular-financing]

Concentration and the crowded trade

The stock market that owns all of this has concentrated to a degree not seen in a generation.

The 2026 rotation — leadership broadens

Approximate year-to-date return through mid-July 2026; group and ex-group calculations differ by methodology



Source: public index and constituent data; rounded analytical calculation. Direction is more robust than the exact spread.

The Magnificent Seven make up roughly 32.5% of the S&P 500's market value, and the top ten companies about 40% of the index, a level Apollo's chief economist Torsten Sløk calls more extreme than the 1990s peak, with a path toward a 50% top-ten weighting. Because passive and target-date flows mechanically bid up the largest names, this makes the index a momentum machine and leaves passive investors structurally overweight AI whether they intend to be or not. The fragility showed in 2026: the Magnificent Seven fell about 3.7% year-to-date while the other 493 S&P names rose about 12.9% ⚠️, a

sharp reversal of the leadership that had roughly doubled index returns over 2023–25. Individual tells accumulated: Alphabet beat but sold off on its capex raise, Palantir fell about 26% year-to-date despite guiding to 71% revenue growth, and short interest clustered in the neoclouds (CoreWeave, Nebius), Super Micro, and the small-modular-reactor names. The rotation into the "S&P 493" is the most important tape change in this cluster, and it argues for fading the crowding.[^atlas-chart-market-rotation-2026]

From equity to credit

Because capex exceeds cash flow, the build is increasingly funded by debt, and this is the most under-appreciated systemic vector in the whole atlas. The five hyperscalers issued roughly \$121B of investment-grade bonds in 2025, more than four times the prior five-year average, with Meta's \$30B October 2025 deal the largest non-M&A investment-grade bond ever; 2026 year-to-date AI-related debt issuance reached about \$489B. Alongside the public bonds, a wave of private credit arrived through vehicles like Meta's roughly \$30B off-balance-sheet financing with Blue Owl and PIMCO, the largest private-credit data-center deal ever, and GPUs themselves are being packaged into asset-backed securities, with CoreWeave amassing over \$14B in GPU-collateralized facilities and its first \$8.5B deal rated investment-grade. The private-credit market has grown from about \$100B in 2010 to roughly \$2.2T in 2026, and an estimated \$800B of it is needed for AI infrastructure through 2028. Both the BIS and the IMF have named this a financial-stability concern, because GPU collateral depreciates quickly and the leverage is partly hidden in less-regulated corners of the market. The leading indicators to watch are credit spreads on hyperscaler and data-center paper and any downgrade of GPU-backed securities.

The ROI and bubble debate

Underneath the spending sits the question of whether it will pay, and the debate is fierce and evenly matched. The bears have three exhibits. First, MIT's Project NANDA found roughly 95% of enterprise generative-AI pilots delivered no measurable profit. Second, Michael Burry argued in late 2025 that hyperscalers understate depreciation by roughly \$176B across 2026–28 by extending the assumed useful life of chips that obsolesce in a few years, which he estimated overstates Oracle's earnings by about 27% and Meta's by about 21% by 2028. Third, Sequoia's David Cahn, author of the original "\$600B

question," now pegs 2026 AI-infrastructure spend around \$1.5T needing roughly \$3T of revenue to justify, while Bain estimates about \$2T of new AI revenue is needed by 2030. The bulls, led by Coatue, argue simply that AI is not in a bubble and that monetization lags capex with a normal delay. As one observer put it, your framework decides your conclusion before the numbers do. The debate is migrating from "is there ROI" to "whose depreciation is honest and whose revenue is real," which is why quality-of-earnings analysis has become a genuine source of edge.

Private megarounds and the IPO wave

The private market is pricing this layer as aggressively as any in history. OpenAI closed a \$122B round at an \$852B valuation in March 2026, the largest private raise ever; Anthropic leapt past it to \$965B on a \$65B round in May, the first challenger to exceed OpenAI; xAI was valued at \$250B inside SpaceX's roughly \$1.25T after an all-stock merger; and the coding startup Anysphere (Cursor) went from a \$29B round to talks above \$50B. These are marks, not prices, set with liquidation preferences, and they price AGI optionality. 2026 is also the most AI-concentrated IPO year on record: CoreWeave has been public since 2025, Cerebras listed in a large 2026 offering, China's CXMT staged the year's biggest A-share IPO at about \$86B, Databricks pushed its listing to late 2026, and an OpenAI S-1 was reported for mid-2026. Retail access has proliferated through vehicles like Destiny Tech100 (DXYZ), which trades at a 50–200% premium to its net asset value 🚩, ARK Venture, and Robinhood Ventures, a sign of how much investors will pay simply for access.

The talent wars and reverse acqui-hires

A quieter but revealing feature is how the giants acquire capability without triggering merger review. Google, Microsoft, Amazon, Meta, and Nvidia deployed more than \$40B through "reverse acqui-hires" from 2024 to mid-2026: Microsoft licensed Inflection's team for about \$650M, Google took Character.AI's founders for about \$2.7B and Windsurf's leadership for about \$2.4B after OpenAI's deal collapsed, and Meta paid about \$14B for 49% of Scale AI and its founder. Compensation escalated to match, with Meta's Superintelligence Labs reportedly offering packages up to \$300M over four years. In February 2026, Senators Warren, Wyden, and Blumenthal urged the FTC and DOJ to probe these structures as "de facto mergers" designed to bypass review, and inquiries

were reportedly opened. For an investor the structures are cheap optionality for the incumbents but carry a new regulatory tail, and the comp inflation is a margin headwind for the labs and a signal that human capital, not just GPUs, is a binding constraint.

China finances it differently

The capital story is largely an American one, and that is itself the point. The circular financing, the private megarounds, the crowded mega-cap trade, and the GPU-backed debt are all features of the US market. China builds its AI on a different capital structure: state guidance funds, the national semiconductor "Big Fund," state-enterprise balance sheets, and government subsidies, rather than venture capital and public equity. That brings its own risks, which the China dossier details: pledged shares (股权质押) that can force selling, subsidies that inflate reported profit (so read the non-GAAP line), state-fund selling overhangs (大基金减持), and fabricated order rumors (小作文). A Chinese AI champion is less likely to be caught in a circular-financing unwind and more likely to be exposed to a shift in state policy. The two systems concentrate their financial risk in different places.

The refinancing and cash-return tests

The signals to watch are financial rather than technological. The first is whether any hyperscaler cuts or pauses its 2027 capex guidance; none has, and the first to do so would be the regime-change signal for the entire supplier complex. The second is credit: spreads on hyperscaler and data-center debt, and any downgrade of GPU-backed securities, are the leading indicators the equity market will lag. The third is quality of earnings, specifically whether any hyperscaler shortens its depreciation assumptions, which would validate the bears and cut reported profits at the most exposed names. And the fourth is the private market, where a down-round or a delayed IPO at OpenAI or Anthropic would transmit through the circular loop quickly.

Positioning: finance the build, screen the earnings

One way to obtain exposure to this financing cycle without owning the most levered operators is through private-credit and alternative-asset managers such as Blue Owl

(OWL), Apollo (APO), Blackstone (BX), and Ares (ARES), which collect fees for originating and managing debt. They exchange direct operating risk for credit-cycle, fundraising and underwriting risk. Among operating companies, the discipline is to compare cash earnings, depreciation policy, customer financing and related-party revenue before treating reported growth as external demand. Portfolio construction should also account for the fact that hyperscalers, neoclouds, electrical suppliers and lenders can all be exposed to the same capex reversal.

The dominoes

This is the layer whose risks would, if they materialized, invalidate much of the rest of the atlas. A hyperscaler cutting capex, a wave of credit-spread widening or a GPU-ABS downgrade, an OpenAI or Anthropic down-round, or a hyperscaler shortening depreciation, any one of these would mark the turn from expansion to contraction and hit the whole supplier chain at once, faster than a normal cycle because the circular loop transmits shocks so quickly. In the other direction, the bullish case is validated if AI revenue growth visibly closes the gap with capex, if the enterprise ROI of Chapter 2.2 broadens, and if the circular deals resolve into real, external, cash-paying demand rather than reflexive marks. The tell that matters most, restated from the executive summary, is whether disclosed hyperscaler AI revenue begins to cover the depreciation on what has been built.

Data and sourcing: the capital, circular-financing, and rotation sidecars

(charts/sidecars/2.11-capex.json, charts/sidecars/2.11-circular.json, and charts/sidecars/2.11-rotation.json), the chart-source registry (audit/chart-sources.json), research/4-trends-forward-evidence.md, and the master figures table (\$A–C, \$H). Tier-1 anchors: company guidance, Goldman Sachs, BIS, Menlo Ventures, and reputable financial reporting. Aggregate capex, circular-financing totals, the analytically derived S&P-493 comparison, and closed-end-fund premiums remain explicitly qualified. Company access and evidence: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-chart-bis-ai-finance] [[^]atlas-chart-company-circular-financing] [[^]atlas-chart-goldman-capex-2027] [[^]atlas-chart-hyperscaler-capex] [[^]atlas-chart-market-rotation-2026]

[[^]atlas-chart-bis-ai-finance]: The AI investment race (<https://www.bis.org/publ/work1367.htm>). Bank for International Settlements, 2026-07-14; accessed 2026-07-25.

[[^]atlas-chart-company-circular-financing]: From OpenAI to Nvidia, firms channel billions into AI infrastructure as demand booms (<https://www.investing.com/news/stock-market-news/factboxfrom-openai-to-nvidia-firms-channel-billions-into-ai-infrastructure-as-demand-booms-4605465>). Reuters, undated; accessed 2026-07-25.

[[^]atlas-chart-goldman-capex-2027]: What the capex boom means for stock investors (<https://www.goldmansachs.com/insights/articles/what-the-capex-boom-means-for-stock-investors>). Goldman Sachs, 2026-07-21; accessed 2026-07-25.

[[^]atlas-chart-hyperscaler-capex]: Big Tech Is on Track to Spend \$750 Billion on AI This Year (<https://www.forbes.com/sites/aliciapark/2026/04/30/big-tech-is-on-track-to-spend-750-billion-on-ai-this-year/>). Forbes, 2026-04-30; accessed 2026-07-25.

[[^]atlas-chart-market-rotation-2026]: Why the Mag 7 stocks are underperforming the S&P 500 in 2026 (<https://m.investing.com/news/stock-market-news/why-the-mag-7-stocks-are-underperforming-the-sp-500-in-2026-93CH-4796219?ampMode=1>). Investing.com, 2026-07-18; accessed 2026-07-25.

Chapter 2.12 — Geopolitics, Policy & Risk

Every layer in this atlas operates inside rules written by governments, and this chapter is about those rules and the risks they create. A major recent development is that, for a strategic slice of demand, Beijing is using procurement and market access to force domestic adoption while the United States combines restrictions with negotiated licensing and fees. The policy regime is increasingly transactional, raising political volatility around every name in the industry. Underneath it sits the risk that can be sized but not readily hedged: Taiwan.

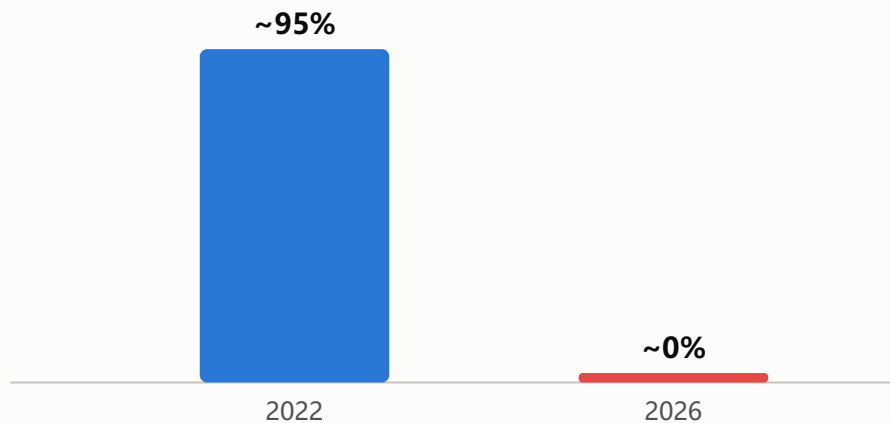
The AI industry is not a free market; it is a managed one, shaped at every layer by export controls, tariffs, industrial policy, and the slow split of one global supply chain into two rival blocs. Understanding this layer means holding several things at once: the whipsaw of US chip policy, the localization drive on the Chinese side, the sovereign-AI demand pool now being courted by both, and the concentration risks, Taiwan above all, that a single event could turn into a crisis. This chapter draws together threads from every earlier one, because policy is the force acting on all of them.

The controls inverted

For years the story was American denial: Washington restricting what China could buy. That story has flipped.

Nvidia's China collapse

Nvidia's share of China's AI-chip market — the final step taken by Beijing, not Washington



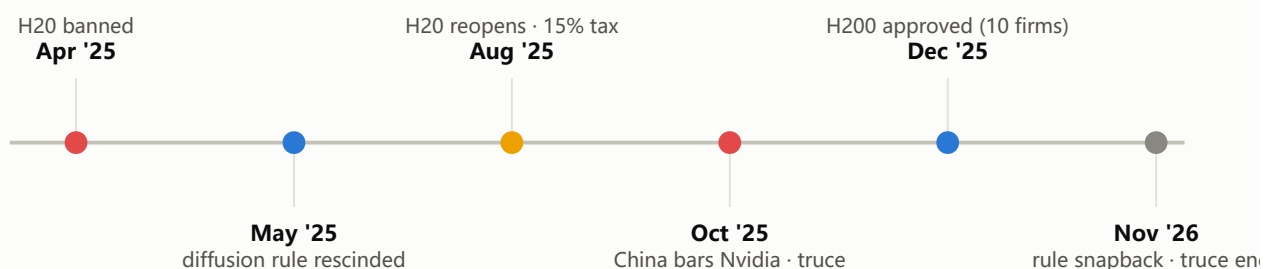
Source: Brookings, Caixin (2026). The US banned the H20, then taxed re-exports; China then barred domestic firms from buying.

Nvidia's share of the Chinese AI-chip market fell from around 95% in 2022 to effectively zero by 2026, but the final step was taken by Beijing, not Washington. After the US banned the high-end H20 in early 2025, then reopened exports under an unprecedented arrangement that routed 15% of the revenue to the US Treasury, and finally approved sales of the more capable H200 to a handful of Chinese firms, it was China's own regulators who told domestic companies not to buy American chips and barred foreign accelerators from state-funded data centers. The result is that the binding constraint on Nvidia in China is now Chinese policy, and the genuinely unresolved question is who is restricting whom. The mandate regime behind this, the "autonomous and controllable" (自主可控) localization drive, is detailed in its own deep-dive; the point here is that it has made China's demand for domestic chips a matter of law rather than competitiveness. [^atlas-chart-nvidia-china-share][^atlas-chart-us-export-controls-timeline]

The policy tools have also changed character, from a rules-based system of published controls to a series of discretionary, case-by-case deals. The timeline of the past eighteen months shows the whipsaw plainly.

The policy whipsaw

US-China chip policy, 2025–26 — rules gave way to case-by-case deals



Analytical synthesis from BIS/Commerce rules and dated notices; cutoff 25 July 2026. Later dates are scheduled policy events.

Policy, Trade, and Regulatory Instruments

The government hand now reaches into the industry through a recognizable set of levers, each with its own status and its own expiry.

Lever	What it is	Status
Export controls / Entity List	trade bans on chips and tools	H20 ban → H200 approval; the "diffusion rule" rescinded
15% revenue-share tax	a pay-to-export levy on China chip sales	in force since Aug 2025
Section 232 tariff	25% on certain advanced chips	in force since Jan 2026 (100% threatened)
US equity stake in Intel	~10% government ownership	since Aug 2025; grant-to-equity template
50% affiliates rule	extends controls to majority-owned subsidiaries	snaps back Nov 10, 2026
China mineral controls	rare-earth / gallium / antimony export limits	truce until ~Nov 27, 2026
Sovereign-AI chip diplomacy	government-to-government GPU allocations	Gulf, and expanding

Two stacks, and the fight for the third bloc

The cumulative effect of controls and localization is bifurcation: the emergence of two largely separate technology stacks, an American one built on Nvidia and CUDA and a Chinese one built on Huawei's Ascend and CANN, each with its own hardware, software, and increasingly its own standards, right down to rival AI-governance forums. This is now the base case rather than a risk, and the important contest has moved to the middle. The prize is the "third bloc", the Gulf states, Southeast Asia, and the Global South, that belong to neither camp yet and whose choice of stack is up for grabs. The United States is using chip diplomacy to court them, while China offers subsidized, string-free Ascend systems to buyers who cannot or will not access US chips. The decisive question of the whole geopolitical layer has become which stack the uncommitted world adopts.

Sovereign AI: the new demand pool

A distinct new source of demand has emerged from all this: sovereign AI, national compute funded by governments, projected above \$100B in 2026. It is anchored by US chip diplomacy. Following a 2025 Gulf tour, Washington authorized the UAE's G42 and Saudi Arabia's HUMAIN to buy Nvidia Blackwell; the Stargate UAE campus (with G42, Nvidia, OpenAI, Oracle, and Cisco) targets 5 GW with a first 200 MW cluster live in 2026, and HUMAIN is targeting up to 600,000 Nvidia GPUs over three years. Europe is backing up to five "AI gigafactories" with around €20B, India and Japan are funding national models and compute, and essentially every major economy now has a program. Sovereign AI is the demand pool that most offsets Nvidia's lost China revenue, but it is lower-quality demand than the hyperscalers': it is policy-driven, oil-price-sensitive, and comes with security strings (no Huawei, US-vetted operators), so it deserves more skepticism than a hyperscaler backlog, and the Gulf is explicitly hedging by courting China as well.

Tariffs and industrial policy

The domestic-policy tool has shifted from grants to tariffs and equity. The Section 232 semiconductor tariff landed at 25% on certain advanced chips in January 2026, narrower than the "100%" once threatened, which functions mainly as a lever to extract US-investment pledges. The marquee industrial-policy move is the US government taking a roughly 10% equity stake in Intel, converting unpaid CHIPS Act grants into ownership, a template that could extend to other recipients. And allied coordination has deepened, with Washington pressing the Netherlands (ASML) and Japan (Tokyo Electron) to restrict not just sales but the servicing of installed tools in China, though the allies push back where their own revenue is at stake. The through-line is a move from carrots to sticks and ownership, and it taxes and reshapes the whole build.

The Taiwan tail, and the risk register

Above every specific policy sits the structural risk that Chapter 2.8 introduced: the concentration of leading-edge manufacturing in Taiwan. It is the one exposure that cannot be diversified away before roughly 2030, and a disruption in the Taiwan Strait is the single event that would invalidate the entire AI build at once, upstream of Nvidia,

Apple, AMD, and Broadcom alike. It can be sized and respected, but not hedged. Around it sit the more manageable risks this chapter has cataloged: the possibility that export controls escalate and cut the toolmakers' China revenue, that the minerals truce lapses and re-weaponizes rare earths, that tariffs raise the cost of the whole build, and that the transactional policy regime produces sudden, headline-driven swings in individual names. These are the risks to price into any position in the industry.

Licences, procurement rules and the Taiwan clock

The near-term calendar is unusually concrete. Two dates in late 2026 matter most: the scheduled snapback of the US "50% affiliates" rule extending controls to majority-owned subsidiaries, and the expiry of the US-China minerals truce, either of which could re-tighten the screws. Beyond those, watch whether a formal replacement for the rescinded chip-diffusion framework ever lands or whether case-by-case licensing becomes permanent; whether any Blackwell-class chip is approved for China; whether the Gulf and Southeast Asian sovereign projects energize on their US-supplied hardware or defect toward Ascend; and, above all, the temperature in the Taiwan Strait, which no calendar can predict.

Positioning for a transactional regime

Policy is a lens on the other chapters more than a set of standalone trades. For Nvidia (NVDA) and AMD (AMD), the read is that China is now optionality rather than a base case, already written toward zero, while sovereign AI is the offsetting demand to watch. The ex-China rare-earth names, MP Materials (MP) and Lynas (LYC.AX), are direct expressions of the minerals leverage. Trade-compliance and reshoring beneficiaries, including the US-footprint expanders of Chapters 2.8 and 2.9, gain from the complexity itself. And TSMC (TSM) carries the Taiwan tail that sits behind the whole complex. The overarching discipline is that the transactional, deal-driven policy regime makes the industry more headline-sensitive than its fundamentals alone would imply, so position sizing and a tolerance for policy whipsaw matter as much as stock selection.

The event that overrides everything

One risk dominates all others: a Taiwan Strait crisis would break not just this chapter's thesis but the entire atlas, and it is the reason to hold any AI exposure with humility about tail risk. Short of that, a sharp escalation of export controls would hurt the toolmakers and Nvidia's residual China business, while a durable détente would relieve both. A lapse of the minerals truce would spike input costs across defense and semiconductors. And the bifurcation thesis itself would soften if the two stacks found reasons to re-converge, or harden further if the third bloc splits cleanly between them, which is the direction the evidence currently points.

For the physical U.S.–China system comparison, including what is confirmed, reported, inferred or still roadmap-only, see §2.13.

Data and sourcing: the policy and market-share sidecars (charts/sidecars/2.12-timeline.json and charts/sidecars/2.12-nvidia-china.json), the chart-source registry (audit/chart-sources.json), the canonical claims database, and the master figures table (§F–G). Tier-1 anchors: BIS/Commerce announcements and USGS, supplemented by reputable financial reporting. Sovereign-AI spending, Gulf chip tranches, and the reported Chinese grid plan remain explicitly qualified. Company access and restriction evidence: Part V.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-chart-nvidia-china-share] [[^]atlas-chart-us-export-controls-timeline]

[[^]atlas-chart-nvidia-china-share]: Nvidia CEO says the company's China market share has fallen to zero (<https://apnews.com/article/1ae6228c4928ddbb43f984e9b38f49dd>). Associated Press, 2026-06-29; accessed 2026-07-25.

[[^]atlas-chart-us-export-controls-timeline]: Department of Commerce Revises License Review Policy for Semiconductors Exported to China (<https://www.bis.gov/press-release/departments-commerce-revises-license-review-policy-semiconductors-exported-china>). U.S. Bureau of Industry and Security, 2026-01-13; accessed 2026-07-25.

Chapter 2.13 — From GPU to AI Factory: Four Supply-Chain Teardowns

The investable unit of AI infrastructure is no longer the accelerator alone. It is a chain that begins with logic dies and high-bandwidth memory, passes through packages, trays, switches, power shelves and liquid-cooling loops, and ends only when a commissioned data center can deliver electricity, reject heat and keep the system utilized. Across the four anchor platforms in this chapter, the clearest current-to-next signals are more memory bandwidth, twice the scale-up fabric per U.S. rack, much larger Chinese system scale, warmer liquid cooling and a broader power-quality burden. The less comfortable conclusion is just as important: only Blackwell Ultra has enough disclosed system-power data to translate honestly into systems per 100 MW. Everywhere else, “N/A” is more decision-useful than an invented number.

Executive Summary

The AI-factory supply chain has three accounting boundaries, and an investor should never mix them. The **accelerator package** contains the compute dies, HBM and package-level interconnect. The **complete system** adds CPUs, scale-up switching, scale-out networking, power conversion, cooling distribution and physical integration. The **100 MW commissioned IT-nameplate facility** adds upstream electrical and cooling plant, building works and recurring power, water and maintenance—but, in this model, excludes active IT. A rack price is therefore not a facility price, and a 160-cabinet SuperPoD is not comparable with a one-rack NVL72 until both are normalized to the same power boundary.

Five findings matter most.

1. **The U.S. generation delta is bandwidth- and fabric-heavy, not GPU-count-heavy.** GB300 NVL72 and Vera Rubin NVL72 both contain 72 GPUs and 36 CPUs and both expose 20.736 TB of GPU memory. Rubin raises rack HBM bandwidth from 576 to about 1,580 TB/s, NVLink bandwidth from 130 to 260 TB/s and scale-out bandwidth from 0.8 to 1.6 Tb/s per GPU. It also raises CPU memory from 17.28

to 54 TB. That shifts addressable content toward HBM4, host memory, switch silicon, networking and power-quality components without requiring more accelerators per rack.[^factory-nvidia-gb300][^factory-nvidia-rubin]

2. **China's next-generation disclosure is a scale-out roadmap, not yet a shipped like-for-like replacement.** The installed Atlas 900 A3 is a 16-cabinet, 384-Ascend-910C system. Huawei physically demonstrated a 1,024-card Atlas 950 SuperPoD on 17 July 2026, but the separately disclosed roadmap maximum is 8,192 cards across 160 cabinets with a Q4 2026 target. The demo and the roadmap maximum are different configurations. Neither should be described as a currently shipping 8,192-card product.[^factory-huawei-a3][^factory-huawei-950]
3. **Power disclosure is the comparability gate.** NVIDIA specifies up to 142 kW for GB300 NVL72. At an 80 MW compute allocation inside a 100 MW IT-nameplate facility, that supports about **563 racks**, with **493–620** across the 70%–88% compute-share sensitivity. NVIDIA has not published Rubin rack power, and Huawei has not published complete-system power for the A3 or Atlas 950 configurations. No normalized system count is shown for those platforms.
4. **The facility is a second BOM, not a rounding error.** The modeled installed, facility-side capex for a 100 MW IT-nameplate site is **\$1.118–\$1.814 billion in the United States** and **\$0.661–\$1.174 billion in China**, with base cases of \$1.372 billion and \$0.895 billion. These are national benchmark cases, not quotes and not proof that any named supplier serves a platform. At the common base case of 1.25 PUE and 65% utilization, each site consumes about **711.75 GWh annually**. Electricity is about \$61.35 million per year in the U.S. base case and \$60.15 million in the China base case; facility maintenance adds about \$54.88 million and \$35.80 million respectively.
5. **Evidence quality and investability are different questions.** A confirmed design-in may still offer poor public-market access; an accessible supplier may have unreported allocation; a roadmap may point to structural demand without supporting near-term revenue. The matrix at the end therefore separates content growth, access, evidence, customer concentration and substitution risk instead of collapsing them into a buy score.

Evidence status is displayed throughout as Confirmed , Reported , Inferred , Roadmap or Speculative . Confirmed and reported facts may enter calculations when their boundaries are compatible; inferred values enter only when the arithmetic is transparent;

roadmap values remain forward-looking; speculative values never enter totals or investment comparisons.

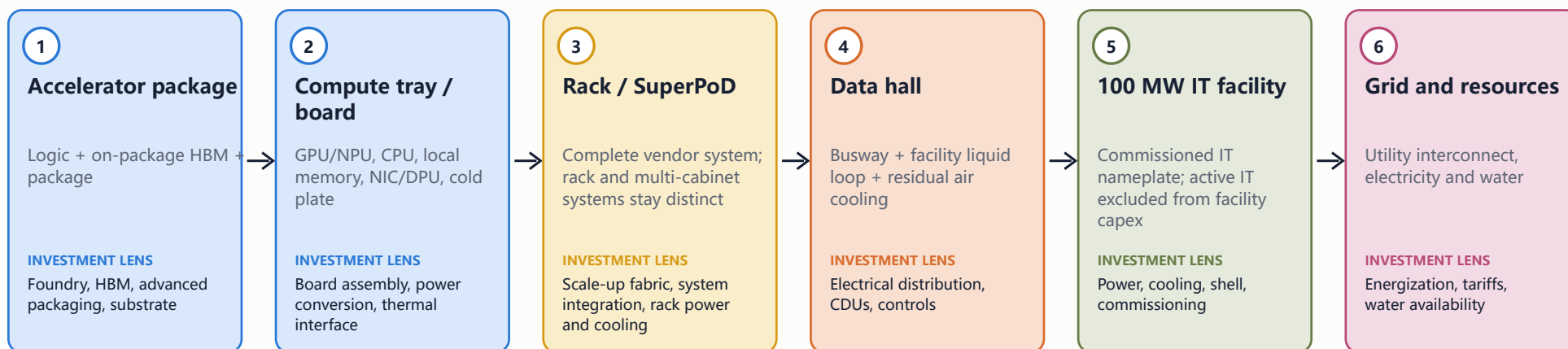
Source quality is a separate dimension. A first-party product specification can be high-quality evidence that a vendor made a claim while the claim itself remains forward-looking; a reputable reported teardown can be a sound secondary source while the supplier relationship remains only “Reported.” The badges describe certainty of the relationship or value, not a score for the publisher.

From chip to grid: where the BOM actually ends

An accelerator becomes useful compute only after six nested physical layers work together.

Chip-to-grid AI factory anatomy

Physical and accounting boundaries used in §2.13; evidence cutoff 2026-07-25



Two accounting paths meet at the rack boundary

Power: grid → switchgear/UPS → busway → rack power shelves

Heat: chip → cold plate → CDU → facility water loop → heat rejection

Sources: platform vendors, U.S. DOE, Turner & Townsend, Cushman & Wakefield; synthesis by AI Investment Atlas.

Layer	What is physically added	Principal constraint	Accounting treatment here
Accelerator package	Logic dies, HBM, substrate, package interconnect	Yield, HBM supply, advanced packaging	Platform BOM
Tray or board	CPUs, GPUs, local memory, NICs/DPUs, voltage conversion, cold plates	Signal integrity, power density, thermal transfer	Platform BOM
Rack / supernode	Scale-up switch trays, cabling, power shelves, manifolds, management	Fabric bandwidth, rack power, serviceability	Platform BOM
Data hall	Busway, CDUs and heat exchangers, row networking, physical deployment	Concurrent maintainability, hydraulic and electrical distribution	Facility BOM except active IT
Campus	Substation, switchgear, UPS/storage, generation, central cooling, shell	Interconnection lead time, water, construction and commissioning	Facility BOM
Grid	Generation and transmission outside the site boundary	Available capacity, price and delivery timing	Excluded remote grid works; electricity is recurring opex

There are two flows. Electricity runs from grid connection to transformer and substation, medium-voltage switchgear, UPS or energy storage and standby generation, low-voltage distribution, busway and finally the rack power shelf. Heat runs in reverse: chip to cold plate, technology loop, CDU and heat exchanger, facility water loop, then chiller, dry cooler or cooling tower to ambient. The rack is the handoff point. A cold plate or power shelf shipped as part of the compute system belongs to the platform BOM; a facility-side CDU, busway or chiller belongs to the facility BOM. The physical architectures are described in the U.S. Department of Energy data-center design guide

(<https://www.energy.gov/sites/default/files/2024-07/best-practice-guide-data-center-design.pdf>).

The chapter therefore reports three different answers for every platform:

- **What is in one accelerator package?**

- **What is in the disclosed complete system boundary?**
- **What can be said at 100 MW commissioned IT nameplate without guessing?**

For deeper treatment of generation, transmission, cooling and data-center constraints, see §2.3. Accelerator architecture belongs in §2.5; scale-up and scale-out fabrics in §2.6; HBM and advanced packaging in §2.7; foundry concentration in §2.8; equipment and EDA in §2.9; materials in §2.10; capital formation in §2.11; and export controls and stack bifurcation in §2.12. Part V remains the canonical home for company profiles, listings and security-level access.

Four platform boundaries and reference configurations

Package and full-system boundaries; system power shown only when published; cutoff 2026-07-25

US • CURRENT

SHIPPING

NVIDIA GB300 NVL72

SYSTEM BOUNDARY

1 liquid-cooled rack

ACCELERATORS

72 packages/cards

SYSTEM POWER

142 kW published maximum

Status date: 2026-07-25

solid = shipping/ramping

US • NEXT

RAMPING

NVIDIA Vera Rubin NVL72

SYSTEM BOUNDARY

1 liquid-cooled compute rack

ACCELERATORS

72 packages/cards

SYSTEM POWER

Not publicly disclosed

Status date: 2026-07-25

solid = shipping/ramping

CHINA • CURRENT

SHIPPING

Atlas 900 A3 SuperPoD

SYSTEM BOUNDARY

16 cabinets: 12 compute + 4 bus

ACCELERATORS

384 packages/cards

SYSTEM POWER

Not publicly disclosed

Status date: 2026-07-17

solid = shipping/ramping

CHINA • NEXT

ANNOUNCED

Atlas 950 SuperPoD

SYSTEM BOUNDARY

160-cabinet roadmap maximum: 128 compute + 32 communications

ACCELERATORS

8,192 packages/cards

SYSTEM POWER

Not publicly disclosed

Status date: 2026-07-17

dashed = roadmap maximum

Sources: NVIDIA and Huawei product/architecture disclosures. Blank power means not publicly disclosed.

Anchor	Package boundary	Complete-system boundary	100 MW normalization
Blackwell Ultra	One B300 package	One GB300 NVL72 rack	Supported from disclosed 142 kW rack maximum
Vera Rubin	One Rubin package	One Vera Rubin NVL72 rack	N/A: rack power undisclosed
Ascend 910C	One 910C accelerator	One 16-cabinet Atlas 900 A3	N/A: system power undisclosed
Ascend 950DT	One 950DT roadmap accelerator	1,024-card demonstration and separate 8,192-card roadmap maximum	N/A: both system powers undisclosed

U.S. installed: Blackwell Ultra / GB300 NVL72

Identity and lifecycle. Confirmed • shipping NVIDIA announced Blackwell Ultra on 18 March 2025 and described GB300 NVL72 as available by the 25 July 2026 cutoff. The system boundary used here is one NVL72 compute rack: 72 Blackwell Ultra GPUs and 36 Grace CPUs. It does not include a surrounding data hall, external storage or campus networking. NVIDIA's GB300 product page (<https://www.nvidia.com/en-us/data-center/gb300-nvl72/>) and launch announcement (<https://nvidianews.nvidia.com/news/nvidia-blackwell-ultra-ai-factory-platform-paves-way-for-age-of-ai-reasoning>) are the lifecycle anchors.

Package and system anatomy

One Blackwell Ultra B300 is a unified package containing two reticle-limited compute dies linked by NVIDIA's 10 TB/s NV-HBI. The disclosed package has 208 billion transistors, 288 GB of HBM3E in eight 12-high stacks, up to 8 TB/s of HBM bandwidth, 1.8 TB/s bidirectional NVLink 5 bandwidth and a maximum total graphics power of 1.4 kW. The 1.4 kW value is a ceiling for the GPU package, not an average rack-load assumption. NVIDIA identifies the process as TSMC 4NP. NVIDIA's Blackwell Ultra architecture note (<https://developer.nvidia.com/blog/inside-nvidia-blackwell-ultra-the-chip-powering-the-ai-factory-era/>) supplies the package boundary.

The complete rack contains 18 compute trays, each with four GPUs and two Grace CPUs; nine NVSwitch trays with two NVSwitch ASICs each; 72 ConnectX-8 adapters; 18 BlueField-3 DPUs; and eight 33 kW power shelves. NVIDIA specifies up to 142 kW for the rack. Direct liquid cooling is part of the physical system boundary. NVIDIA's GB300 reference architecture (<https://docs.nvidia.com/enterprise-reference-architectures/nvl72-ai-factory/latest/components.html>) provides the tray and rack counts.

GB300 boundary	Exact disclosed content	Evidence
One B300 package	2 compute dies; 208B transistors; 288 GB HBM3E; 8 TB/s HBM; 1.8 TB/s NVLink; 1.4 kW maximum TGP	Confirmed
One compute tray	4 B300 GPUs; 2 Grace CPUs; 4 ConnectX-8 adapters; 1 BlueField-3 DPU	Confirmed
One NVL72 rack	18 compute trays; 72 GPUs; 36 CPUs; 9 switch trays; 18 NVSwitch ASICs	Confirmed
Rack power and cooling	Up to 142 kW; direct liquid cooling	Confirmed
Aggregate rack GPU memory	20.736 TB	Inferred: 72 × 288 GB
Aggregate rack HBM bandwidth	576 TB/s	Inferred: 72 × 8 TB/s

Facility implication: the one platform that can be normalized

The base facility reserves 80 MW of its 100 MW IT nameplate for compute, with 10 MW for networking, 7 MW for storage and 3 MW for management and security. Dividing 80 MW by the disclosed 142 kW maximum rack load yields **563.38 GB300 NVL72 racks**. A 70%–88% compute-share sensitivity produces **492.96–619.72 racks**. These are engineering occupancy equivalents, not a procurement forecast: hall geometry, stranded capacity, maintenance reserve and actual utilization will reduce or redistribute deployable counts.

100 MW IT-nameplate case	Compute share	Compute power	GB300 power used	Normalized NVL72 racks
Low compute allocation	70%	70 MW	142 kW maximum per rack	492.96
Base	80%	80 MW	142 kW maximum per rack	563.38
High compute allocation	88%	88 MW	142 kW maximum per rack	619.72

The count does not include the cost of those racks in facility capex. Active accelerators, CPUs, memory, storage, rack networking, optics and active-IT spares are excluded from the 100 MW facility BOM.

Supplier geography and evidence

The current rack is globally manufactured even when final integration moves closer to U.S. demand. NVIDIA (NVDA) designs the platform. TSMC (TSM) fabricates Blackwell Ultra on 4NP; NVIDIA has separately confirmed Blackwell-family production at TSMC Arizona, but public evidence does not allocate GB300 wafer volume between Arizona and Taiwan. Micron Technology (MU) says its 36 GB 12-high HBM3E is designed into GB300 and that its SOCAMM memory was co-developed for the platform. SK hynix (000660.KS) has displayed a GB300 module using its 36 GB HBM3E. Foxconn (Hon Hai Precision, 2317.TW) identifies itself as a pilot-build supplier and its Ingrasys unit as a liquid-cooling design participant. Wistron (3231.TW) says its Fort Worth D1 plant mass-produces GB300 boards.

Role	Named entity	Geography evidenced	Relationship	What is <i>not</i> evidenced
Platform design	NVIDIA	United States	Confirmed	Customer or geography mix
Logic fabrication	TSMC	4NP confirmed; Arizona Blackwell-family production separately confirmed	Confirmed process	GB300 Arizona/Taiwan allocation
HBM3E / SOCAM M	Micron	U.S.-headquartered; manufacturing allocation not disclosed here	Confirmed design-in	Unit or revenue share
HBM3E	SK hynix	Korea-headquartered; manufacturing allocation not disclosed here	Confirmed module evidence	Unit or revenue share
Pilot build / cooling	Foxconn / Ingrasys	Taiwan-centered manufacturing network	Confirmed participation	Rack allocation
Board production	Wistron	Fort Worth, Texas	Confirmed	Platform-wide allocation

What changes, who gets paid and when

For the current generation, the most direct commercial exposure is already in production: NVIDIA's platform silicon and systems, HBM3E from named memory partners, TSMC 4NP wafers, switch and networking silicon within the NVIDIA rack, and board/system integration. The important qualification is allocation.

Participation is confirmed; the percentage of GB300 units, wafers or HBM stacks supplied by each partner is not. The chapter therefore does not convert a design-in into supplier revenue.

The 12–36 month question is how long GB300 remains a volume bridge while Rubin ramps, and whether named memory and integration suppliers retain or expand content during the transition. The 3–7 year question is whether annual platform cadence makes each installed generation a short-lived revenue wave or

expands the service, replacement, networking and facility-content pool enough to offset faster obsolescence.

Risks

- The maximum 142 kW specification may exceed average operating draw, so normalized rack counts are conservative occupancy equivalents rather than energy forecasts.
- A confirmed supplier relationship does not establish supplier share, pricing or gross margin.
- TSMC Arizona production is a family-level statement; assigning a GB300 geographic share would overstate the evidence.
- Rapid Rubin adoption can compress the GB300 revenue window even while installed-base support and facility spending continue.
- Taiwan concentration, export controls and the capital cycle are cross-layer risks; see §§2.8, 2.11 and 2.12.

U.S. ramping: Vera Rubin / NVL72

Identity and lifecycle. **Ramping • broad availability targeted H2 2026** NVIDIA said Rubin silicon was in full production in January 2026 and named system builders manufacturing Vera Rubin systems by 31 May 2026. Broad availability remained an H2 2026 milestone at the cutoff. “Ramping” is therefore more accurate than either “roadmap only” or “fully installed.” The boundary is one Vera Rubin NVL72 compute rack. NVIDIA’s broader five-rack POD adds separate Vera CPU, Groq 3 LPX, BlueField-4 STX and Spectrum-6 SPX racks; those companion racks are excluded from the NVL72 BOM below. NVIDIA’s Rubin launch (<https://nvidianews.nvidia.com/news/rubin-platform-ai-supercomputer>) and production-ramp announcement (<https://nvidianews.nvidia.com/news/vera-rubin-full-production-agentic-ai-factory>) anchor the status.

Package and system anatomy

One Rubin GPU contains two reticle-limited compute dies connected by NV-HBI. NVIDIA’s preliminary specification is 336 billion transistors, 288 GB of HBM4, 22 TB/s of HBM bandwidth, 50 PFLOPS of NVFP4 inference compute and 3.6 TB/s bidirectional NVLink 6 bandwidth. The process node, HBM stack count, package supplier and GPU TGP were not disclosed by the cutoff and remain blank rather than being populated from analyst reports. The Rubin GPU architecture note (<https://developer.nvidia.com/blog/inside-nvidia-rubin-gpu-architecture-powering-the-era-of-agentic-ai/>) and Vera Rubin specification page (<https://www.nvidia.com/en-us/data-center/vera-rubin-nvl72/>) define the package.

The rack preserves 72 GPUs and 36 CPUs. NVIDIA shows two Vera Rubin Superchips per compute tray, or four GPUs and two CPUs, implying 18 compute trays. It discloses 36 NVLink 6 switch ASICs and four ASICs per switch tray, implying nine switch trays. The 18- and nine-tray counts are arithmetic inferences, not separately published line items. The platform introduces BlueField-4, 1.6 Tb/s per-GPU scale-out connectivity, 45°C warm-water single-phase direct liquid cooling, rack-level power smoothing and roughly six times Blackwell Ultra’s local energy buffering. NVIDIA says thermal performance nearly doubles in the same footprint, but it does **not** publish a rack-kW value. NVIDIA’s platform architecture

(<https://developer.nvidia.com/blog/inside-the-nvidia-rubin-platform-six-new-chips-one-ai-supercomputer/>) supports these relationships.

Vera Rubin boundary	Exact or derived content	Evidence
One Rubin package	2 compute dies; 336B transistors; 288 GB HBM4; 22 TB/s HBM; 3.6 TB/s NVLink	Preliminary first-party specification
One NVL72 rack	72 Rubin GPUs; 36 Vera CPUs	Confirmed configuration
Compute trays	18	Inferred from 4 GPUs and 2 CPUs per tray
NVLink switch ASICs / trays	36 ASICs / 9 trays	36 disclosed; 9 inferred at 4 per tray
Rack GPU memory	20.736 TB	72 × 288 GB
Rack HBM bandwidth	About 1,584 TB/s; displayed as 1,580 in rounded platform comparison	72 × 22 TB/s
Rack power	Not disclosed	N/A — no estimate enters totals

Facility implication: demand signal without a rack count

Rubin’s facility impact is directionally clear but not numerically normalizable at the system level. Warm-water liquid cooling can change heat-rejection design; power smoothing and larger local buffering add rack-level power-quality content; and doubled scale-up and scale-out bandwidth increase the networking burden. None of these disclosures supplies rack power. Reusing GB300’s 142 kW, or adopting a third-party Rubin estimate, would create a false systems-per-100-MW comparison. The normalized count is therefore **N/A**.

Facility question	What the evidence supports	What remains unavailable
Cooling architecture	45°C warm-water, single-phase direct liquid cooling; higher-flow manifold	Facility WUE, flow requirement and cooling-plant capex per rack
Power quality	Rack-level smoothing; about 6× Blackwell Ultra local energy buffering	Rack input power and buffering component value
Density	Nearly 2× thermal performance in the same rack footprint	Rack kW and hall-level deployable count
100 MW normalization	Facility-side national cost and opex scenarios still apply	Number of Rubin NVL72 racks per 100 MW

Supplier geography and evidence

NVIDIA names Dell Technologies (DELL), Hewlett Packard Enterprise (HPE), Lenovo (0992.HK), Super Micro Computer (SMCI), Foxconn, Quanta/QCT (2382.TW), Wistron and Wiwynn (6669.TW) among Vera Rubin system builders in production. This confirms participation, not share. Micron says 36 GB 12-high HBM4 and SOCAMM2 designed for Vera Rubin entered volume shipment from Q1 2026. Samsung Electronics (005930.KS) describes mass-production HBM4 and SOCAMM2 as designed for Vera Rubin. SK hynix describes a multi-year co-development and supply relationship with NVIDIA for Vera Rubin memory, but the cited disclosure does not allocate a specific Rubin HBM SKU. Wistron says its Fort Worth site will produce Vera Rubin Superchips.

Role	Named entity or group	Relationship	Evidence boundary
Platform design and rack fabric	NVIDIA	Confirmed	Product specification; commercial allocation not applicable
HBM4 / SOCAMM2	Micron	Confirmed design and volume shipment	Vera Rubin design-in; supplier share undisclosed
HBM4 / SOCAMM2	Samsung	Confirmed design relationship	Supplier share undisclosed
Memory partnership	SK hynix	Reported co-development and supply	No cited SKU allocation
Superchip production	Wistron	Named future Fort Worth production	Volume and start date undisclosed
System building	Dell, HPE, Lenovo, Supermicro, Foxconn, QCT, Wistron, Wiwynn	Participation confirmed	No unit allocation
Foundry / package supplier	Not disclosed	N/A	No name enters exposure analysis

What changes, who gets paid and when

The strongest content-growth signals are inside the unchanged 72-GPU rack boundary. HBM bandwidth rises about 174%; NVLink rack bandwidth doubles; the switch-ASIC count doubles from 18 to 36; scale-out connectivity per GPU doubles; and host CPU memory rises from 17.28 TB to 54 TB. The likely beneficiaries are therefore the platform owner, qualified HBM4 and SOCAMM2 vendors, scale-up and scale-out silicon, high-speed interconnect components, power-smoothing hardware, liquid-cooling components and system integration. Only the named, disclosed relationships are attributed to companies; category demand is not converted into supplier revenue.

The near-term commercial window runs from memory qualification and system-builder production in 2026 through broad availability and customer deployment. The structural window is longer: hotter, more bandwidth-intensive racks pull facility electrical and cooling content forward and make networking a larger fraction of system value. The key risk for supplier modeling is cadence. Qualification can be economically meaningful before end-system volume, but an announcement date is not a revenue-recognition date.

Risks

- Published specifications are preliminary and can change before broad availability.
- No disclosed rack power means no rack count, power cost or active-IT capex per 100 MW.
- Foundry, package supplier, HBM allocations, system-builder shares and BlueField-4 count are undisclosed.
- A named design relationship does not reveal price, margin, duration or exclusive status.
- The five-rack Vera Rubin POD and one-rack NVL72 are different boundaries; blending them double-counts companion infrastructure.

China installed: Ascend 910C / Atlas 900 A3

Identity and lifecycle. Confirmed • shipping Ascend 910C is the accelerator; Atlas 900 A3 is the physical 16-cabinet SuperPoD; CloudMatrix384 is Huawei Cloud's service deployment of the architecture, not a separate hardware BOM. Huawei reported more than 300 A3 SuperPoDs shipped to more than 20 customers by September 2025 and more than 750 commercial deployments by 17 July 2026. Vendor-reported deployment counts establish commercial status but are not an independently audited installed-base series. Huawei's A3 product page (<https://e.huawei.com/cn/products/computing/ascend/atlas-900-a3-superpod>), September 2025 launch disclosure (<https://www.huawei.com/en/news/2025/9/hc-superpod-innovation>) and July 2026 update (<https://www.huawei.com/cn/news/2026/7/atlas-950-superpod>) define the boundary and status.

Package and system anatomy

The disclosed Ascend 910C package boundary is sparse: 128 GB of on-package memory per NPU and 3.2 TB/s of memory bandwidth. Huawei does not disclose HBM generation or supplier, die count, process node, packaging vendor, accelerator power or price. Reuters reported that Huawei sought SMIC N+2 production for the 910C; that relationship remains Reported, not confirmed allocation. Reuters reporting (<https://theprint.in/tech/exclusive-huawei-aims-to-mass-produce-newest-ai-chip-in-early-2025-despite-us-curbs/2366643/>) is therefore kept separate from Huawei's product specifications.

The complete A3 comprises 12 liquid-cooled compute cabinets and four air-cooled bus cabinets: 384 Ascend 910C NPUs, 192 Kunpeng 920 CPUs, 1,536 DDR5 DIMMs and 480 2.5-inch drives. It exposes 48 TB of aggregate on-package memory, 784 GB/s bidirectional die-to-die bandwidth and 288.7–307.2 PFLOPS of FP16 compute, depending on the published operating/configuration range. The input architecture is dual three-phase 380 V AC. Total and per-cabinet power are not disclosed.

Atlas 900 A3 boundary	Disclosed content	Evidence
One Ascend 910C	128 GB on-package memory; 3.2 TB/s memory bandwidth	Confirmed
Physical system	16 cabinets: 12 compute + 4 bus	Confirmed
Accelerators / CPUs	384 Ascend 910C / 192 Kunpeng 920	Confirmed
Host memory / storage	1,536 DDR5 DIMMs / 480 2.5-inch drives	Confirmed counts
Aggregate on-package memory	48 TB	Published aggregate; reconciles with 384 × 128 GB
FP16 compute	288.7–307.2 PFLOPS	Published range
Cooling	Liquid-cooled compute cabinets; air-cooled bus cabinets	Confirmed
System power	Not disclosed	N/A

Facility implication: a deployed system with an undisclosed denominator

The A3 is commercially deployed, but deployment status does not solve the normalization problem. Without total system power, a 16-cabinet footprint cannot be converted into systems per 100 MW. Nor can cabinet count stand in for power: four cabinets are bus cabinets, cooling modes differ, and the disclosed input voltage does not reveal load. The honest 100 MW result is **N/A**.

The national China facility case still describes the surrounding commissioned infrastructure. A base 100 MW IT-nameplate facility has 125 MW gross design demand at 1.25 PUE, consumes about 711.75 GWh annually at 65% utilization, and carries modeled facility-side capex of \$894.92 million. Those values do not determine how many A3 systems fit inside the compute allocation.

Supplier geography and evidence

Huawei and its HiSilicon design arm own the accelerator and system architecture; Kunpeng is also a Huawei processor family. The current architecture therefore concentrates visible design and platform value inside a private, inaccessible-to-most-investors corporate group. SMIC's foundry relationship is reported, not confirmed by a platform allocation. Huawei does not identify the memory, advanced-packaging, optics, storage, liquid-cooling or manufacturing suppliers in the cited product disclosures.

Role	Named entity	Relationship	Investment-access implication
Accelerator and system design	Huawei / HiSilicon	Confirmed	Huawei is private; direct public-equity access is unavailable
Host CPU	Huawei Kunpeng	Confirmed	Same concentrated private platform exposure
Logic fabrication	SMIC (0981.HK / 688981.SS)	Reported for 910C	Public access exists, but platform allocation is not confirmed
On-package memory	Not disclosed	N/A	No supplier attribution
Packaging, optics and cooling	Not disclosed	N/A	Category demand is visible; company allocation is not

What changes, who gets paid and when

The installed A3 base supports current demand for Huawei-designed accelerators, Kunpeng CPUs, large DDR5 and storage counts, bus-cabinet interconnect, liquid cooling and system integration. Yet the disclosed supplier map ends quickly. The most investable-looking external link—SMIC—is also the least certain named link in the platform record because it is reported rather than confirmed and carries no disclosed allocation.

For the next 12–36 months, the commercially relevant evidence is the reported deployment base and continued domestic substitution. For the 3–7 year horizon, the architecture demonstrates a Chinese path that compensates for weaker individual accelerators through system scale and a tightly controlled fabric. That

can enlarge domestic memory, optics, packaging and cooling demand, but it does not identify which listed supplier captures it. Supplier qualification, localization policy and Huawei procurement disclosures are more useful leading indicators than a top-down market-share assumption.

Risks

- Huawei-reported deployments are not an audited installed-base measure.
- The package process, memory supplier, packaging chain, system power and price are undisclosed.
- Reported SMIC participation cannot be treated as confirmed supplier allocation.
- U.S. Bureau of Industry and Security guidance specifically addressed Ascend 910C in May 2025, increasing policy and transaction risk; see BIS General Prohibition 10 guidance (<https://www.bis.gov/media/documents/general-prohibition-10-guidance-may-13-2025.pdf>) and §2.12.
- A3's 16-cabinet system cannot be compared directly with an NVL72 rack or Atlas 950 roadmap maximum.

China next: Ascend 950DT / Atlas 950 SuperPoD

Identity and lifecycle. **Roadmap with physical 1,024-card demonstration** Ascend 950DT is the Q4 2026 accelerator target. Atlas 950 is a family of system configurations. Huawei showed a physical **1,024-card** Atlas 950 SuperPoD at the World Artificial Intelligence Conference on 17 July 2026. Separately, it describes an **8,192-card, 160-cabinet roadmap maximum** targeted for Q4 2026. The demonstrated system and roadmap maximum are not interchangeable. As of the cutoff, the 1,024-card system was demonstrated; the 8,192-card maximum was not shipping. Huawei's WAIC disclosure (<https://www.huawei.com/cn/news/2026/7/atlas-950-superpod>) and 2025 roadmap keynote (<https://www.huawei.com/en/news/2025/9/hc-xu-keynote-speech>) anchor the distinction.

Package, demonstrated system and roadmap anatomy

Huawei's Ascend 950DT roadmap specifies 144 GB of HiZQ 2.0 HBM, 4 TB/s of memory bandwidth, 2 TB/s of interconnect bandwidth per chip, 1 PFLOPS FP8 and 2 PFLOPS FP4. Process node, die count, package supplier, HBM manufacturer, package power and price remain undisclosed.

The July 2026 physical demonstration used 1,024 accelerator cards and disclosed 1 EFLOPS FP8, 2 EFLOPS FP4, 256 TB of globally addressed memory, terabyte-class NPU interconnect and 3 μ s round-trip latency. The exact accelerator-card variant was not disclosed. The memory figure is an important reconciliation warning: 256 TB divided by 1,024 cards equals 256 GB per card, which does not match the separate 144 GB 950DT package roadmap. The demonstration is therefore stored as its own configuration rather than silently labeled "1,024 $\times$ 950DT."

The roadmap maximum is 8,192 accelerator cards across 160 cabinets and more than 1,000 square metres: 128 compute cabinets at 64 cards each plus 32 communications cabinets. Huawei describes an all-optical UnifiedBus fabric, 8 EFLOPS FP8, 16 EFLOPS FP4, 1,152 TB of aggregate memory and roughly 16.3 PB/s of interconnect bandwidth. Cooling, system power, CPUs, host memory, storage, exact optics counts and system price are not disclosed.

Atlas 950 boundary	Disclosed content	Status
Ascend 950DT package roadmap	144 GB HiZQ 2.0 HBM; 4 TB/s memory; 2 TB/s interconnect; 1 PF FP8; 2 PF FP4	Q4 2026 roadmap
Physical demonstration	1,024 cards; 1 EF FP8; 2 EF FP4; 256 TB global memory; 3 μ s RTT	Demonstrated 17 July 2026
Demonstration card identity	Not disclosed; 256 GB/card implied by aggregate memory	Do not label as 950DT
Roadmap maximum	8,192 cards; 160 cabinets; >1,000 m ² ; 1,152 TB memory	Q4 2026 target
Roadmap fabric	128 compute + 32 communications cabinets; all-optical UnifiedBus; ~16.3 PB/s	Roadmap
System power and cooling	Not disclosed	N/A

Facility implication: physical scale without power normalization

The roadmap's more-than-1,000-square-metre footprint and 160 cabinets show that the system is campus-scale relative to a rack, but floor area is not an energy denominator. Without accelerator-card, cabinet or total-system power, neither the 1,024-card demonstration nor the 8,192-card roadmap maximum can be converted into systems per 100 MW. A power estimate would propagate into cooling, electrical capex, annual energy and supplier content, multiplying one unsupported assumption across the whole model. All such outputs remain **N/A**.

Configuration	Physical evidence	100 MW result
1,024-card demonstration	Real hardware shown; exact card variant and total power undisclosed	N/A
8,192-card roadmap maximum	Cabinet, area, memory, compute and interconnect targets disclosed; not shipping	N/A
China base facility	\$894.92M facility-side capex; 125 MW gross design; 711.75 GWh/year at common base assumptions	Valid national benchmark, but cannot yield Atlas 950 system count

Supplier geography and evidence

Huawei owns Ascend design, HiZQ branding and UnifiedBus architecture. That establishes architectural control, not the manufacturing chain behind them. The HBM fabricator, foundry, advanced-packaging provider, optical transceiver and component suppliers, cooling vendors and final system manufacturers were not publicly allocated in the cited disclosures. No outside supplier is therefore assigned Atlas 950 revenue or share.

Role	Named entity	Evidence	What remains open
Accelerator and system design	Huawei / HiSilicon	 Roadmap owner	Manufacturing allocation
HBM architecture/brand	Huawei HiZQ 2.0	 Roadmap specification	Memory fabricator and packaging chain
Scale-up fabric	Huawei UnifiedBus	 Architecture owner confirmed	Actual optics, cable and switch-component suppliers
Foundry	Not disclosed	N/A	Process, location, yield and allocation
Cooling and system integration	Not disclosed	N/A	Vendors, topology, power and economics

What changes, who gets paid and when

Against the A3 package, the 950DT roadmap raises memory capacity from 128 to 144 GB, a 12.5% increase, and bandwidth from 3.2 to 4 TB/s, a 25% increase. At the published complete-system extremes, accelerators rise from 384 in A3 to 8,192 in the Atlas 950 roadmap maximum, cabinet count rises from 16 to 160, and aggregate memory rises from 48 to 1,152 TB. These 21.33×, 10× and 24× ratios describe disclosed configuration boundaries; they are not a shipment forecast.

The demand direction is toward much more accelerator silicon, HBM capacity, optical scale-up connectivity, power conversion and liquid cooling per deployed maximum-size system. “Who gets paid” is unresolved beyond Huawei because external suppliers are unnamed. The near-term catalyst is qualification and actual delivery against the Q4 2026 target; the structural question is whether UnifiedBus and domestic memory/packaging capacity can scale with acceptable yield, power and cost. Investors should require evidence of purchase orders, supplier qualification or recognized revenue before mapping category growth to a listed vendor.

Risks

- The 1,024-card demonstration and 8,192-card roadmap maximum are different systems; merging them creates a false shipping claim.
- The demonstration’s implied 256 GB per card does not reconcile with the 950DT’s separately disclosed 144 GB, so the card identity remains unknown.
- System power, cooling topology, CPU and storage counts, foundry, HBM supplier, packaging provider and optics allocation are undisclosed.
- Q4 2026 is a roadmap target, not evidence of volume shipment.
- Export controls may slow domestic manufacturing inputs while simultaneously strengthening localization demand; see §§2.9, 2.10 and 2.12.

What changes from current to next

The generation delta is more useful than a static BOM because it identifies where addressable content is growing before supplier shares are known. It also exposes the asymmetry between the ecosystems. NVIDIA keeps the accelerator count fixed and

intensifies memory, fabric, host memory, cooling and power quality inside one rack. Huawei's disclosed next step expands both the package and the maximum system boundary, but leaves more of the manufacturing and power chain undisclosed.

Current-to-next component and evidence changes

Comparable physical quantities and architecture changes only; no dollar content or supplier allocation inferred

Domain	Current	Next	Physical delta	Evidence
US GPU count / rack boundary	72 Blackwell Ultra GPUs	72 Rubin GPUs	Flat count; architecture changes	● Confirmed
US On-package memory	288 GB HBM3E; 8 TB/s per GPU	288 GB HBM4; 22 TB/s per GPU	Capacity flat; bandwidth/interface rises	● Confirmed; Rubin preliminary
US Scale-up switch silicon	18 NVSwitch ASICs	36 NVLink 6 switch ASICs	Count doubles at same compute-rack boundary	● Confirmed
US Host memory / rack	17.28 TB LPDDR5X	54 TB LPDDR5X SOCAMM2	Capacity rises 212.5%	● Confirmed; Rubin preliminary
CHINA On-package memory	128 GB; generation undisclosed	144 GB HiZQ 2.0 HBM	Roadmap capacity rises 12.5%	□ Current confirmed; next roadmap
CHINA Maximum system accelerator count	384 NPUs across 16 cabinets	8,192 cards across 160 cabinets	System boundary expands 21.3×	□ Current confirmed; next roadmap
CHINA System interconnect cabinets	4 bus-equipment cabinets	32 communications cabinets; all-optical	Architecture and physical scope change	□ Current confirmed; next roadmap
CHINA Complete-system power	Not publicly disclosed	Not publicly disclosed	No valid 100 MW system-count delta	○ Evidence gap

Sources: NVIDIA and Huawei disclosures. Analytical synthesis; counts are compared only at the stated rack or full-system boundary.

Ecosystem metric	Current	Next	Delta	Evidence and comparability
U.S. GPUs per NVL72	72	72	0%	Same one-rack boundary
U.S. GPU memory per rack	20.736 TB	20.736 TB	0%	Same capacity; HBM3E → HBM4
U.S. HBM bandwidth per rack	576 TB/s	~1,580 TB/s	+174.3%	Derived from disclosed per-GPU values; next rounded
U.S. NVLink bandwidth per rack	130 TB/s	260 TB/s	+100%	Platform comparison
U.S. scale-out bandwidth	0.8 Tb/s/GPU	1.6 Tb/s/GPU	+100%	Platform comparison
U.S. NVLink switch ASICs	18	36	+100%	Current disclosed; next disclosed count
U.S. CPU memory	17.28 TB	54 TB	+212.5%	Same rack boundary
U.S. rack power	Up to 142 kW	N/A	N/A	No Rubin rack-power disclosure
China package memory	128 GB	144 GB	+12.5%	910C actual vs 950DT roadmap
China package memory bandwidth	3.2 TB/s	4 TB/s	+25%	910C actual vs 950DT roadmap
China accelerators per compared system	384	8,192	21.33×	A3 shipping vs Atlas 950 roadmap maximum; not a unit-sales forecast
China cabinet count	16	160	10×	Same caution: current vs roadmap maximum
China aggregate accelerator memory	48 TB	1,152 TB	24×	Current vs roadmap maximum
China system power	N/A	N/A	N/A	No normalized comparison

Three investment readings follow. First, **capacity can stay flat while content rises**: Rubin's unchanged GPU count masks large increases in bandwidth, switching and host memory. Second, **system scale is itself a strategy**: Huawei's roadmap uses many more cards and an all-optical fabric to create a competitive aggregate system. Third, **disclosure quality determines model quality**: the apparently larger China delta cannot be converted into facility demand or supplier revenue until power and allocations are published.

The normalized 100 MW facility

The common facility boundary is 100 MW of commissioned IT nameplate. Low, base and high scenarios use 1.15, 1.25 and 1.35 PUE and 50%, 65% and 80% utilization. Gross design power is therefore 115, 125 and 135 MW, while annual energy is 503.70, 711.75 and 946.08 GWh. The physical ranges are held constant across the U.S. and China to avoid embedding an unsupported efficiency advantage. National differences enter through construction, electricity and water costs.

Installed facility capex per 100 MW of commissioned IT nameplate

US\$ million, low/base/high national screening ranges; active IT, land, taxes and financing excluded



Sources: Turner & Townsend; Cushman & Wakefield; atlas facility model. Intervals are scenario ranges, not confidence intervals.

100 MW IT-nameplate metric	United States low / base / high	China low / base / high	Interpretation
Installed facility-side capex	\$1,118.15M / \$1,371.87M / \$1,814.12M	\$660.71M / \$894.92M / \$1,174.05M	Excludes land, active IT, taxes and financing
Gross design power	115 / 125 / 135 MW	115 / 125 / 135 MW	IT nameplate × PUE
Annual site energy	503.70 / 711.75 / 946.08 GWh	503.70 / 711.75 / 946.08 GWh	Gross power × utilization × 8,760 hours
Electricity	\$25.19M / \$61.35M / \$132.45M	\$25.54M / \$60.15M / \$106.60M	National/regional price sensitivities
Direct water volume	87,600 / 256,230 / 490,560 m ³	Same physical range	WUE 0.2 / 0.45 / 0.7 L/kWh of IT energy
Water and wastewater	\$0.18M / \$0.90M / \$3.92M	\$0.03M / \$0.22M / \$0.66M	Cost is small relative to power; availability can still bind
Facility maintenance	\$22.36M / \$54.88M / \$108.85M	\$13.21M / \$35.80M / \$70.44M	2% / 4% / 6% of installed facility capex
Normalized systems	GB300 only: 492.96 / 563.38 / 619.72 racks across compute-share sensitivity	N/A	Other platform power is undisclosed

The base facility capex is a leaf-only sum, preventing the common error of adding subtotals to their children.

Base facility-side capex leaf	United States	China
Electrical plant	\$567.672M	\$370.312M
Cooling plant	\$390.275M	\$254.589M
Shell and architectural	\$106.439M	\$69.433M
Contractor preliminaries and fees	\$118.265M	\$77.148M
Utility extension allowance	\$94.612M	\$61.719M
Owner, commissioning and facility-spares allowance	\$94.612M	\$61.719M
Total	\$1,371.874M	\$894.920M

The cost gap is a benchmark-geography result, not evidence that a U.S. or Chinese facility is operationally superior. Construction benchmarks are especially sensitive to region, labor market, utility scope and procurement timing. The model uses the Turner & Townsend 2025 cost-index methodology (<https://reports.turnerandtownsend.com/data-centre-construction-cost-index-2025/methodology>) and liquid-cooling allocation, and the Cushman & Wakefield APAC construction guide (<https://cushwake.cld.bz/asiapacificdatacentreconstructioncostguide-01-2025-apac-regional-en-content-datacentres>) for the China range. PUE and utilization bounds are anchored to LBNL's U.S. data-center energy report (<https://eta.lbl.gov/publications/2024-lbnl-data-center-energy-usage-report>), Chinese national efficiency disclosures and Uptime Institute's giant-data-center analysis (https://intelligence.uptimeinstitute.com/sites/default/files/2026-01/UI%20Field%20report%20194_Giant%20data%20center%20power%20plans%20reach%20extreme%20levels.pdf).

From qualification to revenue: timing is not the same as participation

Supplier evidence arrives in a sequence. Architecture announcements create category demand; qualification or design-in identifies a potential supplier; manufacturing start establishes physical participation; shipment establishes a commercial product; customer commissioning creates utilization and recurring facility opex; reported revenue finally establishes economic capture. Most public supply-chain claims stop somewhere before the final step.

Supplier qualification-to-revenue evidence sequence

Analytical gate sequence, not a calendar-duration forecast; evidence cutoff 2026-07-25



STAGE	EVIDENCE GATE	REVENUE READING	EXAMPLE
2. Supplier alignment	Supplier names a designed-for or co-developed product	Addressability, not allocation	HBM4 / SOCAMM2 designed for Vera Rubin
3. Qualification / allocation	Qualified SKU and customer allocation are disclosed	Allocation required for content estimates	No reviewed source discloses supplier shares
5. Customer deployment	Shipping, available-now, or deployed system evidence	Installed-base evidence, not supplier mix	GB300 available; Atlas 900 A3 deployed commercially

Sources: NVIDIA, Huawei, memory suppliers and Wistron. Sequence is an atlas synthesis; stage duration is not implied.

Evidence stage	U.S. current / next examples	China current / next examples	What an investor may conclude
Architecture / roadmap	Rubin HBM4, NVLink 6, warm-water cooling; H2 2026 availability target	950DT and 8,192-card Atlas 950 Q4 2026 targets	Directional category demand only
Design-in / qualification	Micron and Samsung describe Vera Rubin memory designs; SK hynix relationship reported	HiZQ and UnifiedBus are Huawei architectures; external suppliers unnamed	Named participation only where explicitly disclosed
Manufacturing	GB300 Wistron Fort Worth production; Vera Rubin system builders named	A3 commercially deployed; Atlas 950 physical 1,024-card demonstration	Product exists or is entering build; no supplier share
Shipment / availability	GB300 shipping; Rubin ramping toward broad availability	A3 shipping; 8,192-card Atlas 950 not shipping	Revenue window may be open, but value and allocation remain unknown
Commissioning / utilization	Customer-specific deployment not modeled	>750 A3 deployments reported by Huawei	Installed status; no inference to useful throughput
Recognized supplier revenue	Not provided at platform-line level in this dataset	Not provided at platform-line level in this dataset	Requires filings or supplier disclosure; do not back-solve from BOM counts

The timing discipline prevents two errors. A roadmap designation is not backlog, and a confirmed design-in is not a revenue-share estimate. For the 12–36 month horizon, watch component qualification, manufacturing starts, broad availability, system power disclosure and commissioned deployments. For the 3–7 year horizon, watch whether memory, networking and facility content keep rising faster than accelerator counts, and whether domestic Chinese suppliers become visible enough to separate category growth from economic capture.

Investability without a buy score

The matrix below is a screening tool, not a recommendation. “Access” asks whether an investor can obtain direct public-equity exposure. “Evidence” asks whether the platform relationship is documented. “Concentration” asks how dependent the exposure is on one customer or platform. “Substitution” asks how readily another qualified supplier or architecture could take the content. None of these dimensions should be averaged into an opaque composite.

Investability matrix: exposure without a composite score

Content direction, access, evidence and separate risk dimensions; customer concentration stays N/A when undisclosed

Guardrail: confirmed exposure is not supplier share, revenue attribution or an investment recommendation.

ENTERPRISE	CONTENT DIRECTION	MARKET ACCESS	EVIDENCE	CUSTOMER CONC.	SUBST.	CAPACITY	PLATFORM TIMING
Hon Hai / Foxconn	current → next · 4 roles	direct · TWSE: 2317	confirmed / high	not disclosed at platform-line level	HIGH	MEDIUM	shipping → ramping shipping/ramping = solid; announced = dotted
Huawei Technologies / HiSilicon	current → next · 3 roles	non investable directly · no direct equity	confirmed / high	not disclosed at platform-line level	LOW	HIGH	shipping → announced shipping/ramping = solid; announced = dotted
Micron Technology	current → next · 5 roles	direct · NASDAQ: MU	confirmed / high	not disclosed at platform-line level	MEDIUM	HIGH	shipping → ramping shipping/ramping = solid; announced = dotted
NVIDIA	current → next · 4 roles	direct · NASDAQ: NVDA	confirmed / high	not disclosed at platform-line level	MEDIUM	HIGH	shipping → ramping shipping/ramping = solid; announced = dotted
SK hynix	current → next · 2 roles	direct · KRX: 000660	confirmed and reported / medium	not disclosed at platform-line level	MEDIUM	HIGH	shipping → ramping shipping/ramping = solid; announced = dotted
SMIC	current only · 1 roles	direct with restrictions · SSE STAR: 688981	reported / medium	not disclosed at platform-line level	LOW	HIGH	shipping shipping/ramping = solid; announced = dotted
Samsung Electronics	next only · 2 roles	direct · KRX: 005930	confirmed / high	not disclosed at platform-line level	MEDIUM	HIGH	ramping shipping/ramping = solid; announced = dotted
TSMC	current only · 1 roles	direct · TWSE: 2330	confirmed / high	not disclosed at platform-line level	LOW	HIGH	shipping shipping/ramping = solid; announced = dotted
Wistron	current → next · 2 roles	direct · TWSE: 3231	confirmed / high	not disclosed at platform-line level	HIGH	MEDIUM	shipping → ramping shipping/ramping = solid; announced = dotted

Evidence: blue = confirmed; orange = reported. Risk labels remain separate and must not be summed.

Sources: platform vendors and named suppliers; access from Part V. Analytical synthesis; N/A concentration cells are an evidence gap, not a low-risk label.

Exposure	12–36 month catalyst	3–7 year structural position	Public-market access	Platform evidence	Concentration / substitution risk
NVIDIA platform silicon, NVLink and networking	GB300 shipments; Rubin broad availability and production ramp	Integrated rack architecture can capture a larger system BOM	Direct: NVDA	High for both U.S. anchors	High platform concentration; architectural substitution is difficult inside NVL72 but customer alternatives exist
U.S. HBM and SOCAMM — Micron	GB300 HBM3E/SOCAMM M and Rubin HBM4/SOCAMM 2 design-ins	Memory bandwidth and host-memory content grow faster than GPU count	Direct: MU	High for named design-ins; allocation unknown	Cyclical memory pricing; Samsung and SK hynix substitution
Korean HBM — SK hynix and Samsung	GB300 module evidence; Rubin relationships and designs	HBM4 complexity, bandwidth and capacity raise value per system	Direct in Korea	Mixed high/reported; allocation unknown	Multi-sourcing, qualification timing and memory-cycle risk
TSMC leading-edge fabrication	Confirmed GB300 4NP demand	Advanced logic and packaging remain difficult to substitute	Direct: TSM / 2330.TW	High for GB300; Rubin foundry undisclosed	Extreme customer and Taiwan concentration; process leadership lowers substitution

Exposure	12–36 month catalyst	3–7 year structural position	Public-market access	Platform evidence	Concentration / substitution risk
U.S./Taiwan system builders	GB300 production and Rubin manufacturing ramp	More system-level integration, cooling and power content	Direct through named listed builders	Participation high; unit shares unknown	Low transparency and potentially substitutable assembly capacity
Facility electrical and cooling categories	100 MW projects require about \$958M of electrical plus cooling plant in the U.S. base case; \$625M in China	Denser liquid-cooled systems deepen grid-to-chip content	Access through Part V category screens	Provider-neutral facility evidence; no platform supplier attribution	Project timing, regional pricing, customer concentration and vendor competition
Huawei / HiSilicon platform	A3 installed base; Atlas 950 delivery against roadmap	Domestic full-stack control and UnifiedBus ecosystem	No direct public equity; Huawei private	High for architecture, low for external allocation	Extreme single-platform concentration; policy support and policy risk coexist
SMIC reported 910C fabrication	Evidence of continued current-generation production would be the catalyst	Domestic leading-edge substitution has strategic value	Direct: 0981.HK / 688981.SS	Reported, not confirmed allocation	Yield, tools, export controls and customer concentration
China memory, packaging, optics and cooling categories	Supplier qualification or purchase-order disclosure	Atlas-scale systems could expand domestic content sharply	Potentially accessible, but no issuer is assigned here	Low / undisclosed	High attribution risk; do not equate category demand with issuer revenue

The highest-evidence U.S. exposures are also the most obvious and often the most concentrated. The largest apparent China content delta has the weakest external-supplier disclosure and the least direct access. That is not a reason to ignore it; it is a reason to demand a higher evidence threshold. Part V contains the canonical company-level profiles and should be used to test valuation, liquidity, listing access and company-specific risks after this physical screen.

What to ask next

The following disclosures would materially improve the model:

- What is Vera Rubin NVL72's maximum and typical rack input power, and what boundary does each figure use?
- Which foundry process and advanced-packaging flow produce Rubin, and what is the geographic allocation?
- What are the Rubin HBM4 allocations among Micron, Samsung and SK hynix, and when do qualified shipments become recognized platform revenue?
- How many BlueField-4 devices, power shelves and local energy-buffering modules are in one Vera Rubin NVL72?
- What are Atlas 900 A3 total and per-cabinet power, cooling flow, optics counts and supplier allocations?
- Which accelerator card was used in the 1,024-card Atlas 950 demonstration, and why does its implied 256 GB per card differ from the 144 GB 950DT roadmap?
- What is the power, cooling topology and active component count of the 8,192-card Atlas 950 maximum?
- Who fabricates and packages Ascend 950DT and HiZQ 2.0 memory, and which optical suppliers qualify for UnifiedBus?
- At the facility layer, how much regional variation sits behind national construction and electricity cases, and which costs are customer-direct rather than contractor-delivered?
- How do useful training and inference throughput per commissioned megawatt compare once software, utilization and model workload are held constant?

The last question is deliberately outside v1. A common 100 MW denominator normalizes physical occupancy and facility economics; it does not normalize useful compute, model quality or revenue.

Caveats and assumptions

- **Cutoff and lifecycle:** evidence is frozen at 25 July 2026. Status labels describe that date, not a timeless product state.
- **System boundaries:** GB300 and Vera Rubin use one NVL72 compute rack. Atlas 900 A3 uses one 16-cabinet SuperPoD. Atlas 950 retains separate 1,024-card demonstrated and 8,192-card roadmap-maximum configurations.
- **Evidence policy:** confirmed, reported and transparently inferred values may appear in factual tables. Roadmap values appear only with a roadmap label. Speculative values never enter totals, deltas or the investability matrix.
- **Ranges:** low/base/high values express scenario uncertainty, not statistical confidence intervals.
- **Facility scope:** included capex runs from site transformer/substation boundary through electrical distribution, facility-side cooling and heat rejection, shell/core, contractor costs, utility extension allowance, owner costs, commissioning and facility spares.
- **Excluded from facility capex:** land, taxes, financing, remote grid works, active accelerators, CPUs, memory, storage, rack networking, optics, rack cold plates, active-IT spares and project revenue.
- **Recurring scope:** electricity, direct water and routine facility maintenance are modeled separately. Active-IT service contracts, depreciation and major mid-life renewal are excluded.
- **Physical assumptions:** compute share is 70%/80%/88%; PUE is 1.15/1.25/1.35; utilization is 50%/65%/80%; direct WUE is 0.2/0.45/0.7 litres per kWh of IT energy. Critical-support capacity ratios of 1.0×/1.1×/2.0× describe design sensitivity and are not multiplied into capex a second time.
- **Currencies:** facility outputs are presented in U.S. dollars. China electricity and construction inputs are translated using the documented model FX convention; they are scenarios, not forecasts of exchange rates.
- **Supplier allocation:** a named participant, design-in or manufacturing site does not imply exclusivity, unit share, price or margin.
- **Investment use:** this is supply-chain and investment research, not personalized investment advice.

Methods and source note

Markdown is the narrative source; the versioned supply-chain dataset is the quantitative source of truth. Each displayed value maps to a platform state, system configuration, BOM line, facility assumption, supplier relationship, claim and source. Derived fields retain their formulas and dependencies. Facility totals are recomputed from leaf rows only; capex and opex remain separate; system counts require disclosed system power; supplier shares remain null when undisclosed.

The publication snapshot uses methodology version 1.0, evidence cutoff 25 July 2026 and immutable snapshot hash `3647c9....`. The complete checksum is stored with the published supply-chain manifest rather than retyped here. Core first-party anchors are NVIDIA's GB300 and Vera Rubin product and architecture disclosures, Huawei's Atlas 900 A3 and Atlas 950 disclosures, DOE's facility design guide, LBNL and Chinese national efficiency sources, Turner & Townsend and Cushman & Wakefield construction benchmarks, EIA and Chinese electricity sources, EPA and municipal water tariffs, and GAO/NASA maintenance benchmarks. Reuters is used only where the relationship is explicitly labeled reported.

This chapter is the cross-layer capstone, not a replacement for the underlying layers. Use §2.3 for the power and cooling market, §2.5 for accelerator competition, §2.6 for networking, §2.7 for memory and packaging, §2.8 for foundry geography, §2.9 for production tools and EDA, §2.10 for materials, §2.11 for financing and §2.12 for policy. Use Part V for canonical company access, fundamentals and risk profiles.

Chapter endnotes

[^factory-nvidia-gb300]: NVIDIA GB300 NVL72 (<https://www.nvidia.com/en-us/data-center/gb300-nvl72/>). NVIDIA, undated; accessed 2026-07-25.

[^factory-nvidia-rubin]: NVIDIA Vera Rubin NVL72 (<https://www.nvidia.com/en-us/data-center/vera-rubin-nvl72/>). NVIDIA, undated; accessed 2026-07-25.

[^factory-huawei-a3]: Atlas 900 A3 SuperPoD product specifications (<https://e.huawei.com/cn/products/computing/ascend/atlas-900-a3-superpod>). Huawei Enterprise, undated; accessed 2026-07-25.

[^factory-huawei-950]: Huawei's SuperPoD Portfolio Creates New Option for Global Computing at MWC Barcelona 2026 (<https://www.huawei.com/en/news/2026/3/mwc-superpod>

computing). Huawei, 2026-02-28; accessed 2026-07-25.

Part III — Systems in the Round

Part II examined national advantage inside each layer. This part steps back to compare four interacting systems — the United States, China, the allied chokepoint economies, and the third-bloc demand markets — including their binding constraints, capital structures, strengths and vulnerabilities. It is the system view after thirteen chapters of component analysis.

3.1 The United States as a system

The shape of American power

Read from a distance, the American position in AI has a clear and consistent shape: it leads the top of the stack and has offshored the bottom, and its binding constraint has moved from technology to physics and finance. The United States owns the layers where value and margin concentrate — chip design, cloud, models, applications, and the software that ties them together — and it owns the capital markets that fund the whole build. What it does not own is the making. The leading-edge fabrication sits in Taiwan, the one lithography machine that makes it possible sits in the Netherlands, the memory sits in Korea, and the raw materials trace back to China. America leads the parts of the industry that can be done anywhere and has rented out the parts that can only be done in a few places. That is a comfortable position right up until one of those places is disrupted.

Where it leads, and where it is exposed

The pattern is worth making precise, layer by layer, because the American strength is uneven. At the very top — models and applications — the US leads decisively, holding the closed frontier (Chapter 2.1) and the enterprise-monetization winners (Chapter 2.2), though even here China has quietly taken the open-weight and developer-token layer. In

cloud and the CUDA software moat (Chapter 2.4), the lead is close to total. In AI chips (Chapter 2.5), interconnect (Chapter 2.6), and the design tools and EDA (Chapter 2.9), American companies — Nvidia, Broadcom, AMD, Synopsys, Cadence — are dominant. But descend to fabrication (Chapter 2.8), memory (Chapter 2.7), lithography (Chapter 2.9), and materials (Chapter 2.10), and the American name disappears, replaced by TSMC, ASML, the Korean memory makers, and the Japanese materials houses. And at the level of raw power (Chapter 2.3), the US is not a leader but a laggard, constrained where China is abundant. The stack diagram of §1.2 captures it: America is strong at the top and the toolchain, absent in the neutral-chokepoint middle, and weak at the physical base.

The constraint that changed: from chips to electrons

The American system's binding constraint in 2026 has broadened from chips to available, deliverable electricity. The country can design and buy more compute than it can power. Grid-interconnection queues stretch past eight years in the busiest regions, capacity prices have risen roughly tenfold in two years, and the transformers and gas turbines needed to energize a site carry multi-year waits (Chapter 2.3). This is a physical constraint that money alone cannot quickly solve: it is the consequence of long-running under-investment in grid capacity meeting a sudden, enormous new load. The country that leads the world in designing AI compute cannot, for the moment, plug all of it in.

The financial fragility: reflexive and levered

Layered on top of the physical constraint is a financial one. The build is being funded increasingly by debt and by circular arrangements in which the industry finances its own demand (Chapter 2.11). Capital spending now exceeds the operating cash flow of the companies doing it; the largest names issued record bonds and turned to private credit and GPU-backed securitization; and the stock market that owns all of it has concentrated to a degree not seen in a generation, with the top ten companies about 40% of the index. Both the BIS and the IMF have named this a financial-stability concern. The American system, in other words, has moved its risk from "will the technology work" — a question it has largely answered — to "will the financing hold," which it has not. This is a distinctly American fragility, because the circular loops, the private megarounds, the crowded mega-cap trade, and the levered credit are all features of the US capital market, not the Chinese one.

The two dependencies it cannot diversify away

Two exposures sit outside the American system's control and cannot be hedged, only sized. The first is Taiwan, where the leading-edge manufacturing that the entire build depends on is concentrated ninety miles from China; TSMC's Arizona diversification is real but back-loaded past 2030, and a Strait disruption would halt the whole complex at once (Chapters 2.8, 2.12). The second is China-controlled materials, the rare earths, gallium, and graphite that Beijing has repeatedly weaponized and that the West is years from replacing (Chapter 2.10). America's answer to the first is chip diplomacy and slow reshoring; its answer to the second is a policy-driven scramble to stand up MP Materials and Lynas. Neither answer removes the dependency this decade.

The policy posture: from denial to deal-making

The American state has shifted from a rules-based posture of denying China chips to a transactional one of selling and taxing them (Chapter 2.12): a 15% revenue-share tax on exports, a 25% tariff on advanced chips, a roughly 10% equity stake in Intel, and government-to-government GPU deals that route Nvidia hardware to sovereign-AI projects in the Gulf. The logic is to keep the world hooked on the American stack, offset the lost China revenue with sovereign demand, and use market access as leverage. It makes the industry more headline-sensitive and more politically entangled than its fundamentals alone would suggest.

The American balance sheet

The United States as a system is therefore technologically dominant, physically power-constrained, financially reflexive, and geopolitically exposed at two undiversifiable points. Its strengths are the deepest capital markets on earth, the best chip and model designers, the CUDA and cloud moats, and the ability to attract the world's compute demand. Its vulnerabilities are the grid, the leverage in the financing, and the Taiwan-and-materials dependencies. For an investor, owning "America's AI" means owning the top-of-stack designers and the software moats for their quality, owning the domestic power-and-electrical complex for the constraint, and respecting that the whole edifice rests on a Taiwanese foundation and a Chinese mineral base that no American company controls.

3.2 China as a system

The shape of Chinese power

China is the mirror image. It was denied the top of the stack and is being forced to build the base it was cut off from, and its constraint is the reverse of America's: it has the power and the policy will but lacks the leading-edge manufacturing and memory. Where the American system rents the base it cannot make, the Chinese system is building the base it was refused, at subsidized cost and behind a wall of mandates, while remaining throttled at exactly two points.

The mandate flywheel: "autonomous and controllable" localization

The defining feature of the Chinese system is that its AI demand is manufactured by the state rather than won in the market. Beijing has barred foreign accelerators from state-funded data centers, required domestic-chip quotas, and run a localization program — *autonomous and controllable* (自主可控, zìzhǔ kěkòng) delivered through the *IT-application-innovation* procurement regime (信创, Xinchuang) — that makes buying Chinese silicon a matter of law for government, state-enterprise, finance, telecom, and military buyers. Inclusion on the procurement catalogue carries roughly a 20% price preference in tenders and, by 2027, central enterprises are meant to complete a near-total replacement of foreign IT. This creates a protected, guaranteed demand pool for the national champions — Huawei and Cambricon above all — regardless of whether their chips can compete on the open market. It is a powerful flywheel: mandated demand funds domestic capacity, which improves the product, which justifies more mandates. The full mechanism is the subject of a dedicated deep-dive.

The "domestic-content" (国产) hierarchy: real versus nominal

A crucial subtlety that the headline localization numbers hide is that not all "domestic" (国产, guóchǎn) is equal, and the strictest state buyers know it. There is a hierarchy. At the top sit the truly self-controlled designs that own their instruction-set architecture and owe no foreign license: Loongson's LoongArch, and the self-designed accelerator architectures of Huawei's Ascend and Cambricon. Below them sit the chips that are "domestic" only in the sense of being made by a Chinese company while still resting on a foreign instruction set — Hygon's x86 (from a frozen AMD license) and the ARM-based designs of Phytium and Huawei's Kunpeng. For mainstream government procurement

these licensed parts qualify, but for the defence and secrecy tiers they are discounted as not "truly" self-controllable, and the 2026 catalogue explicitly set out to purge "pseudo-Xinhuang" (伪信创) products. For an investor the lesson is to ask *how* domestic a champion really is, because the market sometimes prices "domestic" (国产) exposure without distinguishing the real from the nominal.

The two ceilings: fabrication and memory

The Chinese flywheel spins against two hard limits that policy cannot wish away. The first is fabrication: SMIC produces 7nm-class chips using older deep-ultraviolet tools and multi-patterning, but at yields estimated in the 20–30% range, and pushing below 7nm on purely domestic equipment is judged unlikely before the end of the decade (Chapter 2.8). The second is high-bandwidth memory: China trails badly in HBM specifically, and Huawei's push to build its own is an attempt to close a gap it has not yet closed (Chapter 2.7). These two ceilings are the reason a mandated market is not the same as a competitive one: Ascend and Cambricon can only ship in the volumes SMIC can yield and the memory China can supply, and their headline performance claims rest on using several times the chips and power of a Western rack. Everything else in the Chinese stack is strong or fast-closing; these two points are where it is genuinely stuck.

The one clear advantage: power and "east-data, west-compute"

Where the American system is power-constrained, the Chinese one is power-abundant, and this is China's single clearest structural advantage in the whole contest (Chapter 2.3). China carries well over twice the installed generating capacity of the United States, adds more new capacity in a year than most countries operate in total, and can build transmission on a timescale American permitting cannot match. Its "east-data, west-compute" program (东数西算) deliberately routes data centers to the western interior where hydro, wind, and coal are cheap, and runs the fiber back east. The thing throttling the American build is close to a non-issue for the Chinese one, which is why the power trade is an American trade and the Chinese AI thesis lives in compute and localization instead.

The access maze: what a foreign investor can and cannot own

The Chinese system poses a problem the American one does not: much of it is difficult or impossible for a foreign investor to own, and the "market cap" of some of it does not

mean what it appears to. The categories matter.

Access category	Examples	What it means
Freely accessible	Alibaba (BABA/9988.HK), Tencent (0700.HK), SMIC (0981.HK)	own via HK or ADR listing
Connect-eligible, institutional-only	Cambricon, Hygon, Naura, AMEC, Montage	reachable via Stock Connect, but STAR names are professional-investor-only
Onshore-only / not yet Connect	CXMT (new STAR listing)	hard for offshore investors until index inclusion
Private / inaccessible	Huawei, DeepSeek, Biren, Moore Threads, ByteDance	the crown jewels cannot be bought at all
Barred to US persons (NS-CMIC)	Hikvision, SenseTime	US persons may not own the security

The practical consequence is that a foreign investor's China-AI universe is far smaller than the industry itself. The best assets (Huawei, DeepSeek) are un-buyable; the mandated chip champions are STAR-listed and often institutional-only; and one sanction, NS-CMIC, bars ownership outright for US persons. Screen for access before thesis.

The red flags: reading a Chinese AI stock

Beyond access, the Chinese system carries a set of company-level risks that the American one largely does not, and reading them is a skill of its own. A-shares are heavily retail-driven, so many thematic names trade as 概念股 — "concept stocks" riding a story with little real exposure, at extreme multiples that reflect positioning, not fundamentals. Watch for pledged shares (股权质押), where founders have borrowed against their stock and can be forced to sell; for government subsidies (政府补助) inflating reported profit, which means reading the non-GAAP (扣非) line; for the overhang of the national "Big Fund" selling down its stakes (大基金减持); and for fabricated order rumors (小作文) that move retail-driven names. None of these has an American analog of the same character. The Chinese system's risks are political and structural where the American system's are financial and concentrated.

The Chinese balance sheet

China as a system is therefore power-abundant, demand-guaranteed, and fast-closing at every layer except the two where it is genuinely stuck — leading-edge fabrication and high-bandwidth memory. Its strengths are the state's ability to manufacture demand, fund capacity, and build power at will, plus a genuinely competitive open-weight model layer and an increasingly domestic toolchain. Its vulnerabilities are the two hard ceilings, a dependence on the very foreign tools and memory it is trying to replace, and a capital structure that concentrates risk in state policy and retail sentiment. For an investor, owning "China's AI" means owning the mandated localization winners for the guaranteed demand, accepting the access frictions, distinguishing real domestic from nominal, and pricing in a policy environment that can change everything overnight.

3.3 The allied chokepoint system

The United States and China are not self-contained systems. Between them sits an allied manufacturing system whose members possess the least-substitutable capabilities in the stack: Taiwan in leading-edge foundry and advanced packaging, the Netherlands in EUV lithography, Korea in HBM, and Japan in substrates, wafers, photoresist and precision materials. Calling these countries “neutral chokepoints” understates their agency. Their companies allocate capacity, qualify customers, comply with national policy, choose where to build the next factory and determine how quickly either superpower can reduce its dependence.

The system's strength comes from accumulated process knowledge rather than raw scale alone. TSMC's yield learning, ASML's installed base and service network, Korean memory process integration, and Japanese materials qualification are the product of decades of customer feedback and manufacturing iteration. A subsidy can finance a new factory; it cannot instantly reproduce that learning curve. The resulting scarcity generally has a longer half-life than assembly capacity or model leadership.^[^systems-tsmc-hpc]
^[^systems-asml-products] ^[^systems-skhynix-hbm4]

Its weakness is geographic and political concentration. Taiwan remains the largest common exposure. Korea sits within reach of regional military escalation. ASML and

Japanese suppliers must translate US controls into their own commercial policy while retaining economically meaningful China businesses. These companies can benefit from both stacks' demand, but they are not politically detached toll roads.

For equity investors, the allied system requires two separate judgments. The first is **strategic indispensability**. The second is **security-level capture**: valuation, China revenue, parent-company dilution, currency, cyclicity and location risk. The atlas therefore treats allied chokepoints as high-quality strategic assets, not automatically as low-risk securities.

3.4 The third-bloc demand system

The countries outside the two manufacturing cores form a demand and capital system of their own. Gulf sovereign funds can finance multi-gigawatt campuses and use procurement to secure access to the US-led stack. India combines sovereign-compute policy, a large developer base and a growing data-center market. Southeast Asia offers land, power and regional cloud capacity while navigating US technology controls and Chinese commercial ties. Europe supplies industrial technology and regulation but is constrained by power, permitting and a fragmented capital market. Israel contributes an unusually dense design, networking, cybersecurity and application-startup ecosystem.

This is not a single political bloc. Its members optimize among access to frontier chips, financing, energy cost, data sovereignty, export-control exposure and local industrial development. Their choices determine standards and installed base. A country that builds on CUDA, US cloud and Western networking reinforces the American ecosystem; one that adopts subsidized Chinese systems creates a service, software and replacement market for the Chinese stack.

The Gulf is the clearest capital-led case. Saudi Arabia's HUMAIN and related sovereign partnerships are designed to combine power, capital, models and data-center construction at national scale.^[^systems-humain] India and Southeast Asia are more price-sensitive: they can become important demand pools while placing greater weight on local ownership, lower-cost systems and sovereign control. Indonesia's announced NVIDIA-linked sovereign AI factory is one example of this contest becoming physical infrastructure.^[^systems-indosat]

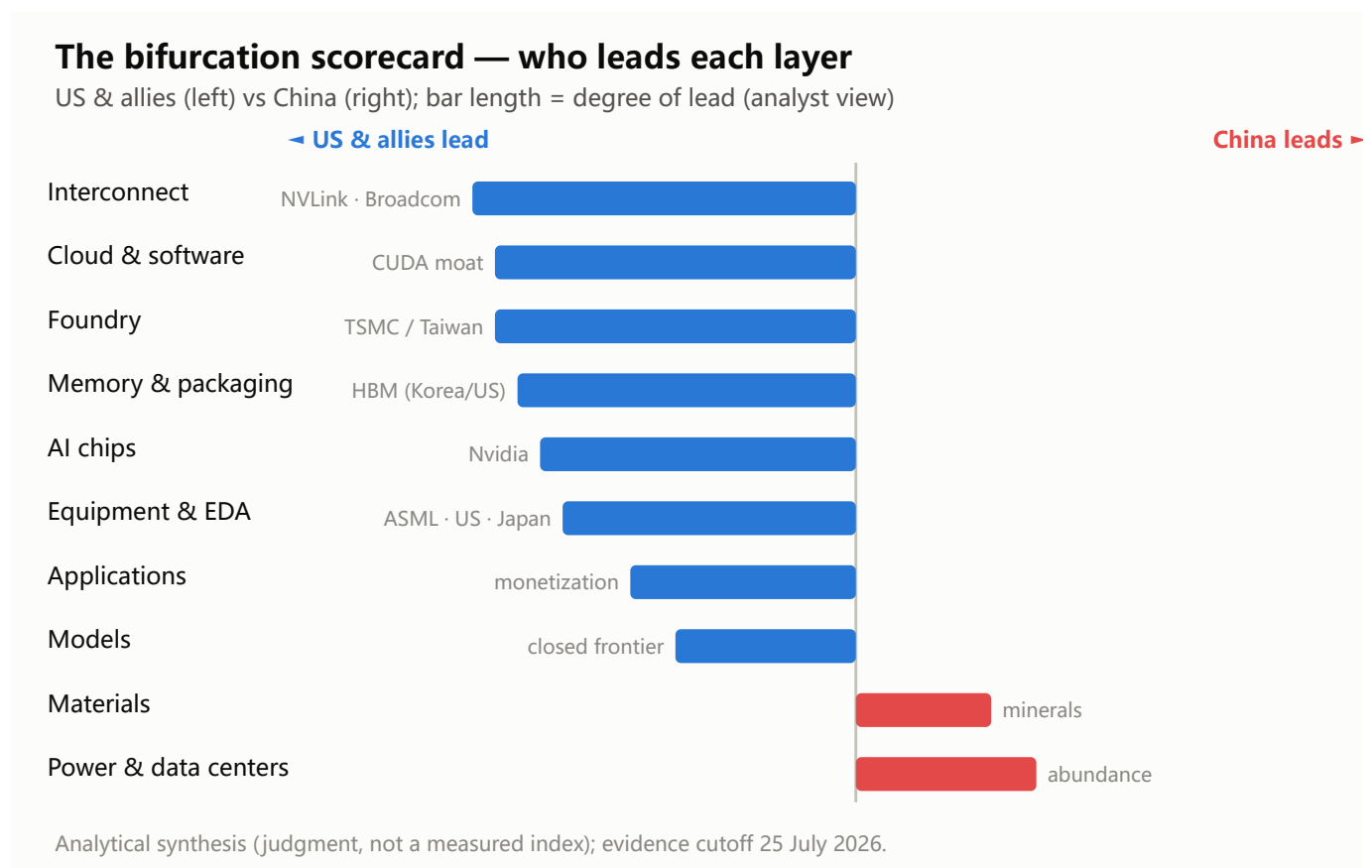
For investors, "sovereign AI" should be decomposed into funded phases, counterparties, energization dates and equipment orders. Announced national capacity is not

commissioned demand. The signals that matter are utility agreements, financing close, named system suppliers, construction progress and recurring cloud utilization.

3.5 Head to head

The scorecard

Put the two systems side by side and the picture is one of near-perfect complementarity. Each is strong exactly where the other is weak.



Who wins each layer

The complementarity is worth laying out explicitly, because it is the core of the whole atlas's geopolitics.

Layer	Advantage	Why
Models	US (closed) / China (open)	US frontier & monetization; China open-weight & usage
Applications	US	direct monetization vs China's free/ad model
Cloud & software	US	CUDA and hyperscaler moats
Power	China	abundance vs US grid constraint
Interconnect	US	NVLink/Broadcom/Arista
AI chips	US	Nvidia; China's Ascend is a mandated substitute
Memory & packaging	US + allies	Korea/US HBM, Taiwan/Japan packaging
Foundry	Neutral (Taiwan)	TSMC; China capped at 7nm
Equipment & EDA	US + allies	ASML/US/Japan; China ~35% self-sufficient
Materials	China (minerals) / allies (engineered)	the mutual chokehold

The allied lead concentrates in the compute-and-software layers and the neutral chokepoints; China leads at the raw base, in power and minerals; and the models and applications layers are the most contested.

The mutual hostage

The most important strategic fact is that this is one interdependent system splitting into two, and the split is incomplete in a way that matters. America cannot manufacture its designs without Taiwan; China cannot manufacture its designs at the leading edge without tools it is denied; both depend on the same neutral chokepoints. China can squeeze the West with minerals; the West can squeeze China with lithography, memory, and design software. That mutual dependence produces a fragile, deterrent stability — each side able to inflict real damage on the other — and a costly, duplicative scramble on both sides to build redundant supply chains. The 2025 minerals episode, in which

China's controls ran the antimony price up 2,600% before a truce, was a live demonstration of the leverage, and of why neither side has yet chosen to use it fully.

Cross-system implications

Sections 3.3 and 3.4 change the bilateral scorecard in two ways. The allied system holds much of the manufacturing leverage attributed to “the US and allies,” while the third-bloc system determines a growing share of marginal demand. The strategic contest is therefore not only about which superpower has the better domestic stack. It is about whether the allied chokepoints remain aligned, and which technical and financing system new buyers install.

Scenarios for the split

Three trajectories are worth holding. In a **managed-interdependence** base case, the two stacks deepen but the neutral chokepoints keep both sides mutually dependent, the truces hold, and the bifurcation is real but not catastrophic. In a **hard-decoupling** case, export controls escalate, the minerals truce lapses, and the two systems wall off further, which raises costs everywhere, strands specific assets, and rewards redundancy and domestic capacity on both sides. And in the **tail** case, a Taiwan Strait disruption collapses the shared foundation and overrides every other consideration at once. The base case is most likely; the tail case is the one that would invalidate the entire atlas.

The mirror-image thesis

For an investor the comparison yields three durable conclusions. The allied chokepoints, which both giants depend on, offer high strategic criticality but retain valuation, cyclicity, access and geopolitical risks. Bifurcation is deepening, so the contest for the uncommitted third bloc is where marginal infrastructure growth may be decided, with sovereign AI as the vehicle. Taiwan concentration sits behind both systems at once — the event that would break the whole structure and the reason to hold any AI exposure with humility about tail risk. The United States rented the base it leads on; China is building the base it was denied; and the two remain, for now, locked together by the chokepoints neither controls.

The physical bridge from this system-level comparison to accelerator packages, complete systems and normalized 100 MW facilities is \$2.13.

This part synthesizes Part II's stack chapters and §2.13 capstone. The scorecard is analytical synthesis, not a measured index. The evidence cutoff is 25 July 2026.

Part III system endnotes

[^systems-tsmc-hpc]: HPC Platform and 3DFabric

(https://www.tsmc.com/english/dedicatedFoundry/technology/platform_HPC_tech_WLSI). TSMC, accessed 2026-07-25.

[^systems-asml-products]: Lithography, metrology and computational products

(<https://www.asml.com/en/products>). ASML, accessed 2026-07-25.

[^systems-skhynix-hbm4]: SK hynix completes HBM4 development and readies mass production (<https://news.skhynix.com/sk-hynix-completes-worlds-first-hbm4-development-and-readies-mass-production/>). SK hynix, 2025; accessed 2026-07-25.

[^systems-humain]: HUMAIN portfolio company (<https://www.pif.gov.sa/en/our-investments/our-portfolio/humain/>). Public Investment Fund, accessed 2026-07-25.

[^systems-indosat]: Indosat and NVIDIA launch sovereign AI factory in Indonesia

(<https://ioh.co.id/portal/en/ioh-about-us-press-release?article=indosat-and-nvidia-launch-sovereign-ai-factory-in-indonesia>). Indosat Ooredoo Hutchison, accessed 2026-07-25.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [^atlas-chart-accelerator-share] [^atlas-chart-usgs-china-minerals-2025] [^atlas-mat-ajinomoto]

[^atlas-chart-accelerator-share]: Eye on the Market: Semiquincententacles

(<https://am.jpmorgan.com/content/dam/jpm-am-aem/global/en/insights/eye-on-the-market/semiquincententacles-amv.pdf>). J.P. Morgan Asset Management, 2026-06-26; accessed 2026-07-25.

[^atlas-chart-usgs-china-minerals-2025]: Mineral Commodity Summaries 2025

(<https://www.usgs.gov/publications/mineral-commodity-summaries-2025>). U.S. Geological Survey, 2025-01-31; accessed 2026-07-25.

[^atlas-mat-ajinomoto]: Ajinomoto Build-up Film

(<https://www.ajinomoto.com/innovation/action/buildupfilm>). Ajinomoto, undated; accessed 2026-07-25.

Part IV — Where It's All Going

This is the forward view, and it is an argument, not a list. From what the evidence shows — the roadmaps with dates, the capacity forecasts, the named projections of Goldman, Morgan Stanley, JPMorgan, the IEA, Epoch, and METR — it reasons toward what happens next and, crucially, where the opportunity sits. The through-line is a single moving target: the binding constraint keeps migrating, and the money follows it.

4.1 The master pattern: the constraint keeps migrating

The most useful thing to understand about the next two years is that the scarce input in AI is not fixed. In 2023 it was the GPU; by 2026 it is memory, packaging, and above all power; and each time the bottleneck moves, the pricing power — and the best trade — moves with it. The evidence for the migration is concrete. On the compute side, Nvidia has more than \$1 trillion of Blackwell-and-Rubin purchase orders visibility through 2027 (Jensen Huang, GTC 2026), with Blackwell sold out into late 2026, so *supply, not demand, is the binding constraint there*. One layer down, TrendForce projects high-bandwidth memory rising from 18% of DRAM wafer input at end-2025 to 30% by end-2027, with HBM contract prices possibly doubling in 2027 — a crowd-out that hands the memory makers pricing power. And one layer down again, transformer lead-times run three to four years, and roughly 7 of 12 gigawatts of US data centers planned for 2026 have already been canceled or delayed for lack of power and gear.

That last number is the tell. When announced compute gets canceled not for lack of chips but for lack of transformers, the constraint has demonstrably left the chip. The investable conclusion that runs through everything below: *own whatever the constraint currently gates*. Today that is power and the electrical supply chain; a year ago it was packaging; and the discipline is to keep asking where the scarcity is, not where it was.

4.2 Compute: the value migrates off the die

Nvidia's own roadmap tells the story of where compute value is going. Vera Rubin ships in 2026, Rubin Ultra in the second half of 2027 (15 exaFLOPS of FP4 inference, 384 GB of HBM per GPU, per TrendForce), and Feynman in 2028 — and Feynman's defining feature is *custom* HBM and 3D die-stacking, not a faster logic core. The architecture itself is telling you that the differentiation is migrating off the logic die into memory and packaging.

The competitive evidence points the same way. Counterpoint and JPMorgan project custom ASIC shipments surpass merchant GPUs in units in 2027 (roughly 12.5M ASIC vs 10.9M GPU of 23.3M total), with Broadcom holding about 60% of AI-server ASIC design share. The reasoning that matters for an investor is that *units are not revenue and neither is computing power*. Nvidia keeps the high-margin frontier-training and merchant markets while ASICs absorb captive, stable inference, so Nvidia's dollars can keep rising even as its unit share falls — the pie is growing fast enough (inference is already two-thirds of compute, and Gartner sees AI-laaS inference spend at 65%+ by 2029) that both can be true.

The opportunity, therefore, is not limited to picking the GPU-versus-ASIC winner; it also includes the layers into which system value is migrating. Three merit particular attention. **Memory is a direct but cyclical expression:** HBM growth can improve mix while crowding conventional DRAM wafer capacity, but high margins invite supply response. **Networking and optics** are the second: Broadcom spans custom ASICs, AI Ethernet, and co-packaged optics, alongside more concentrated component suppliers with greater customer risk. **Advanced packaging** is the third and a near-term deployment constraint. These are strategic observations; entry valuation and capacity normalization determine the eventual security return.

4.3 Models: capability compounds, ROI lags, and the variable that reconciles them

The models layer presents the sharpest tension in the whole atlas, and it is worth stating as a genuine contradiction. On one side, capability is compounding faster than almost anything in technology: Epoch AI measures frontier training compute growing about fivefold a year (and expects that to continue through 2030), with algorithmic efficiency improving threefold a year *on top*; METR finds the length of task an AI agent can

complete reliably is doubling roughly every three months and accelerating. On the other side, MIT found that about 95% of enterprise AI pilots delivered no measurable profit. Capability is racing; deployed value is lagging.

The reconciling variable — and the thing to watch above all in this layer — is *agent reliability at long horizons*. The METR curve says multi-hour, then multi-day, tasks become reliably automatable within a couple of years; the moment that reliability crosses the threshold enterprises need, the MIT failure rate inverts and the "GenAI divide" closes. Until then, the demand engine is not the pilots but the token explosion: Goldman Sachs projects token demand growing roughly 24-fold to 2030 as inference cost falls 60–70% a year and agentic workflows consume 5–30× the tokens of a chat query. That is a Jevons dynamic — cheaper inference expands, not shrinks, total consumption — and it is why aggregate AI spending (Gartner: \$2.5 trillion in 2026, rising to \$3.3 trillion in 2027) keeps climbing even as per-token economics collapse.[^atlas-chart-goldman-capex-2027]

The opportunities that follow. First, the efficiency curve favors the **inference-infrastructure and agent-orchestration layers** over the frontier-model layer, because the value accrues to whoever serves the exploding token volume, not to whoever trains the marginally-best model that open weights will match within a release cycle. Second, **physical AI is the next compute S-curve**: Nvidia's Cosmos world model, Google's Genie, and the humanoid wave (Goldman sizes the humanoid market at \$38 billion by 2035; Morgan Stanley, far more aggressively, at \$5 trillion by 2050) shift demand toward simulation and robotics compute and toward the humanoid supply chain of actuators, sensors, and batteries. Third, the open-weight convergence on coding compresses the margins of the closed coding models even as it expands the deployable base — a reason to favor the compute and the self-hosting infrastructure over the model vendors. And if Epoch's warned-of compute crunch materializes for long-context agentic workloads, pricing power swings *up* the stack to the compute owners and away from the application startups.

4.4 Power: the physical deployment constraint

Nowhere is the "own the constraint" logic more actionable than power, and the forecasts are unusually aligned on direction. The IEA sees data-center electricity roughly doubling from 415 TWh in 2024 to about 945 TWh by 2030; Goldman Sachs sees global data-center power up 165% by 2030 and US demand doubling from 31 GW in 2025 to 66 GW

by 2027; Morgan Stanley sees a roughly 49-gigawatt US shortfall and data centers reaching ~18% of US electricity by 2030. The disagreement is only about the severity of the shortfall, not its existence.[^atlas-global-iea-energy-ai-2025]

The reasoning that turns this into a trade is that the binding constraint is not generation in the abstract but the *equipment and the lead-times*: transformers at three-to-four years, gas turbines sold out through 2030 (GE Vernova's backlog reached 116 GW), and the 7 GW of canceled 2026 capacity as the first hard proof. The opportunities cascade from that. The **electrical-equipment and turbine makers** — Eaton, Siemens Energy, GE Vernova, Hitachi Energy, Quanta — have multi-year backlog visibility and pricing power, which is why they have already re-rated hard (a caution: much of the good news is priced, so the edge is now in the lead-time and book-to-bill data). The **behind-the-meter bridge** is the breakout second-order winner: Bloom Energy signed roughly \$7.65 billion of binding data-center fuel-cell contracts in about ninety days plus a \$25 billion Brookfield commitment, as solid-oxide fuel cells went from backup to primary generation. The **nuclear operators** — Constellation, Vistra, Talen — monetize existing fleets through twenty-year PPAs *now* (Microsoft–Constellation, AWS–Talen, Meta–Vistra), while the SMR developers are a 2030-plus optionality play the market keeps mispricing as near-term. And **liquid cooling** (Vertiv, and Eaton after its \$9.5 billion Boyd Thermal purchase) rides the rack-density curve. The structural second-order effect is that power has become strategic enough to pull hyperscalers into owning generation, which re-rates a swathe of "boring" electrical-industrial names into AI-adjacent growth.

4.5 Capital: the bubble question, quantified

The capital layer is where the whole thesis is stress-tested, and the useful move is to reduce the "is it a bubble" debate to the specific numbers that will resolve it. On the spending side, JPMorgan estimates \$5.5 trillion of AI capex through 2030, of which \$4.1 trillion is debt-financed; Morgan Stanley sees a \$1.5 trillion financing gap. On the revenue side, Sequoia's David Cahn calculates that 2026's ~\$1.5 trillion of infrastructure spend needs about \$3 trillion of revenue to justify, and Bain estimates the industry needs ~\$2 trillion of annual AI revenue by 2030 and will likely fall ~\$800 billion short. Those two sides are the bubble question, quantified.

The single number that decides it is *realized GPU economic life*. Michael Burry argues hyperscalers understate depreciation by roughly \$176 billion across 2026–28 by depreciating chips over five-to-six years when their economic life is two-to-three;

Goldman's own sensitivity shows that cutting useful life from five years to three swings cumulative 2026–31 depreciation from about \$3 trillion to about \$4 trillion — a \$1 trillion earnings hit. If inference revenue (Goldman's 24×) and enterprise ROI inflect up *before* that depreciation reality bites, the capex is validated; if not, the earnings reckoning arrives. The reasoning also identifies the likely *trigger*: with S&P having cut Oracle to one notch above junk and its credit-default swaps at multi-year highs, the repricing is more likely to start in the credit market than the equity market — so the actionable signal is single-name AI-debt CDS, not equity multiples.

Two opportunities follow. The first is that *debt is the new marginal buyer*, so the private-credit originators — Apollo, Blue Owl, Blackstone — capture the yield premium of financing the build (with the tail risk that opacity hides losses until they surface). The second is the rotation: with the Magnificent Seven at ~34% of the S&P and already lagging the other 493 in the first half of 2026 (the Russell 2000 Growth index up 17%), capital is visibly *de-concentrating* even as spending accelerates — the AI trade is broadening beyond the mega-caps, which is itself the actionable market signal.

4.6 Geopolitics: bifurcation, the sovereign demand pool, and the Taiwan clock

The policy trajectory has settled into an unstable pattern: the US oscillates between denial and monetization (the diffusion rule rescinded, a 15% China-sales tax imposed, H200 sales nominally allowed but with essentially zero China data-center revenue actually flowing), while China builds a parallel stack. The forward evidence suggests the export-control debate may already be losing relevance: domestic accelerators reached 41% of China's market in 2025, and Morgan Stanley projects 86% by 2030. The reasoning is that restriction accelerates the very self-sufficiency it aims to prevent — Nvidia's own warning that continued curbs "hand China's market to Huawei permanently" is being borne out.

Two opportunities emerge, plus one clock. The opportunity on the American side is **sovereign AI as a structural new demand pool** — Saudi Arabia's HUMAIN at over \$100 billion and 2,200 MW, the UAE's 5 GW Stargate, and \$66 billion of sovereign AI-infrastructure deployment in 2025 — which offsets Nvidia's lost China revenue and, because it runs on CUDA, entrenches the American stack (even \$100 billion Gulf budgets "can't buy their way out of Nvidia"). The opportunity on the Chinese side is the **closed parallel supply chain** — Huawei, SMIC, CXMT — a self-contained, capacity-constrained

universe insulated from US names, whose bottleneck is SMIC yields and domestic HBM. The clock is **Taiwan**: with ~70% of sub-2nm capacity staying in Taiwan through 2030 despite the Arizona build-out, the single-point-of-failure does not resolve this decade, and the more meaningful de-risking is CoWoS packaging localization, not wafer fabs. The second-order effect is that the whole regime — a 15% sales tax, equity stakes, case-by-case licensing — turns chips into an instrument of statecraft, which rewards diversified, less-US-exposed toolmakers and packagers.

4.7 Historical base rates: how infrastructure booms resolve

AI is technologically new; capital cycles are not. Four historical patterns provide useful—not mechanical—base rates.

Analogue	What persisted	What disappointed investors	Relevant test for AI
1990s telecom fiber	Data traffic and useful infrastructure	Utilization, pricing, leverage and duplicate networks	Can token growth fill commissioned capacity before debt and depreciation mature?
Cloud build-out, 2010s	Secular workload migration and hyperscaler scale	Falling unit prices and weak economics for undifferentiated hosts	Who retains margin as inference becomes cheaper?
Memory cycles	Rising long-run bit demand	Capacity additions, inventory correction and peak-earnings valuation	Does HBM qualification delay the normal supply response, and for how long?
China solar/EV localization	Rapid scale, falling global cost and strategic independence	Overcapacity, price wars and weak minority-shareholder returns	Does mandated AI capacity create durable economics or only strategic output?

The common lesson is that useful infrastructure can coexist with poor security returns. Demand growth does not protect a supplier whose capacity, leverage or valuation assumed even faster growth. The relevant discipline is to estimate each shortage's

scarcity half-life and compare it with announced capacity, qualification time and the multiple already paid.

4.8 The scenarios

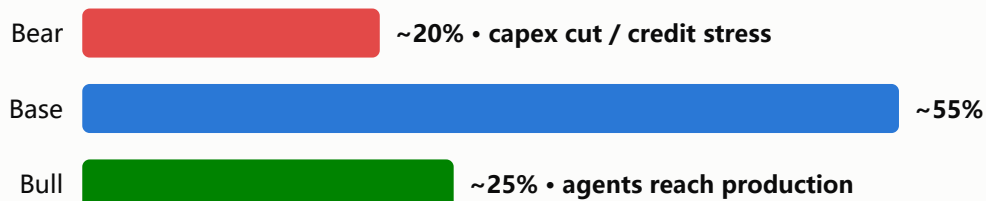
The reasoning above resolves into three scenarios and a tail. The horizon is the two years ending 25 July 2028. Probabilities are analytical weights, not measured frequencies.

Scenario	Weight	Numerical or observable confirmation	Probability moves higher when	Portfolio implication
Bear: financial air pocket	20%	At least one hyperscaler cuts planned capex; AI-linked credit spreads widen; utilization or backlog conversion falls materially	Two signals persist for a quarter	Reduce neoclouds and peak-multiple suppliers; favor balance-sheet strength and diversified cash generation
Base: slower monetization	55%	Capex grows but decelerates; production adoption improves unevenly; shortages normalize at different speeds	Revenue and utilization improve without a broad agent breakthrough	Own qualified bottlenecks selectively; demand valuation discipline and monitor capacity response
Bull: production-agent adoption	25%	Audited customer P&L benefits broaden; production conversion rises; disclosed AI revenue covers more depreciation	Two consecutive reporting periods confirm revenue and margin	Add workflow owners and higher-beta infrastructure while retaining physical chokepoints
Taiwan disruption tail	Unweighted tail	Material interruption to leading-edge production or logistics	Geopolitical indicators move beyond exercises and rhetoric	No listed basket fully hedges the physical interruption; reduce aggregate exposure and liquidity risk

Probability updates should be rule-based. One company announcement is insufficient. A five-percentage-point shift requires either two independent indicators or one audited system-level change, such as a sustained hyperscaler capex reduction or a broad production-ROI dataset.

Three scenarios for the next two years

Illustrative probability weighting (analyst judgment, ⚠️) — triggers and portfolios in the text



Analytical synthesis, not a forecast dataset; evidence cutoff 25 July 2026. Probabilities are scenario weights.

The base case (most likely) is that monetization broadens slowly. Enterprise ROI improves where workflows are deep and stays patchy elsewhere, capex stays high but decelerates, the truces hold, and no financial node breaks. Leadership rotates toward the bottleneck owners — power, memory, packaging, the neutral chokepoints, the credit financiers — while the crowded mega-caps and high-beta names trade volatile. The portfolio is the picks-and-shovels barbell, faded concentration, and quality-of-earnings screens.

The bull case is that agents cross into production. If reliability crosses the threshold and AI captures a share of labor budgets rather than software budgets, the addressable market expands an order of magnitude, MIT's 95% inverts, and demand pulls forward. Everything downstream re-rates, the application layer most of all. The tell is enterprise production deployments with audited P&L and frontier-lab gross margins turning up.

The bear case is financial, not technological. A hyperscaler cuts capex, or credit stress hits the GPU-backed debt, or an efficiency shock lands, or a lab funding round stumbles — and it transmits through the circular loop faster than a normal cycle. It does not require the technology to fail, only the money to get ahead of the earnings. The neoclouds, pre-revenue thematics, and peak-priced electrical orders fall first; the diversified hyperscalers and clean-earnings arms dealers hold up.[^atlas-chart-bis-ai-finance]

The tail is Taiwan — low-probability, catastrophic, un-hedgeable, and the reason to hold any AI exposure with humility.

4.9 What the consensus may be getting wrong

Four places where the crowd looks mispriced. It may be **under-pricing the power constraint's durability** — the physics of transformers and interconnection queues suggests it persists for years, which keeps the electrical complex a longer trade than consensus assumes. It may be **over-pricing the pre-revenue nuclear and humanoid stories**, which are 2030s payoffs trading as if they solve the 2027 crunch. It may be **under-appreciating how far the developer and open-weight layer has already gone Chinese**, which undercuts the "US leads AI" framing and matters for the third-bloc contest. And it is almost certainly **under-weighting the credit vector** — the shift from equity to debt is the least-discussed and most systemic feature of the build, and the regime change is likelier to show up in spreads than in stock prices.

4.10 The one number, and the calendar

If an investor tracks one thing, it should be whether disclosed hyperscaler AI revenue begins to *cover the depreciation* on what has been built — the metric that resolves the bubble question by settling the realized-GPU-life debate. Around it, the near-term calendar is unusually dense with resolving events.

Date / window	Catalyst	Why it matters
Late July 2026	Microsoft, Meta, Amazon earnings	the next capex re-rating event after Alphabet's beat-but-sell
H2 2026	Vera Rubin ships; AMD Helios; co-packaged optics; HBM4 volume	confirms or slips the hardware roadmap the whole build assumes
Late 2026	AI IPO window: rumored OpenAI S-1, SpaceX, Databricks	the most AI-concentrated IPO wave on record; aftermarket sets risk appetite
10 Nov 2026	US "50% affiliates" export rule scheduled to snap back	re-extends Entity-List controls to majority-owned subsidiaries
~27 Nov 2026	US–China minerals truce expires	watch for re-tightening of rare-earth and gallium controls
Through 2026–27	TSMC Arizona 2nm; SMIC yields; transformer/turbine lead-times; single-name AI-debt CDS	the real-time gauges of the manufacturing, power, and credit bottlenecks
2027	Rubin Ultra; ASIC unit-crossover; Intel 14A external test	whether the compute roadmap and the custom-silicon shift play out

The next several quarters will resolve most of the open questions of this atlas — the durability of capex, the trajectory of controls, the pace of the bottleneck easing. An investor who watches these signposts, and above all the depreciation-versus-revenue number and the credit spreads, will see the regime change in either direction before the market fully prices it.

4.11 Canonical monitoring dashboard

No single disclosed metric fully measures AI's economic return. The operating dashboard therefore uses a linked set of indicators.

Indicator	Current interpretation at cutoff	Bull confirmation	Bear warning	Review cadence
Hyperscaler capex and FCF after capex	Spending remains historically high	Capex converts to cloud/AI revenue without sustained FCF deterioration	Planned capex cut or debt-funded acceleration without revenue coverage	Quarterly
Disclosed AI revenue	Partial and inconsistently defined	Comparable disclosure broadens and covers more depreciation	Bundled metrics disappear or margins weaken	Quarterly
Accelerator useful life	Accounting lives generally exceed rapid product cadence	Older fleets retain utilization and pricing	Rental prices or utilization fall before debt maturity	Quarterly
Neocloud utilization and credit spreads	Contracted demand is high; financing remains central	Utilization, cash conversion and spreads improve together	Backlog delays, customer concentration or spread widening	Monthly/quarterly
HBM price, qualification and wafer allocation	Scarcity supports pricing	HBM4 yield and qualification remain tight while demand grows	Capacity and inventory rise faster than qualified demand	Monthly/quarterly
Advanced-packaging capacity	Still a critical delivery constraint	Capacity additions are absorbed without lead-time collapse	Lead times and pricing normalize faster than expected	Quarterly
Transformer and turbine lead times	Multi-year physical constraint	Book-to-bill and margin remain firm as capacity expands	Cancellations, queue withdrawals or rapid lead-time compression	Monthly/quarterly
Agent pilot-to-production conversion	Broad deployment remains limited	Audited production conversion and P&L impact rise	Cancellation rates remain high and contracts stay experimental	Semiannual

Indicator	Current interpretation at cutoff	Bull confirmation	Bear warning	Review cadence
Inference cost at fixed capability	Falling rapidly	Volume and useful-task demand grow faster	Price decline outruns demand and supplier margin	Quarterly
Ethernet versus proprietary fabric	Open networking is gaining relevance	Merchant Ethernet/optics take sockets without severe pricing pressure	Proprietary integration retains the frontier or merchant margins compress	Semiannual
SMIC yield and domestic HBM	China's principal physical ceiling	Verified yield, capacity and HBM qualification improve	Roadmap claims rise without shipment evidence	Quarterly
Export-control and ownership status	Transactional and entity-specific	Stable licensing and access rules	New entity, affiliate or ownership restrictions	Event-driven
Taiwan advanced-capacity distribution	Diversification is real but incomplete	Leading-edge and packaging capacity commissions outside Taiwan	Delays or concentration persists beyond disclosed plans	Semiannual

Every actionable security in Part V maps to at least two of these indicators and a stale-after date. This makes the atlas refreshable without rewriting its entire narrative whenever a quarterly number changes.

This part reasons from the forward-evidence research (research/4-trends-forward-evidence.md), with every projection attributed to its forecaster (IEA, Goldman Sachs, Morgan Stanley, JPMorgan, Sequoia, Bain, Gartner, Counterpoint, Epoch AI, METR, TrendForce, and company roadmaps). Scenario weightings are analyst judgment. Figures are mid-2026 and will move; contested items are flagged ⚠️ in the research.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [^{^atlas-chart-bis-ai-finance}] [^{^atlas-chart-goldman-capex-2027}] [^{^atlas-global-iea-energy-ai-2025}]

[^{^atlas-chart-bis-ai-finance}]: The AI investment race (<https://www.bis.org/publ/work1367.htm>). Bank for International Settlements, 2026-07-14; accessed 2026-07-25.

[^{^atlas-chart-goldman-capex-2027}]: What the capex boom means for stock investors (<https://www.goldmansachs.com/insights/articles/what-the-capex-boom-means-for-stock-investors>). Goldman Sachs, 2026-07-21; accessed 2026-07-25.

[^{^atlas-global-iea-energy-ai-2025}]: Energy and AI (<https://www.iea.org/reports/energy-and-ai>). International Energy Agency, 2025-04-10; accessed 2026-07-25.

Part V — Investment Decision Layer

This is the selective security-analysis companion to the 544-company handbook.

Coverage tier measures strategic importance; this section asks the separate questions of valuation, equity capture, expected return and portfolio overlap.

5A.1 Scope and valuation discipline

The shortlist contains **30 liquid US-listed securities or ADRs**. Prices are the latest market observations on **2026-07-24**, before the evidence cutoff. Observed P/E values are a screening field only: acquisition accounting, losses, stock splits, ADR ratios and different fiscal periods make them non-comparable. N/M means unavailable or not meaningful.

Bull/base/bear returns are **analytical two-year cumulative price-return ranges through 25 July 2028**, before dividends, taxes and currency effects. They are not price targets. The ranges impose valuation discipline on the operating thesis and should be refreshed whenever the price, estimate base or scenario probabilities change.

Strategic importance and security attractiveness are separate axes

Axis	Question	Primary evidence	Failure mode
Strategic criticality	Does the company gate the AI stack?	Substitutability, qualification, market position, policy role	Mistaking an important input for a material parent-company exposure
Equity capture	Does the listed security retain the economics?	AI revenue share, margins, pricing, customer power, capital intensity	Revenue growth without margin or cash conversion
Valuation	What does the price already assume?	Price, normalized earnings/cash flow, scenario multiple	Paying today for an implausibly long shortage
Expected return	Is upside adequate for the downside and common factors?	Scenario returns, catalysts, thesis-break conditions	Owning many securities that are one correlated capex trade

5A.2 Valuation and scenario screen

Ticker	Company	Layer	Price (USD)	Observed P/E	AI exposure	Valuation posture	Bear	Base	Bull
AMD	Advanced Micro Devices	AI chips	521.95	171.1x	High growth exposure; AI accelerators remain a smaller earnings base than CPUs	Very demanding	-60% to -40%	-5% to +15%	+45% to +80%
AVGO	Broadcom	AI chips	381.92	97.5x	High; custom silicon, switching and optics span multiple AI-system value pools	Demanding; GAAP P/E is affected by acquisition accounting	-45% to -25%	+5% to +25%	+40% to +65%
NVDA	NVIDIA	AI chips	206.84	31.5x	Very high; AI compute, networking and systems are the principal earnings engine	Premium but supported by current earnings	-45% to -25%	+5% to +25%	+40% to +70%
AMZN	Amazon	Cloud & platforms	232.11	27.8x	Material through AWS, Trainium and Anthropic; diversified by commerce and advertising	Reasonable with execution dependence	-35% to -20%	+10% to +30%	+40% to +65%
CRWV	CoreWeave	Cloud & platforms	71.88	N/M	Very high; direct GPU-cloud and AI-factory exposure	Speculative and leverage-sensitive	-80% to -55%	-20% to +10%	+45% to 0%

Ticker	Company	Layer	Price (USD)	Observed P/E	AI exposure	Valuation posture	Bear	Base	Bull
GOOGL	Alphabet	Cloud & platforms	319.74	16.1x	Material across advertising, Cloud, Gemini and TPU; diversified cash generation	Least demanding mega-cap valuation in the shortlist	-30% to -15%	+15% to +35%	+45% to +70%
META	Meta Platforms	Cloud & platforms	595.19	21.6x	Material through recommendation, advertising, models and custom infrastructure	Reasonable if advertising returns remain visible	-35% to -20%	+10% to +30%	+40% to +60%
MSFT	Microsoft	Cloud & platforms	381.70	22.7x	Material but diversified across cloud, software and model partnerships	Reasonable relative to cash generation	-30% to -15%	+10% to +30%	+40% to +60%
NBIS	Nebius	Cloud & platforms	187.77	N/M	Very high; AI-cloud capacity and related platform assets	Speculative; execution value dominates current earnings	-80% to -55%	-25% to +5%	+50% to +110%
ORCL	Oracle	Cloud & platforms	114.99	20.6x	High incremental exposure through OCI and large AI contracts; legacy software remains the cash base	Optically moderate, financially leveraged	-60% to -35%	-10% to +15%	+35% to +70%

Ticker	Company	Layer	Price (USD)	Observed P/E	AI exposure	Valuation posture	Bear	Base	Bull
AMAT	Applied Materials	Equipment & materials	536.25	50.5x	High indirect exposure across logic, memory and packaging equipment	Demanding for a cyclical supplier	-50% to -30%	+0% to +20%	+40% to +60%
ASML	ASML ADR	Equipment & materials	1757.09	N/M	High indirect exposure through leading-edge logic and memory capital intensity	Premium monopoly with cyclical and China sensitivity	-40% to -25%	+10% to +30%	+40% to +55%
ENTG	Entegris	Equipment & materials	129.15	74.2x	Medium-high through advanced-node materials and contamination control	Demanding with leverage and cycle sensitivity	-55% to -35%	+0% to +20%	+40% to +65%
KLAC	KLA	Equipment & materials	210.52	N/M	High indirect exposure through process-control intensity at advanced nodes	Premium; observed feed P/E excluded because of comparability inconsistency	-40% to -25%	+5% to +25%	+35% to +55%
LIN	Linde	Equipment & materials	512.28	34.0x	Low-medium direct exposure; semiconductor gases sit within a diversified industrial-gas franchise	Premium defensive compounder	-25% to -10%	+10% to +25%	+30% to +45%

Ticker	Company	Layer	Price (USD)	Observed P/E	AI exposure	Valuation posture	Bear	Base	Bull
LR CX	Lam Research	Equipment & materials	305.21	56.8x	High indirect exposure through memory and advanced-node process intensity	Demanding and memory-cycle sensitive	-55% to -35%	+0% to +20%	+40% to +65%
AL AB	Astera Labs	Interconnect	291.58	197.0x	Very high; connectivity products are concentrated in AI systems	Extremely demanding	-70% to -50%	-10% to +10%	+45% to +90%
AN ET	Arista Networks	Interconnect	173.99	58.8x	High growth exposure through scale-out Ethernet and cloud networking	Demanding	-45% to -25%	+5% to +25%	+40% to +60%
CO HR	Coherent	Interconnect	282.39	133.8x	High growth exposure through lasers and optical components; diversified industrial operations remain	Demanding; accounting earnings understate cash complexity	-55% to -35%	+0% to +20%	+45% to +75%
CR DO	Credo Technology	Interconnect	213.15	117.1x	Very high; active electrical cables and connectivity are AI-cluster driven	Extremely demanding	-70% to -50%	-15% to +5%	+50% to +95%

Ticker	Company	Layer	Price (USD)	Observed P/E	AI exposure	Valuation posture	Bear	Base	Bull
MRVL	Marvell Technology	Interconnect	194.23	66.8x	High incremental exposure through custom silicon, interconnect and optics	Demanding and program-dependent	-55% to -35%	+0% to +20%	+45% to +75%
AMKR	Amkor Technology	Memory & packaging	64.96	37.3x	Medium-high; advanced packaging is material but the portfolio is broader	Moderately demanding	-45% to -25%	+5% to +25%	+40% to +60%
MU	Microchip Technology	Memory & packaging	92.095	20.9x	High incremental exposure through HBM; conventional memory remains cyclical	Mid-cycle multiple on peak-like earnings risk	-60% to -40%	+0% to +20%	+40% to +70%
TSM	TSMC ADR	Memory & packaging	403.41	N/M	High; leading-edge logic and advanced packaging are central to AI systems	Premium quality with a structural location discount	-55% to -35%	+10% to +30%	+45% to +65%
CEG	Constellation Energy	Power & electrical	274.35	26.7x	Material incremental exposure through contracted nuclear power and large-load demand	Moderate premium for scarce existing generation	-40% to -25%	+10% to +30%	+40% to +65%
ETN	Eaton	Power & electrical	404.07	39.5x	Material incremental exposure through electrical distribution and	Demanding quality multiple	-40% to -25%	+5% to +20%	+30% to +50%

Ticker	Company	Layer	Price (USD)	Observed P/E	AI exposure	Valuation posture	Bear	Base	Bull
					data-center content				
GEV	GE Vernova	Power & electrical	1014.75	29.1x	Material incremental exposure through gas generation and grid equipment	Premium after substantial rerating	-50% to -30%	+0% to +20%	+40% to +65%
PWR	Quantum Services	Power & electrical	625.84	85.8x	Medium-high indirect exposure through transmission, substations and large-load construction	Very demanding for project execution risk	-50% to -30%	+0% to +15%	+35% to +55%
VRT	Vertiv	Power & electrical	290.36	73.0x	Very high incremental exposure through rack power and liquid cooling	Very demanding	-60% to -40%	-5% to +15%	+40% to +70%
VST	Vistra	Power & electrical	163.38	27.3x	Material incremental exposure through generation in constrained power markets	Moderate premium with commodity and market exposure	-40% to -25%	+5% to +25%	+35% to +60%

5A.3 Look-through factor concentration

The shortlist is diversified by product but not by ultimate demand. The factor counts below are intentionally overlapping.

Common factor	Securities exposed	Portfolio interpretation
Hyperscaler capex	26	Several layers ultimately depend on the same small group of cloud capital budgets.
Customer concentration	19	A small number of cloud or platform customers can dominate volume and bargaining power.
Capacity normalization	14	Current scarcity and margin can attract supply and shorten the expected shortage.
Taiwan and advanced-node concentration	11	Fabrication, packaging or upstream tool demand remains exposed to Taiwan and leading-edge capacity.
China controls	10	Licensing, tariffs, servicing limits or China revenue can change the earnings path.
Power and permitting	9	Orders require interconnection, permits, equipment, labor and commissioning before they become cash.
AI-linked credit	7	Debt cost, collateral values, refinancing and counterparty quality are central to the thesis.
Mega-cap concentration	5	The security is also a major index and passive-flow position.
Currency and ADR	4	The listed instrument adds currency, jurisdiction or ADR-ratio exposure to the operating thesis.
Memory cycle	2	Pricing and earnings can reverse when capacity and inventory overtake demand.

A portfolio containing semiconductors, memory, optics, ODM-linked suppliers and electrical equipment can still be one hyperscaler-capex position. Position sizing should be performed at the factor level, then at the security level.

5A.4 Actionable security files

AI chips

Advanced Micro Devices (AMD) — Very demanding.

- **What is priced in:** A large, profitable second-source accelerator franchise and strong Helios execution.
- **Variant view:** Open software and customer desire for a second source can create a durable niche without matching NVIDIA's full platform share.
- **Catalyst:** Helios shipments, large repeat deployments and improving accelerator gross margin.
- **Thesis break:** Software friction, roadmap delay or customers limiting AMD to low-margin negotiating leverage.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Accelerator useful life; Advanced-packaging capacity
- **Shared factors:** Hyperscaler capex, Taiwan and advanced-node concentration, China controls, Customer concentration

Broadcom (AVGO) — Demanding; GAAP P/E is affected by acquisition accounting.

- **What is priced in:** Continued custom-ASIC wins and merchant-networking leadership with successful VMware cash conversion.
- **Variant view:** Broadcom can earn content whether merchant GPUs or captive ASICs gain units, making its exposure broader than a single-chip bet.
- **Catalyst:** Additional hyperscaler design wins, AI-revenue conversion and sustained free cash flow.
- **Thesis break:** Customer insourcing, program concentration or acquisition leverage reducing financial flexibility.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Ethernet versus proprietary fabric; Advanced-packaging capacity
- **Shared factors:** Hyperscaler capex, Taiwan and advanced-node concentration, Customer concentration, AI-linked credit

NVIDIA (NVDA) — Premium but supported by current earnings.

- **What is priced in:** Sustained platform leadership, a successful Rubin transition and no severe hyperscaler capex interruption.

- **Variant view:** Networking, systems and software can offset accelerator-share dilution if NVIDIA retains the rack-level economics.
- **Catalyst:** Rubin availability, Spectrum-X growth and guidance strength excluding China.
- **Thesis break:** Material capex cuts, custom-silicon displacement at high-value workloads or a product-transition inventory correction.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Ethernet versus proprietary fabric; Accelerator useful life
- **Shared factors:** Hyperscaler capex, Taiwan and advanced-node concentration, China controls, Customer concentration, Mega-cap concentration

Cloud & platforms

Amazon (AMZN) — Reasonable with execution dependence.

- **What is priced in:** AWS acceleration and custom silicon improving AI economics without a major retail-margin reversal.
- **Variant view:** Trainium can reduce cost and preserve AWS margin even if merchant-GPU demand remains strong.
- **Catalyst:** AWS growth, Trainium utilization and durable operating cash flow after capex.
- **Thesis break:** Capex outruns AWS revenue or custom silicon fails to deliver meaningful cost advantage.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Disclosed AI revenue; Accelerator useful life
- **Shared factors:** Hyperscaler capex, Mega-cap concentration, Customer concentration

CoreWeave (CRWV) — Speculative and leverage-sensitive.

- **What is priced in:** High utilization, enforceable take-or-pay contracts and continued financing access.
- **Variant view:** Backlog is valuable only after energization and cash conversion; the equity is a utilization-and-credit option, not a software multiple.
- **Catalyst:** Commissioned capacity, customer diversification, free-cash-flow improvement and lower funding cost.
- **Thesis break:** Rental-price compression, contract delay, collateral markdown or refinancing stress.

- **Canonical monitors:** Neocloud utilization and credit spreads; Accelerator useful life; Power and permitting
- **Shared factors:** Hyperscaler capex, AI-linked credit, Customer concentration, Power and permitting

Alphabet (GOOGL) — Least demanding mega-cap valuation in the shortlist.

- **What is priced in:** Search disruption remains manageable and cloud growth converts despite heavy capital spending.
- **Variant view:** Proprietary distribution, data and TPU economics can make Alphabet a lower-priced AI platform rather than an incumbent victim.
- **Catalyst:** Cloud margin, AI-search monetization and evidence that capex produces incremental revenue.
- **Thesis break:** Search economics erode faster than cloud and AI products replace them.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Disclosed AI revenue; Inference cost at fixed capability
- **Shared factors:** Hyperscaler capex, Mega-cap concentration, China controls

Meta Platforms (META) — Reasonable if advertising returns remain visible.

- **What is priced in:** AI improves engagement and advertising enough to fund infrastructure and frontier-model spending.
- **Variant view:** The immediate return is recommendation and ad conversion, not a standalone assistant, giving Meta an internal monetization path.
- **Catalyst:** Advertising conversion, engagement and clearer return disclosure on infrastructure spending.
- **Thesis break:** Spending rises while ad economics, engagement or regulatory access deteriorate.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Disclosed AI revenue; Agent pilot-to-production conversion
- **Shared factors:** Hyperscaler capex, Mega-cap concentration, China controls

Microsoft (MSFT) — Reasonable relative to cash generation.

- **What is priced in:** Azure and Copilot growth offsetting depreciation while the broader software franchise remains durable.
- **Variant view:** Distribution and enterprise data may capture more value than frontier-model ownership, while the diversified cash engine limits financing risk.

- **Catalyst:** Comparable AI revenue disclosure, paid Copilot adoption and Azure margin resilience.
- **Thesis break:** AI capex depresses free cash flow without durable workload revenue or Copilot weakens seat economics.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Disclosed AI revenue; Agent pilot-to-production conversion
- **Shared factors:** Hyperscaler capex, Mega-cap concentration, AI-linked credit

Nebius (NBIS) — Speculative; execution value dominates current earnings.

- **What is priced in:** Rapid capacity deployment, strong utilization and financing without severe dilution.
- **Variant view:** The balance sheet and European footprint may differentiate Nebius, but announced capacity must become contracted, energized cash flow.
- **Catalyst:** Commissioning, contracted revenue, customer diversification and funding on improved terms.
- **Thesis break:** Build delays, weak utilization or equity issuance at depressed economics.
- **Canonical monitors:** Neocloud utilization and credit spreads; Power and permitting; Accelerator useful life
- **Shared factors:** Hyperscaler capex, AI-linked credit, Customer concentration, Power and permitting, Currency and ADR

Oracle (ORCL) — Optically moderate, financially leveraged.

- **What is priced in:** Large contracted demand converts on schedule without credit deterioration or customer renegotiation.
- **Variant view:** Backlog can create a scale cloud franchise, but equity value depends more on financing and conversion than headline contract size.
- **Catalyst:** Data-center delivery, OCI revenue conversion and stabilizing credit spreads.
- **Thesis break:** Project delays, concentrated counterparties or debt costs make contracted demand uneconomic.
- **Canonical monitors:** Neocloud utilization and credit spreads; Hyperscaler capex and FCF after capex; Accelerator useful life
- **Shared factors:** Hyperscaler capex, AI-linked credit, Customer concentration, Power and permitting

Equipment & materials

Applied Materials (AMAT) — Demanding for a cyclical supplier.

- **What is priced in:** AI-related WFE and packaging demand keep earnings above historical mid-cycle levels.
- **Variant view:** Broad process exposure captures rising complexity, but the multiple must account for eventual fab-spending normalization.
- **Catalyst:** Advanced-packaging orders, services and leading-edge customer spending.
- **Thesis break:** WFE digestion, China controls or capacity additions compressing utilization and margin.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Export-control and ownership status; HBM price, qualification and wafer allocation
- **Shared factors:** Hyperscaler capex, China controls, Taiwan and advanced-node concentration, Capacity normalization

ASML ADR (ASML) — Premium monopoly with cyclical and China sensitivity.

- **What is priced in:** EUV and High-NA demand support growth while China restrictions remain manageable.
- **Variant view:** Installed-base service and technical monopoly lengthen scarcity, but customer capex timing still controls annual earnings.
- **Catalyst:** High-NA acceptance, bookings and service growth.
- **Thesis break:** Customer delays, control-driven China decline or High-NA economics disappointing customers.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Export-control and ownership status; Taiwan advanced-capacity distribution
- **Shared factors:** Hyperscaler capex, China controls, Taiwan and advanced-node concentration, Currency and ADR, Customer concentration

Entegris (ENTG) — Demanding with leverage and cycle sensitivity.

- **What is priced in:** Advanced-node content, qualification and deleveraging offset semiconductor cyclicalities.
- **Variant view:** Long qualifications can protect share, but the parent must convert content growth into cash after expansion and acquisition.
- **Catalyst:** Advanced-node revenue, margin recovery and debt reduction.

- **Thesis break:** Customer inventory correction, qualification loss or leverage limiting investment.
- **Canonical monitors:** Taiwan advanced-capacity distribution; HBM price, qualification and wafer allocation
- **Shared factors:** Taiwan and advanced-node concentration, AI-linked credit, Capacity normalization

KLA (KLAC) — Premium; observed feed P/E excluded because of comparability inconsistency.

- **What is priced in:** Inspection intensity and service revenue remain resilient through WFE normalization.
- **Variant view:** Yield learning makes process control one of the more durable equipment value pools.
- **Catalyst:** Advanced-node mix, service growth and sustained margin.
- **Thesis break:** Leading-edge delays, control restrictions or customer concentration reducing orders.
- **Canonical monitors:** Taiwan advanced-capacity distribution; Export-control and ownership status
- **Shared factors:** Taiwan and advanced-node concentration, China controls, Customer concentration

Linde (LIN) — Premium defensive compounder.

- **What is priced in:** Stable pricing, project execution and diversified growth rather than a pure AI acceleration.
- **Variant view:** Long contracts and on-site infrastructure offer lower-beta exposure to fab and facility growth.
- **Catalyst:** Project backlog conversion and electronics-volume growth.
- **Thesis break:** Industrial slowdown, project returns or valuation compression overwhelming the modest AI contribution.
- **Canonical monitors:** Taiwan advanced-capacity distribution; Power and permitting
- **Shared factors:** Capacity normalization, Currency and ADR

Lam Research (LRCX) — Demanding and memory-cycle sensitive.

- **What is priced in:** HBM and leading-edge complexity offset a normal memory-equipment cycle.

- **Variant view:** More process steps per bit support structural content, but customer capacity still creates large cyclical swings.
- **Catalyst:** HBM-related orders, installed-base revenue and sustained customer utilization.
- **Thesis break:** Memory capacity overshoot, China restriction or customer digestion.
- **Canonical monitors:** HBM price, qualification and wafer allocation; Export-control and ownership status
- **Shared factors:** Memory cycle, China controls, Taiwan and advanced-node concentration, Capacity normalization

Interconnect

Astera Labs (ALAB) — Extremely demanding.

- **What is priced in:** Rapid content growth, durable platform wins and little pricing or customer pressure.
- **Variant view:** Connectivity complexity can increase content per rack, but the valuation leaves limited room for a normal qualification or platform pause.
- **Catalyst:** New platform sockets, broader customers and sustained gross margin.
- **Thesis break:** A platform transition, customer concentration or integration by larger silicon vendors.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Ethernet versus proprietary fabric
- **Shared factors:** Hyperscaler capex, Customer concentration, Capacity normalization

Arista Networks (ANET) — Demanding.

- **What is priced in:** Ethernet captures a growing share of AI scale-out while major cloud customers sustain spending.
- **Variant view:** Open Ethernet can gain sockets without displacing every proprietary frontier fabric.
- **Catalyst:** 1.6 Tb/s deployments, broader customer mix and continued AI-networking growth.
- **Thesis break:** Customer concentration, proprietary-fabric retention or merchant-switch pricing pressure.
- **Canonical monitors:** Ethernet versus proprietary fabric; Hyperscaler capex and FCF after capex

- **Shared factors:** Hyperscaler capex, Customer concentration, Capacity normalization

Coherent (COHR) — Demanding; accounting earnings understate cash complexity.

- **What is priced in:** AI optics growth and execution offset leverage, integration and non-AI cyclicalities.
- **Variant view:** Optical content can compound as bandwidth and power constraints intensify, but equity capture depends on product mix and debt reduction.
- **Catalyst:** Datacom growth, margin expansion and deleveraging.
- **Thesis break:** CPO timing slips, customer insourcing rises or leverage limits investment.
- **Canonical monitors:** Ethernet versus proprietary fabric; Neocloud utilization and credit spreads
- **Shared factors:** Hyperscaler capex, Customer concentration, AI-linked credit, Capacity normalization

Credo Technology (CRDO) — Extremely demanding.

- **What is priced in:** High-speed connectivity share gains persist with minimal customer or architecture disruption.
- **Variant view:** Copper can remain economically relevant longer than an all-optical narrative implies, but concentration dominates the risk.
- **Catalyst:** Additional customers and sustained content growth at higher speeds.
- **Thesis break:** Customer loss, faster optical substitution or price pressure.
- **Canonical monitors:** Ethernet versus proprietary fabric; Hyperscaler capex and FCF after capex
- **Shared factors:** Hyperscaler capex, Customer concentration, Capacity normalization

Marvell Technology (MRVL) — Demanding and program-dependent.

- **What is priced in:** Large custom programs ramp on time and diversify beyond a small number of customers.
- **Variant view:** Marvell offers high operating leverage to custom AI silicon but less program diversification than Broadcom.
- **Catalyst:** Tape-outs becoming volume revenue and stronger optical/interconnect mix.

- **Thesis break:** Program delay, customer insourcing or lower-than-expected margin on custom silicon.
- **Canonical monitors:** Hyperscaler capex and FCF after capex; Ethernet versus proprietary fabric; Advanced-packaging capacity
- **Shared factors:** Hyperscaler capex, Customer concentration, Taiwan and advanced-node concentration

Memory & packaging

Amkor Technology (AMKR) — Moderately demanding.

- **What is priced in:** Advanced-packaging demand fills new capacity at acceptable returns.
- **Variant view:** Geographic diversification and overflow packaging can gain value as customers seek alternatives to concentrated capacity.
- **Catalyst:** New advanced-package qualifications, capacity utilization and margin improvement.
- **Thesis break:** Customer delays, underutilized expansion or pricing power remaining with foundry customers.
- **Canonical monitors:** Advanced-packaging capacity; Hyperscaler capex and FCF after capex
- **Shared factors:** Hyperscaler capex, Capacity normalization, Customer concentration

Micron Technology (MU) — Mid-cycle multiple on peak-like earnings risk.

- **What is priced in:** HBM qualification and DRAM crowd-out keep pricing strong without a normal capacity-led correction.
- **Variant view:** HBM qualification can extend the upcycle, but conventional-memory supply remains the eventual earnings governor.
- **Catalyst:** HBM4 volume, allocation visibility and free cash flow through capacity expansion.
- **Thesis break:** Inventory and wafer capacity rise faster than demand or HBM yield/mix disappoints.
- **Canonical monitors:** HBM price, qualification and wafer allocation; Advanced-packaging capacity
- **Shared factors:** Hyperscaler capex, Memory cycle, Taiwan and advanced-node concentration, China controls

TSMC ADR (TSM) — Premium quality with a structural location discount.

- **What is priced in:** Leading-edge share, pricing and packaging growth continue while Taiwan risk remains latent.
- **Variant view:** Process leadership and customer breadth justify premium economics, but no valuation removes the common Taiwan tail.
- **Catalyst:** 2nm yield, CoWoS expansion and overseas capacity commissioning.
- **Thesis break:** Material Taiwan disruption, loss of process leadership or returns diluted by structurally higher overseas costs.
- **Canonical monitors:** Taiwan advanced-capacity distribution; Advanced-packaging capacity; Hyperscaler capex and FCF after capex
- **Shared factors:** Taiwan and advanced-node concentration, Hyperscaler capex, China controls, Currency and ADR, Customer concentration

Power & electrical

Constellation Energy (CEG) — Moderate premium for scarce existing generation.

- **What is priced in:** Long-term data-center contracting supports generation value without adverse regulatory intervention.
- **Variant view:** Existing nuclear plants monetize scarcity sooner than new-reactor narratives, but contract and political risk remain.
- **Catalyst:** Additional PPAs, regulatory approvals and fleet availability.
- **Thesis break:** Power-price reversal, regulatory cost allocation or operational underperformance.
- **Canonical monitors:** Power and permitting; Hyperscaler capex and FCF after capex
- **Shared factors:** Power and permitting, Hyperscaler capex, Customer concentration

Eaton (ETN) — Demanding quality multiple.

- **What is priced in:** Electrical scarcity, mix and margins remain elevated through capacity expansion.
- **Variant view:** Installed base and broad content support quality, but scarcity economics should not be capitalized indefinitely.
- **Catalyst:** Data-center orders, capacity commissioning and cash conversion.
- **Thesis break:** Lead times normalize rapidly, pricing fades or capacity is overbuilt.
- **Canonical monitors:** Transformer and turbine lead times; Power and permitting
- **Shared factors:** Power and permitting, Hyperscaler capex, Capacity normalization

GE Vernova (GEV) — Premium after substantial rerating.

- **What is priced in:** Turbine and grid scarcity remain durable and backlog converts at strong margins.
- **Variant view:** Service and installed-base economics can outlast the order spike, but current expectations require disciplined execution.
- **Catalyst:** Capacity expansion, margin conversion and named campus awards.
- **Thesis break:** Project cancellations, quality issues or turbine capacity normalizing faster than expected.
- **Canonical monitors:** Transformer and turbine lead times; Power and permitting; Hyperscaler capex and FCF after capex
- **Shared factors:** Power and permitting, Hyperscaler capex, Capacity normalization, Customer concentration

Quanta Services (PWR) — Very demanding for project execution risk.

- **What is priced in:** Grid investment and backlog convert with sustained labor availability and margin.
- **Variant view:** The grid build is broader than AI, but a record backlog does not guarantee cash or eliminate fixed-price risk.
- **Catalyst:** Large-load awards, electric-segment margin and cash conversion.
- **Thesis break:** Permitting, labor or project overruns weaken backlog economics.
- **Canonical monitors:** Transformer and turbine lead times; Power and permitting
- **Shared factors:** Power and permitting, Hyperscaler capex, Capacity normalization

Vertiv (VRT) — Very demanding.

- **What is priced in:** Rack-density growth, backlog and margins remain exceptional without project delays.
- **Variant view:** Content per megawatt rises structurally, but the multiple prices a long scarcity period.
- **Catalyst:** Liquid-cooling attach, backlog conversion and margin durability.
- **Thesis break:** Campus delays, competitive capacity or cooling architectures reducing expected content.
- **Canonical monitors:** Transformer and turbine lead times; Power and permitting; Hyperscaler capex and FCF after capex
- **Shared factors:** Power and permitting, Hyperscaler capex, Customer concentration, Capacity normalization

Vistra (VST) — Moderate premium with commodity and market exposure.

- **What is priced in:** Large-load demand supports power margins without a major commodity or regulatory reversal.
- **Variant view:** Existing dispatchable generation can capture near-term scarcity, but earnings retain power-price sensitivity.
- **Catalyst:** Data-center contracts, capacity-market strength and capital returns.
- **Thesis break:** Power-price normalization, regulatory intervention or fleet outages.
- **Canonical monitors:** Power and permitting; Hyperscaler capex and FCF after capex
- **Shared factors:** Power and permitting, Hyperscaler capex, Capacity normalization

5A.5 Portfolio construction by scenario

Portfolio sleeve	Base-case role	Bull behavior	Bear behavior	Risk control
Cash-generative platforms	Fund the build from diversified cash flow	Participate with lower operating leverage	Better balance-sheet resilience	Cap aggregate capex and antitrust exposure
Qualified semiconductor chokepoints	Capture content and scarcity	Benefit from faster system growth	Cyclical and Taiwan/control drawdown	Size the shared Taiwan and capex factors
Power and electrical	Monetize regional energization scarcity	Orders and margins persist	Backlog de-rates if projects slip	Track book-to-bill, lead times and cash conversion
High-beta infrastructure	Optionality on utilization and financing	Largest operating leverage	Largest credit and dilution loss	Trading sleeve; no assumption that backlog equals cash

The atlas does not prescribe portfolio weights. A disciplined implementation sets maximum exposure to hyperscaler capex, Taiwan, single-customer revenue, AI-linked credit and illiquidity before choosing individual names.

5A.6 Refresh and stale-after policy

- Prices and valuation posture: refresh monthly or after a $\pm 15\%$ move.
 - Earnings, capex, backlog and cash conversion: refresh quarterly.
 - Sanctions, ownership and exchange access: refresh on every policy event.
 - Scenario ranges: refresh after two confirming indicators or one audited system-level change.
 - Full shortlist review: stale after 25 October 2026.
-

Investment decision layer endnotes

[^decision-market-feed]: Market-price and observed-P/E snapshot supplied by the integrated market-data feed, latest trades on 24 July 2026. Corporate actions and accounting differences can impair comparability; the publication package preserves the raw values in `audit/investment-decision-data.json`.

[^decision-company-evidence]: Operating claims, access treatment and company-specific sources are preserved in the corresponding approved profiles and `universe/atlas.sqlite`. Return ranges, valuation posture and variant views are analytical synthesis.

Evidence links: [^decision-market-feed] [^decision-company-evidence]

Part VI — Reference & Apparatus

The lookup layer: the glossary, how the atlas was built, the tacit knowledge that the fundamentals don't spell out, the sources, and the disclaimers.

6.1 Glossary

Terms are rated for how much they matter: **[Critical]** you cannot follow the argument without it; **[Useful]** it sharpens understanding; **[Context]** helpful background. Chapters also define terms inline as they arise; this is the lookup.

Compute & chips

- **Accelerator / GPU** [Critical] — a chip built to do the massively parallel matrix math that neural networks require; the unit of AI compute.
- **Training vs inference** [Critical] — training is the one-time, expensive run that builds a model; inference is the recurring cost of using it. The shift toward inference-heavy reasoning is the master trend of the models layer.
- **Token** [Useful] — a word-piece; the unit models process and are priced by.
- **FLOPS / FP4 / FP8** [Context] — floating-point operations per second, increasingly at low precision to raise throughput.
- **ASIC** [Useful] — an application-specific chip (e.g., Google's TPU, AWS Trainium) designed for one workload, versus a general-purpose merchant GPU.
- **Rack-scale / "AI factory"** [Useful] — the shift from selling chips to selling integrated racks of dozens of GPUs with networking and cooling.

Memory, packaging, manufacturing

- **HBM (high-bandwidth memory)** [Critical] — stacked DRAM placed next to the logic die; now more than half a top GPU's cost and the binding memory constraint.

- **Advanced packaging / CoWoS** [Critical] — the step that bonds logic and memory into one module; TSMC's CoWoS is the industry's gating bottleneck.
- **Foundry** [Critical] — a contract chip factory (TSMC, Samsung, Intel Foundry, SMIC).
- **Process node / yield** [Useful] — the generation (2nm, 3nm) and the fraction of working chips per wafer.
- **EUV / High-NA lithography** [Critical] — extreme-ultraviolet light patterning; ASML's monopoly, without which no advanced chip exists.
- **GAA / backside power** [Context] — next-generation transistor and power-delivery techniques at 2nm.
- **WFE / EDA** [Useful] — wafer-fab equipment (the machines) and electronic design automation (the design software).

Models & software

- **Scaling laws / test-time compute** [Critical] — the old rule that bigger models improve predictably, and the new axis of making models reason at the moment of use.
- **Agent** [Critical] — a model that takes actions over a long horizon; the 2026 product story, gated by reliability.
- **MoE / distillation / quantization** [Context] — efficiency techniques that lower training and inference cost.
- **Open-weight vs closed** [Useful] — whether a model's parameters are published; China leads open-weight.
- **CUDA / CANN** [Critical] — Nvidia's software moat and Huawei's alternative; the two rival ecosystems.
- **NVLink / UALink / InfiniBand / Ultra Ethernet** [Useful] — the proprietary and open interconnect standards fighting over the networking layer.
- **Co-packaged optics** [Context] — moving optical links onto the switch/chip package to break the bandwidth-power wall.

Capital & markets

- **Capex supercycle** [Critical] — the ~\$725B (2026) of hyperscaler capital spending driving the build.
- **Circular / vendor financing** [Critical] — the industry funding its own demand (Nvidia→OpenAI→Oracle/CoreWeave→Nvidia).
- **RPO / backlog** [Useful] — remaining performance obligations; contracted future revenue, central to the neoclouds.

- **GPU-backed debt / ABS** [Useful] — debt collateralized by GPUs; the credit vector the BIS flagged.
- **Depreciation (useful life)** [Useful] — how fast chips are expensed; extending it flatters earnings (the Burry critique).
- **Jevons paradox** [Useful] — cheaper compute inducing more usage, so total spend rises.
- **Neocloud** [Useful] — a GPU-rental specialist (CoreWeave, Nebius).

Geopolitics

- **Neutral chokepoints** [Critical] — the hardest, least-substitutable points held by third parties both giants depend on (ASML/NL, TSMC/Taiwan, HBM/Korea, materials/Japan).
- **Bifurcation** [Critical] — the split of one global supply chain into a US-led and a China-led stack.
- **Entity List / NS-CMIC / 1260H** [Useful] — distinct US restrictions: trade (Entity List), ownership (NS-CMIC, the only one barring US persons from owning), and procurement (1260H). Never blur them.
- **自主可控 / 信创 (Xinchuang)** [Critical] — China's "autonomous & controllable" localization regime that mandates domestic chips for state buyers.
- **国产 hierarchy** [Useful] — the gradations of "domestic," from truly self-controlled to foreign-IP-dependent.
- **东数西算** [Context] — China's "east-data, west-compute" policy routing data centers to cheap-power regions.
- **Sovereign AI** [Useful] — government-funded national compute (Gulf, EU, India); the new demand pool.
- **Rare earths / gallium / antimony** [Useful] — the minerals China dominates and has used as leverage.
- **亿 (yi)** [Context] — Chinese unit = 100 million; a frequent source of 10× translation errors, corrected throughout.

6.2 Methodology & data hygiene

How it was built. The atlas uses a restart-safe research pipeline rather than a single all-or-nothing draft. The canonical SQLite database holds 544 included companies, normalized identities and access records, 3,184 evidence claims, 782 sources, profile state, chart state, and durable task manifests. Research batches write isolated JSON shards; the main process validates and imports them transactionally. Part V and the combined English manuscript are deterministic builds, so an interruption resumes from completed tasks rather than repeating the project.

How the AI-factory teardowns were built. Section 2.13 uses a separate, versioned evidence store at `supply-chain/atlas.sqlite`; it never mutates the approved company universe. Four dated platform states are modeled at accelerator-package, complete-system, and 100 MW commissioned IT-nameplate boundaries. Undisclosed prices, power, allocation, and facility assumptions use low/base/high ranges with calculation lineage. Source quality and claim certainty are recorded separately, supplier relationships are labeled confirmed, reported, inferred, roadmap, or speculative, and speculative lines are excluded from totals and investability comparisons. Installed capex excludes land, taxes, financing, and project revenue; recurring power, water, maintenance, and utilization assumptions are shown separately.

Three publication products. The deterministic source can be assembled as (1) a reader atlas containing the thesis, stack, system comparison, trends, 30-security decision layer and reference apparatus; (2) a company handbook containing all 544 approved companies; or (3) a full archival edition containing both. The products share the same canonical evidence and cutoff, but answer different questions. Coverage tier measures research depth and strategic importance; it is not an investment rating.

Universe scoring and calibration. Eligible companies pass an evidence gate before scoring. Strategic criticality contributes 30%, AI materiality 25%, market position 20%, investability 15%, and evidence quality 10%. Core coverage starts at 70 and is capped at 150, with documented chokepoint overrides; Extended covers scores of 45–69 or qualifying names displaced by the cap; Directory covers 30–44. The score determines publication depth. It does not incorporate security valuation, portfolio fit or expected return.

Security-analysis method. The Investment Decision Layer applies a second screen to 30 liquid US-listed securities and ADRs. It records the latest price observed before cutoff, an observed P/E when comparable, AI exposure, valuation posture, catalysts, thesis-break

conditions, common factors and analytical two-year bear/base/bull return ranges through 25 July 2028. These ranges are coherent scenario intervals, not price targets or probability-weighted forecasts. Prices, estimates, multiples and scenario weights require refresh before use.

Conflict resolution. When sources disagree, the atlas prefers the most recent primary filing or regulator record that measures the exact entity, period and unit in question. Official operating data and recognized industry bodies follow; reputable financial reporting is used when it adds unavailable context; specialist estimates are retained only when methodology is visible and uncertainty is disclosed. Conflicts are not averaged mechanically. The approved claim records the selected value, the disagreement and the reason for selection; unresolved material conflicts remain labeled contested.

The data-hygiene rules applied throughout:

- **One evidence cutoff (25 July 2026)**, with every figure's observation or forecast period labeled explicitly. The fastest-moving numbers (quarterly capex, private valuations, model leaderboards, export-control status) re-rate constantly and are flagged for re-verification.
- **Tier-1 anchoring.** Company filings, Reuters/CNBC/Bloomberg, and specialist primary sources (SEMI, USGS, IEA, Epoch AI, Menlo Ventures, METR, official leaderboards) are the anchors; aggregator and SEO figures are flagged with a warning mark and never treated as hard. **False precision on unverifiable data is treated as a worse error than honest generality.**
- **Number reconciliation.** Quantitative chart values, units, periods, uncertainty labels, and source IDs live in approved machine-readable sidecars; company and chapter claims live in SQLite. `_foundation/figures.md` is the reader-facing summary and must agree with those canonical records.
- **The 亿 conversion** (Chinese 100-million unit) is applied and known 10× errors corrected.
- **Sanction specificity.** Controls apply to the exact legal entity, and the three regimes (Entity List, NS-CMIC, 1260H) are kept distinct, because only NS-CMIC bars ownership.
- **Marks vs prices.** Private "valuations" are treated as illiquid marks with liquidation preferences, not market caps.
- **Version and staleness.** The edition date, evidence cutoff, price observation date and output hash travel with every build. Fast-moving market prices are stale immediately after cutoff; quarterly financials, capex plans and private marks should

be reviewed on the next reporting event; sanctions and access restrictions must be checked before any transaction. A new cutoff requires a whole-atlas evidence refresh rather than selective use of later facts.

6.3 Reading between the lines (shadow knowledge)

The tacit realities the fundamentals do not state, layered on top of the analysis.

On the US side. Circular financing is an accounting issue as much as a narrative: round-tripping, GPU-backed debt, depreciation-extension that flatters earnings, and a backlog concentrated in a few counterparties. Private "valuations" are marks, not prices. The AI trade is among the most crowded in history, so strong earnings can still sell off and passive flows amplify moves both ways. "AI revenue" is a soft, bundled, company-flattering metric. And the neutral chokepoints are not risk-free: ASML still earns a large share of sales from China, and TSMC's "own it safely" case rests on the unstated assumption that Taiwan stays undisrupted.

On the China side. Demand is two-track: commercial buyers still want Nvidia, but it is government, SOE, and military procurement that must be 自主可控, and that mandated slice, not open-market competitiveness, is what protects the champions. "国产" is a hierarchy, so ask how domestic a name really is. A-shares are retail-driven, and many thematic names are 概念股 riding a story at extreme multiples with little real exposure. Watch the China-specific tells: pledged shares (股权质押), subsidy-inflated profit (check non-GAAP), state-fund selling overhangs (大基金减持), and fabricated order rumors (小作文). On access, many STAR names are Connect-restricted or institutional-only, the crown-jewel private companies are inaccessible, and headline performance claims (an Ascend cluster "beating" a rival) are system-level, using far more chips and power.

6.4 Source index

The evidence base is preserved in the repository for traceability:

- **Canonical registry:** `universe/atlas.sqlite`, with deterministic JSON/CSV exports in `universe/review/`.
- **GPU-to-AI-factory evidence:** `supply-chain/atlas.sqlite`, with immutable snapshot exports, source-access results, validation, and checksums in `supply-chain/published/`.
- **Company evidence:** isolated profile shards in `runs/2026-07-25-profiles/results/` and validated profiles in `en/part-5-universe/profiles/`.
- **Chart provenance:** 36 records in `charts/sidecars/`; the six \$2.13 records include row-level source IDs, dark-SVG checksums and semantic CSV fallbacks, while the legacy chart sources remain normalized in `audit/chart-sources.json`.
- **Chapter claims:** machine-readable packages in `audit/master-claims.json` and the SQLite claim-to-source joins.
- **Research appendices:** focused files in `research/` for models, applications, cloud, equipment/EDA, materials, trends, and the original profile work.
- **Reader summary:** `_foundation/figures.md`, which records approved values, periods, sources, and uncertainty labels.

Principal tier-1 sources cited across the atlas include: company filings and earnings; Reuters, CNBC, Bloomberg; SEMI and USGS; the IEA and DOE/Berkeley Lab; Epoch AI, METR, and Stanford's AI Index; Menlo Ventures; TrendForce, Counterpoint, and SemiAnalysis; the BIS and IMF; and the official Arena and Artificial Analysis model leaderboards.

6.5 Disclaimers

This document is research and structural analysis prepared for a business and investment audience. **It is not investment advice, an offer, or a solicitation**, and it does not account for any individual's circumstances. The evidence cutoff is 25 July 2026. Figures are drawn from public sources believed reliable but not independently audited, and several—especially private valuations, forward capital-spending estimates, market-share splits, and any figure flagged with a warning mark—are contested or fast-moving. Verify

every ticker, financial, sanction status, and valuation against a current primary source before making any decision. Forward-looking statements and scenarios are judgments, not forecasts. The author holds no brief for any company named. Past performance and current positioning are not indicative of future results.

End of Part VI. This release covers the English Markdown reader atlas, company handbook, full archival edition, searchable HTML and paged PDF, plus a native Chinese §2.13 capstone built from the same snapshot. A complete Chinese edition remains a later language wave.

Chapter evidence

Linked evidence for this chapter's figures and load-bearing claims: [[^]atlas-chart-bis-ai-finance] [[^]atlas-chart-usgs-china-minerals-2025]

[[^]atlas-chart-bis-ai-finance]: The AI investment race (<https://www.bis.org/publ/work1367.htm>). Bank for International Settlements, 2026-07-14; accessed 2026-07-25.

[[^]atlas-chart-usgs-china-minerals-2025]: Mineral Commodity Summaries 2025 (<https://www.usgs.gov/publications/mineral-commodity-summaries-2025>). U.S. Geological Survey, 2025-01-31; accessed 2026-07-25.

GPU-to-AI-factory evidence explorer

Filter the canonical package, system, facility and supplier records. Low/base/high values stay together so ranges cannot be mistaken for point forecasts.

66 canonical records shown.

Eco system	Generation	Scale	Geography	Confidence / evidence	Record	Low / base / high
china	current	package	China	high / verified_current	ascend-910c-npu-package	1
china	current	package	China	high / verified_current	on-package-high-bandwidth-memory,-generation-and-supplier-undisclosed	128
china	next	system	China	high / company_announced	atlas-950-communications-cabinet	32
china	next	system	China	high / company_announced	atlas-950-compute-cabinet	128
china	next	system	China	high / company_announced	ascend-950dt-accelerator-card	8,192
china	next	system	China	high / verified_current	ascend-950-family-accelerator-card,-exact-variant-undisclosed	1,024
china	next	package	China	high / company_announced	huawei-hizq-2.0-hbm	144
china	next	package	China	high / company_announced	ascend-950dt-npu-package	1

Eco syst em	Gen erat ion	Sca le	Geography	Confidenc e / evidence	Record	Low / base / high
china	current	system	China	high / verified_current	47u-air-cooled-bus-equipment-cabinet	4
china	current	system	China	high / verified_current	47u-liquid-cooled-compute-cabinet	12
china	current	system	China	high / verified_current	kunpeng-920	192
china	current	system	China	high / verified_current	ddr5-dimm,-64-gb-or-96-gb	1,536
china	current	system	China	high / verified_current	ascend-910c	384
china	current	system	China	medium / verified_current	dual-three-phase-380-v-ac-input	2
china	current	system	China	high / verified_current	2.5-inch-drive	480
us	current	package	United States-led	high / verified_current	blackwell-ultra-gpu-package	1
us	current	package	United States-led	high / verified_current	12-high-hbm3e-stack	8
us	current	package	United States-led	high / verified_current	blackwell-ultra-reticle-limited-compute-die	2

Eco syst em	Gen erat ion	Sca le	Geography	Confidenc e / evidence	Record	Low / base / high
us	curr ent	sys te m	United States-led	high / verified_cu rrent	nvidia-bluefield-3-b3240- dpu	18
us	curr ent	sys te m	United States-led	high / verified_cu rrent	gb300-nvl-compute-tray	18
us	curr ent	sys te m	United States-led	high / verified_cu rrent	nvidia-connectx-8- network-adapter	72
us	curr ent	sys te m	United States-led	high / verified_cu rrent	nvidia-grace-cpu	36
us	curr ent	sys te m	United States-led	high / verified_cu rrent	nvidia-b300-blackwell- ultra-gpu-package	72
us	curr ent	sys te m	United States-led	high / verified_cu rrent	fifth-generation-nvswitch- asic	18
us	curr ent	sys te m	United States-led	high / verified_cu rrent	33-kw-power-shelf	8
us	next	pac kag e	United States-led	high / verified_cu rrent	rubin-gpu-package	1
us	next	pac kag e	United States-led	medium / verified_cu rrent	rubin-hbm4-capacity-gb- per-package	288
us	next	pac kag e	United States-led	high / verified_cu rrent	rubin-reticle-limited- compute-die	2

Eco system	Gen eration	Sca le	Geography	Confidenc e / evidence	Record	Low / base / high
us	next	sys tem	United States-led	high / modeled	vera-rubin-nvl-compute-tray	18
us	next	sys tem	United States-led	high / verified_current	nvidia-vera-cpu	36
us	next	sys tem	United States-led	high / verified_current	nvidia-rubin-gpu-package	72
us	next	sys tem	United States-led	medium / verified_current	rubin-hbm4-capacity-tb-per-system	20.74
us	next	sys tem	United States-led	medium / verified_current	vera-cpu-lpddr5x-socamm2-capacity	54
us	next	sys tem	United States-led	high / verified_current	nvlink-6-switch-asic	36
china	all	facility	China	high / verified_current	Compute, networking and storage active IT reference	0 USD million per 100 MW IT
china	all	facility	China	medium / modeled	General contractor preliminaries, fees, margin and contingency	60.06 / 77.15 / 94.68 USD million per 100 MW IT
china	all	facility	China	medium / modeled	Electrical distribution and resilience from project transformers through rack busway	288.31 / 370.31 / 454.47 USD million per 100 MW IT

Eco syst em	Gen erat ion	Sca le	Geography	Confidenc e / evidence	Record	Low / base / high
china	all	facility	China	medium / modeled	Facility-side liquid cooling, residual air cooling and heat rejection	198.21 / 254.59 / 312.45 USD million per 100 MW IT
china	all	facility	China	low / modeled	Professional services, owner controls, commissioning and initial facility spares	30.03 / 61.72 / 113.62 USD million per 100 MW IT
china	all	facility	China	medium / modeled	Core, shell and architectural fit-out	54.06 / 69.43 / 85.21 USD million per 100 MW IT
china	all	facility	China	low / modeled	Site utility interconnect, substation extension and metering allowance	30.03 / 61.72 / 113.62 USD million per 100 MW IT
us	all	facility	United States	high / verified_ current	Compute, networking and storage active IT reference	0 USD million per 100 MW IT
us	all	facility	United States	medium / modeled	General contractor preliminaries, fees, margin and contingency	101.65 / 118.27 / 146.3 USD million per 100 MW IT
us	all	facility	United States	medium / modeled	Electrical distribution and resilience from project transformers through rack busway	487.92 / 567.67 / 702.24 USD million per 100 MW IT
us	all	facility	United States	medium / modeled	Facility-side liquid cooling, residual air cooling and heat rejection	335.45 / 390.27 / 482.79 USD million per 100 MW IT

Eco system	Generation	Scale	Geography	Confidence / evidence	Record	Low / base / high
us	all	facility	United States	low / modeled	Professional services, owner controls, commissioning and initial facility spares	50.83 / 94.61 / 175.56 USD million per 100 MW IT
us	all	facility	United States	medium / modeled	Core, shell and architectural fit-out	91.49 / 106.44 / 131.67 USD million per 100 MW IT
us	all	facility	United States	low / modeled	Site utility interconnect, substation extension and metering allowance	50.83 / 94.61 / 175.56 USD million per 100 MW IT
china	current	supplier	China	high / verified_current	Huawei Technologies / HiSilicon: system vendor and semiconductor designer	1 share
china	current	supplier	China	medium / independently_reported	Semiconductor Manufacturing International Corporation: reported foundry	Allocation undisclosed
china	next	supplier	China	high / verified_current	Huawei: UnifiedBus system architect	Allocation undisclosed
china	next	supplier	China	high / company_announced	Huawei Technologies / HiSilicon: system vendor and semiconductor designer	1 share
china	next	supplier	China	high / company_announced	Huawei HiZQ: proprietary memory architecture brand	Allocation undisclosed

Eco system	Generation	Scale	Geography	Confidence / evidence	Record	Low / base / high
china	current	supplier	China	high / verified_current	Huawei Kunpeng: cpu product vendor	1 share
us	current	supplier	Taiwan; global network	high / verified_current	Hon Hai Precision Industry (Foxconn): Pilot build, co-design, rack manufacturing and liquid-cooling integration	Allocation undisclosed
us	current	supplier	Undisclosed multi-country network	high / verified_current	Micron Technology: HBM3E and LPDDR5X SOCAMM memory	Allocation undisclosed
us	current	supplier	United States	high / verified_current	NVIDIA: GPU, CPU, NVLink, networking silicon and platform design	Allocation undisclosed
us	current	supplier	South Korea; origin allocation undisclosed	high / verified_current	SK hynix: HBM3E memory	Allocation undisclosed
us	current	supplier	Taiwan; United States	high / verified_current	Taiwan Semiconductor Manufacturing Company: Blackwell Ultra GPU wafer fabrication	Allocation undisclosed
us	current	supplier	United States; Taiwan; global network	high / verified_current	Wistron: Superchip board assembly and test	Allocation undisclosed
us	next	supplier	Undisclosed multi-country network	high / verified_current	Micron Technology: HBM4, LPDDR5X SOCAMM2 and qualified storage products	Allocation undisclosed

Eco syst em	Gen erat ion	Sca le	Geography	Confidenc e / evidence	Record	Low / base / high
us	next	sup plie r	United States	high / verified_cu rrent	NVIDIA: GPU, CPU, NVLink, networking silicon and platform design	Allocation undisclosed
us	next	sup plie r	South Korea; origin allocation undisclosed	high / verified_cu rrent	Samsung Electronics: HBM4 and SOCAMM2 memory	Allocation undisclosed
us	next	sup plie r	South Korea; origin allocation undisclosed	medium / independe ntly_report ed	SK hynix: Co-development and supply of advanced memory for Vera Rubin	Allocation undisclosed
us	next	sup plie r	Global	high / verified_cu rrent	Dell Technologies; Hewlett Packard Enterprise; Lenovo; Supermicro; Foxconn; Quanta Cloud Technology; Wistron; Wiwynn: Named system builders in full-scale production and DSX adopters	Allocation undisclosed
us	next	sup plie r	United States; Taiwan; global network	high / company_ announce d	Wistron: Superchip board assembly and test	Allocation undisclosed