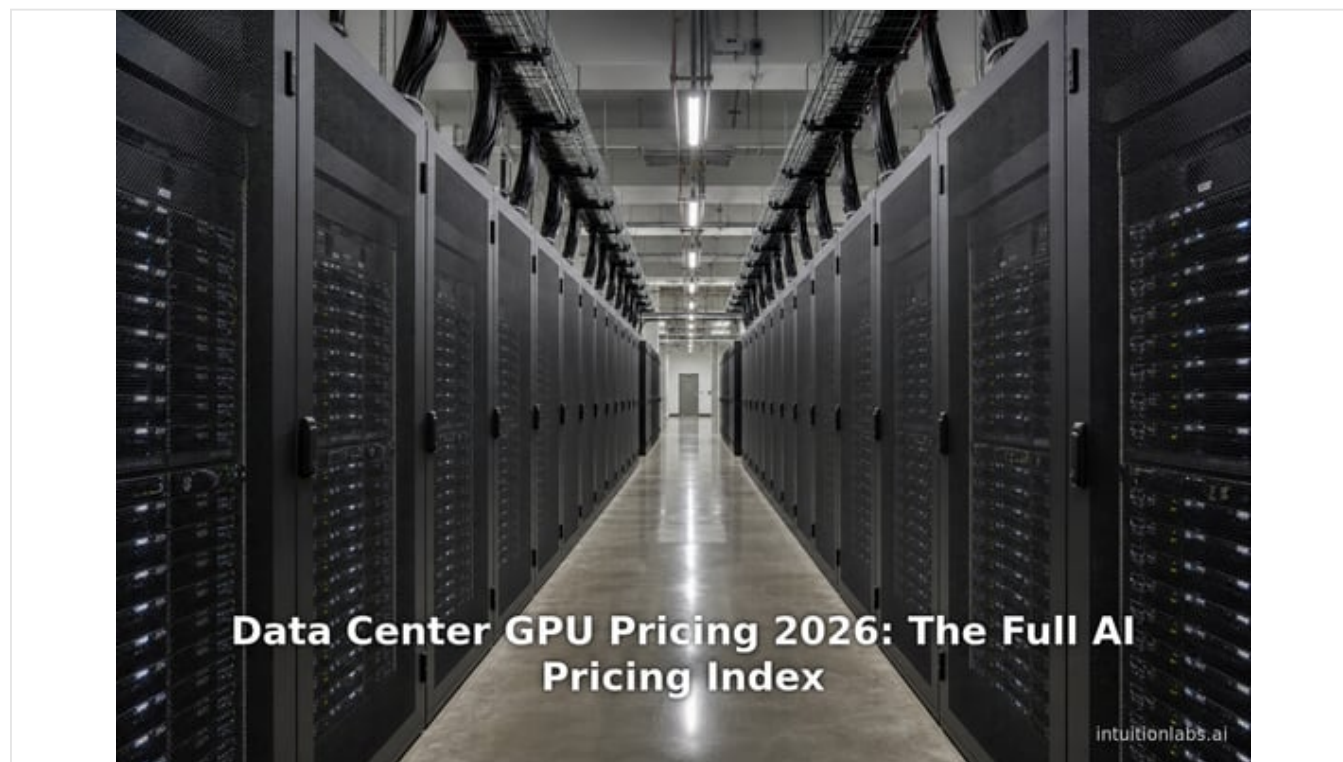


Data Center GPU Pricing 2026: The Full AI Pricing Index

Published July 20, 2026 33 min read



Executive Summary

Data center GPU pricing in 2026 is defined by a wide spread rather than a single number. An NVIDIA **H100** 80GB GPU rents for as little as \$1.38 to \$1.49 per GPU-hour on marketplace and boutique clouds and as much as \$11.68 to \$12.29 per GPU-hour on hyperscaler on-demand instances, an 8x to 9x range for functionally identical silicon (Source: [deploymentbase.ai](#)) (Source: [gpucloudcost.com](#)). Google Cloud's own pricing page lists the 8-GPU **a3-highgpu-8g** H100 instance at \$88.49 per hour on-demand, or roughly \$11.06 per GPU, while Amazon's EC2 P5 page confirms up to 8 H100 or H200 GPUs per instance in clusters scaling to 20,000 GPUs (Source: [cloud.google.com](#)) (Source: [aws.amazon.com](#)). The newer **H200**, with 141GB of HBM3e memory versus the H100's 80GB, carries a 25 to 30 percent premium and starts near \$3.72 per hour on specialist clouds, while NVIDIA's Blackwell-generation **B200** runs \$2.12 to \$6.04 per hour depending on provider tier and contract length, and a full **GB200 NVL72** rack (72 Blackwell GPUs) prices out at \$10.50 to \$27 per GPU-hour, or \$756 to \$1,944 per hour for the whole rack (Source: [gpuaas.com](#)) (Source: [spheron.network](#)).

Purchasing hardware outright follows the same wide bands. A new H100 costs \$25,000 to \$40,000 per card depending on PCIe versus SXM5 form factor, an 8-GPU DGX H100 system runs \$300,000 to \$460,000, and an 8-GPU H200 DGX system is quoted at roughly \$400,000 to \$500,000 (Source: [cloudzero.com](#)) (Source: [intuitionlabs.ai](#)). Used H100 cards that sold for \$40,000 in late 2023 now trade for \$6,000 to \$22,000 on secondary markets, an 85 percent decline that reflects both Blackwell's arrival and a supply base that has broadened from a handful of hyperscalers to more than two dozen specialist clouds (Source: [cloudzero.com](#)). Break-even between renting and owning a single H100 lands around 18 to 24 months of continuous, high-utilization operation once power, cooling, and networking are counted, which is why most [organizations below hyperscaler scale](#) still rent (Source: [cloudzero.com](#)).

The pricing spread is widening, not narrowing, because of a genuine supply squeeze. Contract pricing for H100 and H200 GPUs climbed roughly 40 percent between October 2025 and March 2026 as HBM3e memory costs passed through from Samsung and SK Hynix collided with surging demand, and industry trackers describe finding eight-GPU nodes of Hopper-class capacity on short notice as "no longer routine" (Source: [gpuaas.com](#)). Lead times for Blackwell PRO GPUs stretched to three to seven months by early 2026, and the Financial Times, cited by Reuters, reported that NVIDIA's export-restricted chips more than doubled in price on China's black market in the first half of 2026 (Source: [electropages.com](#)) (Source: [reuters.com](#)). Demand-side pressure is equally stark: NVIDIA reported record Data Center revenue of \$75.2 billion for its fiscal first quarter of 2027 (the three

months ended April 26, 2026), up 92 percent year over year, while the four largest hyperscalers, Amazon, Google, Meta, and Microsoft, plan a combined [\\$630 billion in 2026 capital expenditures](#), a 62 percent jump from 2025's \$388 billion (Source: [nvidianews.nvidia.com](#)) (Source: [datacenter richness.substack.com](#)). Building the data centers that house this compute now costs \$20 million or more per megawatt for AI-optimized facilities, versus \$7 to \$12 million per megawatt for conventional builds, with fully built-out hyperscale campuses modeled at \$45 to \$55 billion per gigawatt (Source: [gigaenergy.com](#)). For buyers, the practical takeaway is that provider selection and contract structure, not the GPU model alone, now determine the bill: reserved multi-year commitments cut 25 to 50 percent off on-demand rates, and specialist "neoclouds" such as CoreWeave, Lambda, and Nebius consistently undercut AWS, Azure, and Google Cloud on-demand pricing by roughly 40 to 70 percent for identical hardware.

Introduction and Background

Twenty-four months ago, a single NVIDIA H100 was a scarce commodity that commanded whatever a hyperscaler chose to charge. By mid-2026, the same chip is sold through more than two dozen distinct channels, from spot marketplaces to reserved hyperscaler capacity blocks to secondary-market resale, at prices that can differ by a factor of eight or more for the identical part (Source: [cloudzero.com](#)). This report is a data-driven index of what data center **graphics processing units (GPUs)**, the specialized processors that perform the parallel matrix math underlying artificial intelligence (AI) model training and inference, actually cost across the market as of July 2026. It draws on official cloud provider pricing pages, NVIDIA's own financial filings, independent pricing trackers that publish verified vendor-by-vendor rate tables, and reporting on the largest AI infrastructure contracts signed to date.

The market this report covers spans three NVIDIA GPU generations still in active commercial use: the **Hopper**-architecture **H100** (80GB HBM3 memory, released 2022) and its memory-upgraded sibling the **H200** (141GB HBM3e), and the **Blackwell**-architecture **B200** (192GB HBM3e) alongside its rack-scale sibling the **GB200 NVL72**, which packages 72 B200 GPUs and 36 Grace CPUs into a single liquid-cooled cabinet (Source: [gpuaas.com](#)) (Source: [spheron.network](#)). Buyers access this hardware through four distinct channels that this report treats separately throughout: **hyperscaler clouds** (AWS, Microsoft Azure, Google Cloud, Oracle Cloud Infrastructure), **specialist or "neocloud" providers** (CoreWeave, Lambda, Nebius, Crusoe), **long-tail and marketplace providers** (Vast.ai, RunPod, ThunderCompute, Hyperstack), and **direct hardware purchase for organizations building owned infrastructure**. Each channel prices the same underlying silicon differently because each bundles a different mix of networking, service-level agreements, compliance certifications, and sales overhead into the hourly rate (Source: [mercatus-ai.com](#)).

NVIDIA does not publish an official price list for its data center chips, which is precisely why third-party price discovery matters and why this report leans on verified, dated vendor quotes rather than manufacturer suggested pricing (Source: [intuitionlabs.ai](#)). Pricing is also unusually volatile in 2026: on-demand H100 rates fell sharply through 2025 as AWS cut prices by roughly 44 percent that June and competition from specialist clouds intensified, then reversed course as a fresh supply squeeze, driven by high-bandwidth memory (HBM) shortages rather than GPU die availability, pushed contract pricing back up through the first half of 2026 (Source: [intuitionlabs.ai](#)) (Source: [electropages.com](#)). Every figure in this report is anchored to a specific measurement date, generally between March and July 2026, because a rate quoted without a date is close to meaningless in a market moving this fast.

This report is written for technology and operations leaders evaluating AI infrastructure budgets, including teams in regulated, data-intensive fields such as life sciences, where GPU-backed workloads now span genomic sequence modeling, molecular simulation, and clinical natural-language processing pipelines that were not economically feasible on general-purpose compute even three years ago. IntuitionLabs, a life-sciences and AI consultancy, tracks GPU pricing as part of its own infrastructure cost research and has independently published H100 rental comparisons and hardware acquisition cost guides that corroborate the figures assembled here, though it is not itself a GPU cloud or hardware vendor and does not appear as a pricing option in the comparisons that follow (Source: [intuitionlabs.ai](#)).

The Cloud GPU Rental Market in 2026: Pricing by Provider and GPU Model

Cloud GPU pricing in 2026 sorts cleanly into four provider tiers, and the gap between the cheapest and most expensive tier is the single largest lever available to a buyer, larger than the choice between GPU generations. Hyperscalers (AWS, Azure, Google Cloud, Oracle) charge the most because their on-demand rate bundles a broad managed-services ecosystem, global compliance certifications, and enterprise support that many AI workloads do not strictly require, while specialist "neoclouds" that run leaner software stacks and price aggressively for reserved-capacity deals routinely undercut hyperscaler on-demand rates by 40 to 70 percent for the same chip (Source: [gpcostcompare.com](#)).

Table 1 below assembles verified on-demand, per-GPU-hour H100 pricing from official provider pages and independent trackers as of mid-2026, normalized to a single-GPU basis where providers only sell multi-GPU instances.

Table 1: NVIDIA H100 SXM on-demand cloud pricing by provider, normalized to \$/GPU-hour (verified March to July 2026)

PROVIDER	TIER	\$/GPU-HOUR (ON-DEMAND)	NOTES
Vast.ai	Marketplace	\$1.49	Interruptible marketplace floor (Source: gpucloudcost.com)
ThunderCompute	Long-tail	\$1.38	Virtual GPU model, shared infrastructure (Source: deploybase.ai)
RunPod	Long-tail	\$1.99 to \$2.89	PCIe Community Cloud to Secure Cloud SXM (Source: intuitionlabs.ai)
Together AI	Specialist, reserved	\$3.09	Reserved cluster, 91 to 180 day term (Source: gpucloudcost.com)
Nebius	Specialist	\$3.85	HGX H100 80GB, 3.2 Tbps InfiniBand (Source: gpcostcompare.com)
Lambda	Specialist	\$3.99	8x H100 SXM5 node, on-demand (Source: gpcostcompare.com)
CoreWeave	Specialist	\$6.16	HGX H100 8x node, per-GPU list rate (Source: gpcostcompare.com)
AWS EC2 (p5)	Hyperscaler	\$6.88	On-demand, post mid-2025 price cut (Source: gpucloudcost.com)
Oracle Cloud (OCI)	Hyperscaler	\$10.00	Bare-metal 8x H100 node (Source: gpcostcompare.com)
Google Cloud (a3-highgpu-8g)	Hyperscaler	\$11.06	On-demand list, official pricing page (Source: cloud.google.com)
Microsoft Azure (ND H100 v5)	Hyperscaler	\$12.29	On-demand, highest tracked list rate (Source: gpcostcompare.com)

Reading the table, the 8.9x spread between Vast.ai's marketplace floor and Azure's list rate is not a data error; it is the defining fact of the 2026 GPU market. Multi-GPU hyperscaler instances also bundle CPU, RAM, high-throughput networking, and NVLink interconnect into the quoted price, which pushes the effective per-GPU cost higher than bare single-GPU rental even before accounting for service overhead (Source: deploybase.ai). A comparable analysis normalizing full 8-GPU instances to \$/GPU-hour found AWS at \$6.88, Oracle at \$10.00, Google Cloud at \$11.62, and Azure at \$12.29, confirming the ranking is consistent across independent trackers even when the exact figures differ by a few cents (Source: gpcostcompare.com).

Answering "**nvidia h100 gpu price 2026**" directly: as of mid-2026, expect \$1.38 to \$3.50 per GPU-hour on marketplace and long-tail clouds, \$3.85 to \$6.16 per GPU-hour on specialist neoclouds, and \$6.88 to \$12.29 per GPU-hour on hyperscaler on-demand instances, an overall range documented consistently by DeployBase (\$1.38 to \$11.68 across 28+ providers, March 2026), GPUCloudCost.com (\$1.49 to \$12.29 across 17 providers, June 2026), and CloudZero (\$1.38 to \$8.00+, with a market median of \$2.29 to \$3.12), with AWS's own instance page confirming the P5 family as the delivery vehicle at hyperscaler scale (Source: deploybase.ai) (Source: cloudzero.com) (Source: aws.amazon.com).

For "**h200 vs b200 pricing**," the two chips occupy different rungs of the same ladder rather than competing directly. The H200 shares the H100's Hopper compute architecture and 1,979 TFLOPS FP8 throughput but nearly triples usable memory to 141GB of HBM3e at 4.8 TB/s bandwidth, a change that eliminates the need for tensor parallelism on models up to roughly 100 billion parameters (Source: gputaas.com). That memory premium translates into roughly 25 to 30 percent higher hourly pricing, with H200 SXM on-demand running \$4.50 to \$7.00 per GPU-hour on hyperscalers, \$3.50 to \$4.50 on specialist clouds, and \$2.50 to \$3.50 on long-tail providers (Source: mercatus-ai.com). Google Cloud's own pricing page confirms this tier directly: the a3-ultragpu-8g instance, built on 8x H200, lists at \$84.81 per hour on-demand, roughly \$10.60 per GPU, while the Blackwell-

based a4-highgpu-8g (8x B200) is not sold as a standard on-demand SKU but carries a published spot rate of \$39.63 per hour, roughly \$4.95 per GPU (Source: cloud.google.com) (Source: cloud.google.com). The B200 itself is an entirely different chip: NVIDIA's Blackwell architecture packs 208 billion transistors in a dual-die design, 192GB of HBM3e at 8.0 TB/s, and native FP4 precision support that roughly 2.5x to 4.5x's inference throughput over Hopper (Source: gputaas.com). B200 on-demand pricing runs \$2.12 spot to \$4.99 to \$6.04 per GPU-hour at specialist providers, with 36-month reserved contracts as low as \$2.25 per GPU-hour, while AWS's Blackwell instance list rate reaches roughly \$14.24 per GPU-hour (Source: gputaas.com). Because B200 supply is still constrained relative to Hopper, most teams in 2026 are advised to buy or reserve H100 or H200 capacity today and evaluate Blackwell on a two-year upgrade cycle rather than wait (Source: mercatus-ai.com).

At the top of the price ladder sits the **GB200 NVL72**, NVIDIA's rack-scale Grace Blackwell system that is not rented as individual GPUs but as Superchip or rack-node allocations. CoreWeave lists GB200 NVL72 on-demand access at roughly \$10.50 per GPU-hour through 4-GPU instances, Oracle Cloud prices bare-metal access around \$16 per GPU-hour, and Azure's ND GB200 v6, generally available since March 2025, runs approximately \$27 per GPU-hour (Source: spheron.network) (Source: spheron.network). Scaled to a full 72-GPU rack, that range implies \$756 to \$1,944 per hour, 1.7x to 4.5x the cost of the same 72 GPUs deployed as nine separate 8xB200 nodes, a premium that only pays back when a workload's inter-GPU communication overhead is large enough that the NVL72's 130 TB/s all-to-all NVLink fabric meaningfully cuts training or serving time (Source: spheron.network).

Legacy Ampere-generation **A100** GPUs, NVIDIA's pre-Hopper flagship, remain in active commercial use for cost-sensitive fine-tuning and inference that does not require FP8 acceleration. A100 80GB on-demand pricing runs \$1.50 to \$3.00 per GPU-hour on hyperscalers, some of which have phased the chip out entirely, \$1.10 to \$1.50 on specialist clouds, and \$0.80 to \$1.10 on long-tail providers (Source: mercatus-ai.com). Cloud A100 rates of roughly \$1.29 to \$2.50 per hour remain cheaper than H100 rates on paper, but the H100 delivers three to five times better throughput on transformer workloads, so an A100 job that takes 100 hours can still cost more than an H100 job that takes 30, even at the lower hourly rate, which is why cost per completed run, not cost per hour, should drive the comparison (Source: cloudzero.com).

Hardware Acquisition Costs: Buying H100, H200, B200, and GB200 Systems Outright

For organizations planning owned infrastructure rather than cloud rental, the acquisition-cost picture answers "cost to build an ai data center" at both the chip and facility level. *Table 2* summarizes verified purchase pricing for individual GPUs, complete 8-GPU systems, and rack-scale Blackwell hardware as of Q1 to Q2 2026.

Table 2: Data center GPU and system purchase costs, new and used (verified Q1 to Q2 2026)

ITEM	PRICE RANGE	SOURCE DETAIL
H100 PCIe 80GB, new	\$25,000 to \$30,000	Standard PCIe slot, no NVLink (Source: deploybase.ai)
H100 SXM5 80GB, new	\$35,000 to \$40,000	NVLink 4.0, 900 GB/s interconnect (Source: cloudzero.com)
H100, used or secondary market	\$6,000 to \$22,000	85% below the 2023 peak of \$40,000 (Source: cloudzero.com)
DGX H100, 8-GPU system	\$300,000 to \$460,000	Full system: CPUs, networking, storage (Source: deploybase.ai)
DGX H200, 8-GPU system	\$400,000 to \$500,000	Quoted range per market analysis (Source: intuitionlabs.ai)
GB200 NVL72 rack (72 GPUs)	\$756 to \$1,944/hour cloud-equivalent	No single retail price published; sold as rack-node cloud allocations (Source: spheron.network)

The 85 percent collapse in used H100 pricing deserves emphasis: cards that traded for \$40,000 in late 2023 changed hands for as little as \$6,000 by mid-2026, a decline driven by Blackwell's arrival making Hopper-class hardware progressively less economical for inference workloads even as demand for AI compute overall keeps climbing (Source: cloudzero.com). At the same time, new H100 street pricing has not collapsed nearly as far, because HBM3 memory and packaging costs remain a hard floor under manufacturing cost regardless of secondary-market discounting. Viewed

across the full product lineup, new-card pricing runs roughly \$10,000 to \$15,000 for an A100, \$25,000 to \$40,000 for an H100, and \$30,000 to \$50,000 for the newest B200 cards, a curve that tracks each generation's memory bandwidth and precision-format gains more closely than it tracks raw compute alone (Source: cloudzero.com).

Buy-versus-rent math depends heavily on utilization. At a representative \$25,000 purchase price and a \$3.00 per hour rental benchmark, the raw break-even point is roughly 8,333 hours, about 347 days of continuous use, but that calculation ignores power, cooling, and rack infrastructure, which shift the realistic break-even to 18 months or more of near-100 percent utilization (Source: cloudzero.com). A single SXM5 GPU pulls up to 700 watts, meaning a full 8-GPU node draws roughly 5.6 kilowatts under load, translating to real annual electricity costs in addition to \$10,000 to \$100,000 in incremental cooling and power infrastructure per rack (Source: deploybase.ai). Purchasing tends to make financial sense only when GPUs run 24/7 for 18 or more months, the buyer already has (or budgets \$400,000-plus for) data center infrastructure, and in-house staff can manage firmware, drivers, and liquid cooling, a profile that describes hyperscalers and large AI labs far more than mid-market engineering teams (Source: cloudzero.com).

Beyond the chips themselves, the facility that houses them is now the larger line item for anyone building at scale. Conventional data centers cost \$7 to \$12 million per megawatt of capacity, but AI-optimized facilities, which must accommodate GPU-dense racks, liquid cooling, and high-voltage power distribution, run \$20 million or more per megawatt, and hyperscale AI campuses at gigawatt scale are modeled at \$45 to \$55 billion per gigawatt for a fully built-out ecosystem (Source: gigaenergy.com) (Source: gigaenergy.com). That construction budget is separate from the IT equipment inside it: at 100 megawatts of capacity, servers, networking gear, and GPUs alone can run \$2.5 to \$4 billion, and transformer lead times from legacy manufacturers can stretch past 50 weeks, a supply constraint that increasingly rivals GPU allocation as the binding factor on how quickly new AI capacity can come online (Source: gigaenergy.com) (Source: gigaenergy.com). Independent trackers segment likely buyers by scale: a small team running one to four GPUs for inference and fine-tuning should budget roughly \$500 to \$5,000 per month, a mid-sized cluster of 8 to 64 H100 or A100 GPUs runs \$15,000 to \$250,000 per month, and a large cluster of 256 or more H100, H200, or B200 GPUs runs \$1 million to \$40 million or more per month (Source: gpucloudcost.com) (Source: gpucloudcost.com).

Commitment Structures: On-Demand, Reserved, and Spot Pricing Economics

Beyond provider choice, the contract structure a buyer selects moves the effective hourly rate by 30 to 80 percent, which is why two organizations renting the identical H100 SXM chip from the identical provider can pay materially different bills. Three structures dominate the 2026 market, and understanding each is central to answering "cloud gpu rental pricing comparison" and "ai gpu cost per hour" in practical terms rather than headline list-price terms.

On-demand pricing carries no commitment and the highest per-hour rate; it suits unpredictable workloads, unproven providers, or short evaluation windows. **Reserved** capacity, typically committed for one to three years, cuts cost 25 to 50 percent versus on-demand at the same provider, with one-year terms averaging 25 to 35 percent savings and three-year terms averaging 35 to 50 percent, rising past 50 percent for custom five-year-plus commitments (Source: mercatus-ai.com). **Spot** or preemptible instances undercut on-demand by 30 to 80 percent in exchange for the risk of short-notice reclamation, typically a two-minute eviction warning on AWS, which makes them suitable for checkpointed training runs and batch workloads but poor fits for production inference with uptime requirements (Source: deploybase.ai). Google Cloud's own preemptible A3 instances illustrate the discount structure directly, published at roughly 60 to 91 percent off on-demand pricing depending on availability (Source: deploybase.ai).

Table 3 breaks out H100 pricing across all three structures and all four provider tiers, drawn from a May 2026 market analysis.

Table 3: H100 SXM5 pricing by provider tier and commitment structure, \$/GPU-hour (May 2026)

PROVIDER TIER	ON-DEMAND	RESERVED (1-YEAR)	RESERVED (3-YEAR)	SPOT
Hyperscaler (AWS p5, Azure ND H100 v5, GCP A3)	\$3.50 to \$5.00	\$2.80 to \$3.80	\$2.20 to \$3.00	\$1.50 to \$3.00
Specialist (CoreWeave, Lambda)	\$2.50 to \$3.50	\$2.00 to \$2.80	\$1.70 to \$2.30	\$1.20 to \$2.00
Long-tail	\$1.99 to \$2.50	\$1.60 to \$2.10	\$1.30 to \$1.80	\$0.80 to \$1.50
Decentralized networks	\$1.50 to \$2.50	n/a	n/a	\$0.60 to \$1.50

(Source: mercatus-ai.com)

The most consequential number in this table is the bottom-left corner: three-year reserved H100 capacity at long-tail providers, \$1.30 to \$1.80 per GPU-hour, sits within 15 to 25 percent of a comparably utilized owned cluster's effective per-hour cost, which is why cloud rental at the right provider has come to dominate the buy-versus-build decision for all but the very largest AI labs; the operational simplicity of cloud is nearly free once the price gap narrows this far (Source: mercatus-ai.com). A cost-component breakdown of a representative long-tail versus hyperscaler H100 rate shows why the gap persists even though hardware amortization is nearly identical across tiers (\$0.55 per GPU-hour either way): hyperscalers absorb \$0.30 to \$0.60 in operations and support overhead versus \$0.05 to \$0.15 at long-tail providers, and \$0.50 to \$1.00 in sales and marketing overhead versus \$0.05 to \$0.10, with margin itself running \$1.00 to \$2.00 per hour at hyperscalers against \$0.20 to \$0.50 at long-tail providers (Source: mercatus-ai.com) (Source: mercatus-ai.com). Headline hourly rates also understate total spend: real-world cloud bills typically run 30 to 50 percent above the quoted GPU sticker once egress, storage, orchestration, and support tiers are added, so the rate in Table 3 is a floor, not a ceiling, on the actual invoice (Source: mercatus-ai.com). In other words, the headline GPU rate at a hyperscaler is roughly 60 to 70 percent hardware and infrastructure cost and 30 to 40 percent overhead and margin that a leaner specialist provider simply does not carry.

This structure also governs **"gpu prices for llm training."** A full fine-tune of a 7-billion-parameter Mistral model on a single H100 SXM at \$2.69 per hour for 20 hours costs \$53.80 in raw compute, while the same task using parameter-efficient LoRA fine-tuning on a cheaper A100 at \$1.39 per hour for the same duration costs \$27.80, roughly half, for broadly comparable output quality (Source: deploybase.ai). At the median H100 cloud rate, a single reserved GPU running continuously costs roughly \$1,964 per month, \$23,567 annually, and \$70,700 over three years, a baseline that scales linearly (before volume discounts) to the multi-hundred-GPU clusters that frontier model training runs require (Source: deploybase.ai). Right-sizing GPU form factor also matters more than most teams expect: an inference workload running on SXM at \$2.69 per hour that does not actually need NVLink multi-GPU interconnect can move to PCIe at \$1.99 per hour and save 26 percent with no measured performance loss, since the SXM premium only earns its keep when multi-GPU communication bandwidth is actually the bottleneck (Source: deploybase.ai).

Data Analysis and Evidence

The quantitative backdrop to every price quoted above is a demand curve that is still steepening in 2026, not flattening, which is the central fact behind **"ai infrastructure cost trends 2026."** NVIDIA's own financial results are the clearest signal: for its fiscal first quarter of 2027, the three months ended April 26, 2026, the company reported record total revenue of \$81.6 billion, up 85 percent year over year and 20 percent sequentially, with GAAP gross margin of 74.9 percent (Source: nvidianews.nvidia.com). Data Center revenue specifically, the segment that includes H100, H200, B200, and GB200 sales, hit a record \$75.2 billion, up 92 percent year over year, and within that, Data Center compute revenue alone reached \$60.4 billion (up 77 percent year over year) while Data Center networking revenue reached \$14.8 billion (up 199 percent year over year), evidence that buyers are increasingly paying as much for the interconnect fabric around GPUs as for the chips themselves (Source: nvidianews.nvidia.com). NVIDIA guided its following quarter to \$91.0 billion in revenue, plus or minus 2 percent, explicitly excluding any Data Center compute revenue from China given ongoing export restrictions (Source: nvidianews.nvidia.com).



On the buyer side, the four largest hyperscalers disclosed a combined \$630 billion in planned 2026 capital expenditures during their respective early-2026 earnings calls, a 62 percent increase over the record \$388 billion the same four companies spent in 2025 (Source: datacenter richness.substack.com). Individually: Amazon guided to \$200 billion for 2026, up from \$125 billion in 2025; Google guided to \$175 to \$185 billion, up from \$91 billion; Meta guided to \$115 to \$135 billion, up from \$72 billion; and Microsoft guided to \$110 to \$120 billion, up from \$90 billion (Source: datacenter richness.substack.com) (Source: datacenter richness.substack.com). Most of that capital converts directly into GPU purchases and the data center shells that house them, which is the mechanical link between hyperscaler earnings calls and the cloud rental prices in Table 1. AWS's own P5 documentation illustrates the resulting scale directly, describing EC2 UltraClusters that interconnect up to 20,000 H100 or H200 GPUs on a petabit-scale nonblocking network capable of up to 20 exaflops of aggregate compute, with Anthropic among the early customers stating publicly that it expected the P5 generation "to deliver substantial price-performance benefits" over the prior GPU generation at that scale (Source: aws.amazon.com) (Source: aws.amazon.com).

Supply-side scarcity compounds the demand story. Industry trackers reported H100 and H200 contract pricing climbing roughly 40 percent between October 2025 and March 2026, driven primarily by HBM3e memory cost pass-throughs from Samsung and SK Hynix rather than by GPU die shortages, with roughly half of tracked specialist providers reporting no Hopper-class capacity coming off contract at all during that window (Source: gpuaas.com). This is the substance behind "**gpu shortage pricing impact**": lead times for Blackwell PRO-class GPUs stretched to three to seven months by early 2026, RTX 6000 Ada 48GB cards rose to \$7,400 to \$7,800, and RTX PRO 6000 96GB cards climbed to \$9,450 to \$9,800, all reflecting a memory-constrained rather than silicon-constrained bottleneck (Source: electropages.com). The scarcity is broad-based rather than confined to the newest chips: sourcing data cited by trade press shows NVIDIA's H200 NVL PCIe and H100 NVL PCIe among the most requested parts even as buyers simultaneously widen their search to older Ada-generation and legacy cards to find any available inventory (Source: electropages.com). Analysts distinguish this cycle explicitly from the 2021-2022 GPU shortage, which was driven by speculative cryptocurrency mining demand that eventually collapsed; the 2026 tightness is described as structural, tied to persistent AI infrastructure investment cycles rather than a speculative bubble that self-corrects (Source: electropages.com).

Export controls add a further pricing distortion. The Financial Times, in reporting cited by Reuters, found that NVIDIA's export-restricted AI chips more than doubled in price on China's black market over the first half of 2026, based on interviews with multiple Chinese chip traders, illustrating that geopolitical restrictions on legitimate sales channels create their own parallel pricing structure entirely disconnected from the on-demand and reserved rates in this report's core tables (Source: reuters.com).

Case Studies and Real-World Examples

xAI's Colossus Cluster: Speed as a Pricing and Deployment Strategy

Elon Musk's xAI built the clearest illustration of what pricing looks like when speed, not unit cost, is the binding constraint. Quoted 18 to 24 months for conventional data center construction, xAI instead repurposed a former Electrolux appliance factory in Memphis, Tennessee, and deployed 100,000 NVIDIA H100 GPUs there in 122 days, then doubled the cluster to 200,000 GPUs in a further 92 days (Source: [introl.com](https://www.introl.com)) (Source: [introl.com](https://www.introl.com)). By December 2025, the cluster comprised 150,000 H100, 50,000 H200, and 30,000 GB200 GPUs, described by the deployment's infrastructure partner as the largest fully operational, single-coherent AI training cluster in the world (Source: [introl.com](https://www.introl.com)). Powering that density required 35 gas turbines capable of 420 megawatts alongside 208 Tesla Megapack battery units and a 150-megawatt utility substation built in 97 days versus the roughly 2.5 years such projects normally require, with the facility drawing approximately 250 megawatts by late 2025 and xAI separately committing \$80 million to a wastewater recycling facility to support cooling at that scale (Source: [introl.com](https://www.introl.com)). Networking relied on NVIDIA's Spectrum-X Ethernet fabric rather than the InfiniBand standard used by most large training clusters, achieving 95 percent data throughput against the roughly 60 percent typical of standard Ethernet at this scale, evidence that Ethernet-based fabrics can now support frontier-scale training economically (Source: [introl.com](https://www.introl.com)).

CoreWeave's OpenAI and Meta Contracts: Pricing Frontier-Scale Capacity

CoreWeave, a specialist GPU cloud provider, illustrates how frontier AI labs now lock in capacity years in advance through multibillion-dollar bilateral contracts rather than spot-market rental. CoreWeave's relationship with OpenAI began with an \$11.9 billion, five-year agreement in March 2025, expanded by \$4 billion in May 2025, then expanded again by \$6.5 billion in September 2025, bringing the cumulative contract value to \$22.4 billion (Source: [cnbc.com](https://www.cnbc.com)) (Source: [cnbc.com](https://www.cnbc.com)) (Source: [cnbc.com](https://www.cnbc.com)). Separately, in April 2026, CoreWeave and Meta Platforms announced an expanded, long-term agreement to provide AI cloud capacity through December 2032 for approximately \$21 billion, with the dedicated capacity to include some of the earliest commercial deployments of NVIDIA's next-generation Vera Rubin platform, a deal CoreWeave itself described as "a clear signal of the industry's" accelerating demand for high-performance AI infrastructure capacity (Source: [coreweave.com](https://www.coreweave.com)) (Source: [coreweave.com](https://www.coreweave.com)) (Source: [coreweave.com](https://www.coreweave.com)). By the time of its 2025 public listing prospectus, CoreWeave operated 32 data centers powered by more than 250,000 NVIDIA GPUs across the United States and parts of Europe, and separately committed \$6 billion to a 100-megawatt data center in Lancaster, Pennsylvania (Source: [cnbc.com](https://www.cnbc.com)). These contracts illustrate that at the frontier-lab scale, GPU access is priced and negotiated as multi-year committed capacity rather than the hourly rates that dominate Tables 1 through 3, but the underlying economics, reserved capacity beating on-demand, are the same principle scaled up by roughly six orders of magnitude.

Stargate: A \$500 Billion Commitment to Physical AI Infrastructure

OpenAI, Oracle, and SoftBank's Stargate initiative demonstrates how GPU demand cascades into gigawatt-scale physical construction commitments. First announced in January 2025 at \$500 billion and 10 gigawatts of planned capacity, the partnership announced five additional U.S. data center sites in October 2025 that brought total planned capacity to nearly 7 gigawatts and over \$400 billion in committed investment over three years, putting the project ahead of schedule to secure its full \$500 billion commitment by the end of 2025 (Source: [openai.com](https://www.openai.com)) (Source: [openai.com](https://www.openai.com)). A separate July 2025 agreement between OpenAI and Oracle to develop up to 4.5 additional gigawatts of Stargate capacity was itself valued at more than \$300 billion over five years (Source: [openai.com](https://www.openai.com)). At the flagship Abilene, Texas campus, Oracle began delivering the first NVIDIA GB200 racks in June 2025 and had already started early training and inference workloads on that capacity by the time of the announcement, a concrete example of GB200-generation hardware moving from announcement to production inside roughly a year (Source: [openai.com](https://www.openai.com)). Separately, NVIDIA itself announced a \$100 billion investment in OpenAI to support the buildout, and Microsoft committed \$4 billion to a second Wisconsin data center in the same period, illustrating how capital now flows in multiple directions between chip suppliers, cloud providers, and AI labs simultaneously (Source: [cnbc.com](https://www.cnbc.com)).

China's Export-Restricted Chip Black Market: Pricing Under Sanctions

The fourth case illustrates what GPU pricing looks like when the legal supply channel is closed by policy rather than by capacity. With U.S. export controls barring the sale of NVIDIA's most advanced AI chips to China, a parallel black market has formed, and the Financial Times, in reporting relayed by Reuters in June 2026, found that prices for these restricted chips had more than doubled over the preceding six months based on interviews with multiple Chinese chip traders (Source: [reuters.com](https://www.reuters.com)). This case matters for the broader pricing index because it demonstrates that the supply-and-demand dynamics documented in Tables 1 through 3, driven by memory shortages, hyperscaler capex, and reserved-versus-spot economics, operate even more acutely once legitimate market access is restricted, and it explains why NVIDIA's own Q2 fiscal 2027 revenue guidance explicitly assumed zero Data Center compute revenue from China rather than model a highly uncertain, sanctioned market (Source: [nvidianews.nvidia.com](https://www.nvidianews.nvidia.com)).

Implications and Future Directions

Three structural forces will keep shaping data center GPU pricing through the remainder of 2026 and into 2027. First, memory, not GPU die fabrication, has become the binding supply constraint: HBM3e cost pass-throughs from Samsung and SK Hynix were the primary driver of the roughly 40 percent H100 and H200 contract price increase between October 2025 and March 2026, and unless HBM output expands materially, price pressure on Hopper and Blackwell-class GPUs alike is likely to persist even as GPU die supply itself improves (Source: gputaas.com). Second, the buy-versus-rent calculus continues to favor rental for all but the largest, most utilization-disciplined organizations, since the reserved, long-tail rate of \$1.30 to \$1.80 per GPU-hour now sits within striking distance of owned-cluster economics without the operational burden of managing firmware, cooling, and depreciation directly (Source: mercatus-ai.com). Third, physical infrastructure, power procurement, transformer lead times exceeding 50 weeks, and construction costs of \$20 million or more per megawatt, is increasingly the true bottleneck on how fast new GPU capacity can be deployed, a shift visible in xAI's turn to gas turbines and Tesla Megapacks and in Stargate's framing of its commitment in gigawatts rather than GPU counts (Source: gigaenergy.com) (Source: openai.com).

For buyers outside the hyperscaler tier, including mid-market AI teams, research labs, and specialized enterprise deployments in fields such as life sciences that increasingly run genomic foundation models, molecular dynamics simulations, and large-scale clinical text pipelines on rented GPU capacity, three practical shifts follow from this data. Reserved and long-tail capacity now beats hyperscaler on-demand pricing by wide enough margins, 40 to 70 percent in many comparisons, that provider selection alone can swing an annual GPU budget by hundreds of thousands of dollars at modest cluster sizes (Source: mercatus-ai.com). Workload-to-hardware matching is the second lever: a team serving a 70-billion-parameter model in FP16 needs roughly 140GB of VRAM, which a single H200 covers at \$3.72 to \$6.00 per hour where an equivalent H100 deployment would need two GPUs billed at double the rate, and choosing PCIe over SXM for single-GPU inference that does not need NVLink saves a further 20 to 26 percent with no measurable performance cost (Source: gputaas.com) (Source: deploybase.ai). Consultancies advising on AI infrastructure procurement, including firms serving regulated life-sciences organizations, are increasingly treating GPU cost modeling as a first-class part of technology strategy rather than an afterthought delegated entirely to engineering, given that the same workload can cost 8 to 9 times more or less depending purely on which of the four provider tiers is selected; one such firm's own independent review of the H100 rental market, conducted separately from the trackers cited elsewhere in this report, found rates spanning \$1.49 to \$6.98 per hour across more than fifteen providers, a range consistent with every other source in this index despite being compiled independently (Source: intuitionlabs.ai).

Looking forward, the Blackwell Ultra generation (B300) and NVIDIA's announced Vera Rubin platform are already entering early commercial deployment. At least one specialist provider prices B300 access at \$8.55 per hour on-demand or \$3.67 per hour spot, and CoreWeave and Meta's April 2026 agreement specifically calls out Vera Rubin as part of the contracted capacity, suggesting the pricing ladder documented in this report will gain at least one more rung, with correspondingly higher price ceilings, before 2026 is over (Source: spheron.network) (Source: coreweave.com). Whether that new generation compresses Hopper and early Blackwell pricing further, the way B200 arrival already discounted secondary-market H100s by 85 percent, or simply adds a new premium tier on top of an already-expensive market, will depend largely on whether HBM memory supply, the binding constraint identified throughout this report, expands fast enough to keep pace with hyperscaler capital spending that shows no sign of decelerating in 2026.

Conclusion

Data center GPU pricing in 2026 is not one number but a matrix defined by GPU generation, provider tier, and contract commitment, and every axis of that matrix moved this year. An H100 costs anywhere from \$1.38 to \$12.29 per GPU-hour to rent and \$25,000 to \$40,000 to buy outright, an H200 carries a 25 to 30 percent premium for nearly triple the memory, and a B200 or GB200 NVL72 rack sits meaningfully higher still, with rack-scale Blackwell capacity running \$756 to \$1,944 per hour. Reserved, long-tail contracts at \$1.30 to \$1.80 per GPU-hour now approach owned-cluster economics closely enough that cloud rental dominates buyer decisions outside the hyperscaler tier, while hyperscaler on-demand rates, still the highest in the market by a factor of roughly four to nine, buy managed services and compliance infrastructure that many workloads do not strictly need.

Two forces will keep this market volatile through the rest of 2026. On the supply side, HBM memory scarcity, not GPU die availability, is now the primary price driver, pushing H100 and H200 contract pricing up roughly 40 percent in a six-month window even as older-generation secondary-market pricing collapsed 85 percent over the same broader period. On the demand side, NVIDIA's record \$75.2 billion quarterly Data Center revenue and the four largest hyperscalers' combined \$630 billion 2026 capital expenditure plan show no evidence of near-term deceleration, and multibillion-dollar bilateral contracts between labs like OpenAI and Meta and infrastructure providers like CoreWeave suggest the largest buyers are now locking in capacity years ahead rather than transacting at spot rates at all.

For any organization budgeting AI infrastructure in the second half of 2026, the actionable conclusion is straightforward: benchmark provider tier and contract length before benchmarking GPU generation, since the gap between the cheapest and most expensive channel for identical hardware routinely exceeds the price difference between an H100 and an H200. Providers, contract structures, and even geography now separate the same silicon by an order of magnitude, and that spread, more than any single vendor's list price, is the defining fact of the 2026 data center GPU market.

Frequently Asked Questions (FAQs)

What does an NVIDIA H100 cost per hour in 2026? Verified on-demand rates range from \$1.38 to \$1.49 per GPU-hour on marketplace and long-tail clouds up to \$11.06 to \$12.29 per GPU-hour on Google Cloud and Azure, with specialist clouds such as CoreWeave, Lambda, and Nebius priced between \$3.85 and \$6.16 per GPU-hour (Source: [deploybase.ai](#)) (Source: [gpuccostcompare.com](#)).

How much does it cost to rent versus buy a GPU cluster for LLM training? A single H100 at the median cloud rate of \$2.69 per hour costs roughly \$1,964 per month or \$23,567 per year to rent continuously; buying the same GPU outright costs \$25,000 to \$40,000 plus power and cooling, with realistic break-even at 18 months or more of near-full utilization (Source: [deploybase.ai](#)) (Source: [cloudzero.com](#)).

Is the H200 worth the price premium over the H100? The H200 costs roughly 25 to 30 percent more per hour but delivers 76 percent more memory and 43 percent more bandwidth at identical compute throughput, which pays off for models above roughly 40 billion parameters or workloads with long context windows, while smaller models remain more cost-effective on H100 (Source: [mercatus-ai.com](#)).

What does an NVIDIA A100 cost compared to newer GPUs? A100 80GB on-demand pricing runs \$0.80 to \$3.00 per GPU-hour depending on provider tier, cheaper than H100 or H200, but the H100 delivers three to five times better throughput on transformer workloads, so the A100 remains attractive mainly for fine-tuning and inference tasks that do not need FP8 acceleration (Source: [mercatus-ai.com](#)) (Source: [cloudzero.com](#)).

Why is there such a large gap between the cheapest and most expensive GPU cloud providers? Hardware amortization costs are nearly identical across provider tiers, but hyperscalers layer on far higher sales, support, and margin overhead, roughly \$1.00 to \$2.00 per GPU-hour in margin alone versus \$0.20 to \$0.50 at long-tail providers, which is why identical silicon can differ in price by 8 to 9 times (Source: [mercatus-ai.com](#)).

How much does it cost to build an AI data center in 2026? AI-optimized facilities run \$20 million or more per megawatt of capacity, versus \$7 to \$12 million per megawatt for conventional data centers, with hyperscale campuses at gigawatt scale modeled at \$45 to \$55 billion per gigawatt, plus \$2.5 to \$4 billion in servers, networking, and GPUs per 100 megawatts of capacity (Source: [gigaenergy.com](#)) (Source: [gigaenergy.com](#)).

Is the current GPU shortage likely to ease in 2026? Available evidence suggests no near-term easing. The shortage is driven by structural HBM memory constraints tied to sustained AI infrastructure investment rather than a speculative demand spike, which analysts distinguish explicitly from the 2021-2022 cryptocurrency-driven shortage that eventually collapsed on its own (Source: [electropages.com](#)).

Tags: data center gpu pricing, nvidia h100 price, gpu cloud pricing, h100 vs h200 vs b200, gpu rental cost per hour, ai infrastructure cost, gb200 nvl72 pricing, gpu shortage 2026, cloud gpu comparison

IntuitionLabs - Industry Leadership & Services

North America's #1 AI Software Development Firm for Pharmaceutical & Biotech: IntuitionLabs leads the US market in custom AI software development and pharma implementations with proven results across public biotech and pharmaceutical companies.

Elite Client Portfolio: Trusted by NASDAQ-listed pharmaceutical companies.

Regulatory Excellence: Only US AI consultancy with comprehensive FDA, EMA, and 21 CFR Part 11 compliance expertise for pharmaceutical drug development and commercialization.

Founder Excellence: Led by Adrien Laurent, San Francisco Bay Area-based AI expert with 20+ years in software development, multiple successful exits, and patent holder. Recognized as one of the top AI experts in the USA.

Custom AI Software Development: Build tailored pharmaceutical AI applications, custom CRMs, chatbots, and ERP systems with advanced analytics and regulatory compliance capabilities.

Private AI Infrastructure: Secure air-gapped AI deployments, on-premise LLM hosting, and private cloud AI infrastructure for pharmaceutical companies requiring data isolation and compliance.

Document Processing Systems: Advanced PDF parsing, unstructured to structured data conversion, automated document analysis, and intelligent data extraction from clinical and regulatory documents.

Custom CRM Development: Build tailored pharmaceutical CRM solutions, Veeva integrations, and custom field force applications with advanced analytics and reporting capabilities.

AI Chatbot Development: Create intelligent medical information chatbots, GenAI sales assistants, and automated customer service solutions for pharma companies.

Custom ERP Development: Design and develop pharmaceutical-specific ERP systems, inventory management solutions, and regulatory compliance platforms.

Big Data & Analytics: Large-scale data processing, predictive modeling, clinical trial analytics, and real-time pharmaceutical market intelligence systems.

Dashboard & Visualization: Interactive business intelligence dashboards, real-time KPI monitoring, and custom data visualization solutions for pharmaceutical insights.

Leading Claude & ChatGPT AI Enablement and Training: Help biotech and pharmaceutical teams adopt Claude and ChatGPT through practical training, governed workflows, use-case development, and implementation designed for regulated environments.

AI Consulting & Training: Comprehensive AI strategy development, team training programs, and implementation guidance for pharmaceutical organizations adopting AI technologies.

Contact founder Adrien Laurent and team at intuitionlabs.ai/contact for a consultation.

DISCLAIMER

The information contained in this document is provided for educational and informational purposes only. We make no representations or warranties of any kind, express or implied, about the completeness, accuracy, reliability, suitability, or availability of the information contained herein.

Any reliance you place on such information is strictly at your own risk. In no event will IntuitionLabs.ai or its representatives be liable for any loss or damage including without limitation, indirect or consequential loss or damage, or any loss or damage whatsoever arising from the use of information presented in this document.

This document may contain content generated with the assistance of artificial intelligence technologies. AI-generated content may contain errors, omissions, or inaccuracies. Readers are advised to independently verify any critical information before acting upon it.

All product names, logos, brands, trademarks, and registered trademarks mentioned in this document are the property of their respective owners. All company, product, and service names used in this document are for identification purposes only. Use of these names, logos, trademarks, and brands does not imply endorsement by the respective trademark holders.

IntuitionLabs.ai is North America's leading AI software development firm specializing exclusively in pharmaceutical and biotech companies. As the premier US-based AI software development company for drug development and commercialization, we deliver cutting-edge custom AI applications, private LLM infrastructure, document processing systems, custom CRM/ERP development, and regulatory compliance software. Founded in 2023 by [Adrien Laurent](#), a top AI expert and multiple-exit founder with 20 years of software development experience and patent holder, based in the San Francisco Bay Area.

This document does not constitute professional or legal advice. For specific guidance related to your business needs, please consult with appropriate qualified professionals.

© 2026 IntuitionLabs.ai. All rights reserved.