

How to Build an AI Datacenter: GPU Cluster Engineering Guide

Published July 21, 2026 33 min read



Executive Summary

Building an artificial intelligence (AI) datacenter in 2026 means engineering a facility around three interlocking constraints that did not exist a decade ago: [rack-level power densities that can exceed 130 kilowatts](#) (kW), a networking fabric that must move data between accelerators faster than most 2015-era supercomputers could move data at all, and an electrical grid connection process that now takes longer than the construction of the building itself. The reference point for the current generation is **NVIDIA's GB200 NVL72**, a liquid-cooled rack that connects 36 **Grace** central processing units (CPUs) and 72 **Blackwell** graphics processing units (GPUs) into a single 130 terabyte-per-second (TB/s) NVLink domain rated at roughly 120 to 132 kW (Source: [nvidia.com](#)) (Source: [journal.uptimeinstitute.com](#)). By comparison, the Uptime Institute's own global survey found that only about 1% of operators run any rack above 100 kW, which is why direct liquid cooling (DLC) has moved from a specialty technique to what Uptime Institute analyst Daniel Bizo calls a near-mandatory baseline for frontier AI hardware (Source: [journal.uptimeinstitute.com](#)).

Cost has scaled with density. JLL Research reports that the global average shell-and-core construction cost for a data center rose from \$7.7 million per megawatt (MW) in 2020 to \$10.7 million per MW in 2025, and forecasts a further 6% rise to \$11.3 million per MW in 2026, before tenant technology fit-out, which can add up to \$25 million per MW for AI infrastructure, is even included (Source: [jll.com](#)) (Source: [jll.com](#)). Goldman Sachs Research puts the all-in figure for next-generation AI facilities at \$15 million to \$20 million per MW, roughly double the \$10 million per MW typical of a conventional hyperscale build (Source: [goldmansachs.com](#)). That premium buys electrical switchgear, medium-voltage distribution, and cooling plant capable of handling loads that the International Energy Agency (IEA) says can reach 20 times the consumption of a typical AI-focused data center, itself already comparable to 100,000 households (Source: [iea.org](#)).

On networking, the industry has split into two camps rather than converging on one standard. NVIDIA's **Quantum-2 InfiniBand** platform delivers 400 gigabits per second (Gb/s) per port and up to 51.2 TB/s of aggregate switch throughput (Source: [nvidia.com](#)), remaining the reference fabric behind NVIDIA's flagship supercomputing cluster designs. At the same time, Ethernet-based fabrics such as NVIDIA's own **Spectrum-X** and Broadcom's **Tomahawk 6**, which ships 102.4 terabits per second (Tbps) of switching capacity in a single chip built for clusters exceeding one million accelerators

(XPUs), have closed most of the historical performance gap while offering open, multi-vendor economics (Source: investors.broadcom.com). NVIDIA itself reports that its Spectrum-X fabric sustained 95% data throughput on xAI's 100,000-GPU Colossus cluster versus roughly 60% for standard Ethernet (Source: nvidianews.nvidia.com).

The hardest constraint is now the [electrical grid](#) rather than the chips themselves. The IEA finds that data centers consumed about 415 terawatt-hours (TWh) of electricity in 2024, roughly 1.5% of world consumption, but that this figure will more than double to around 945 TWh by 2030 (Source: iea.org) (Source: iea.org). The agency estimates that roughly 20% of planned data center projects worldwide could face delays because of grid connection constraints, and separately notes that transformer and cable lead times have doubled over the past three years (Source: iea.org). In Texas alone, the ERCOT grid operator's large-load interconnection queue reached over 410 gigawatts (GW) by April 2026, dominated by data center requests (Source: rtoinsider.com). This is why the fastest builders, including xAI's Colossus facility in Memphis and Meta's Hyperion and Prometheus campuses, are increasingly generating power on-site with gas turbines rather than waiting in the interconnection queue, a strategy detailed further in the case studies below.

This report walks through the full stack of decisions involved in building an AI datacenter: rack and power architecture, cooling method selection, network topology and fabric choice, facility-level power and redundancy design, and the real-world economics and timelines observed at hyperscale sites operated by xAI, Meta, Microsoft, Google, and CoreWeave. It closes with data on rack density trends, [water consumption](#), and [equipment supply constraints](#) that any organization planning a build in 2026 and beyond needs to factor into its budget and schedule.

Introduction and Background

An AI datacenter is a purpose-built facility engineered to house large numbers of GPUs or equivalent accelerators (such as Google's Tensor Processing Units, or TPUs) as a single, tightly coupled system rather than as isolated servers. The distinction from a conventional enterprise datacenter is not cosmetic. A general-purpose CPU rack draws up to about 12 kW, while a **NVIDIA H100** air-cooled rack draws around 40 kW, and the current flagship **GB200 NVL72** rack draws approximately 120 kW, ten times the density of ordinary enterprise infrastructure (Source: newsletter.semianalysis.com). That order-of-magnitude jump in density cascades into every other design decision: power distribution, [cooling architecture](#), floor loading, and the physical topology of the network that connects the racks. The United States alone hosts nearly half of the world's data center capacity within just five regional clusters, according to the IEA, which concentrates both the opportunity and the grid strain of new AI-focused construction in a handful of markets (Source: iea.org).

The scale of the buildout underway is unprecedented in the history of computing infrastructure. Global investment in data centers has nearly doubled since 2022, reaching half a trillion dollars in 2024 according to the IEA (Source: iea.org). Individual companies are now committing capital at a scale once reserved for national infrastructure projects. **The Stargate Project**, backed by OpenAI, SoftBank, Oracle, and MGX, intends to invest \$500 billion over four years, with \$100 billion deployed immediately, and has since grown to nearly 7 gigawatts of planned capacity and over \$400 billion in committed investment across sites including Abilene, Texas (Source: openai.com) (Source: group.softbank). Meta raised its 2026 capital expenditure guidance to a range of \$125 billion to \$145 billion, up from \$115 billion to \$135 billion, and up from \$72.2 billion actually spent in 2025 (Source: fortune.com).

Building an AI datacenter today requires coordinating five interdependent workstreams that traditional datacenter projects handled largely in sequence: compute hardware selection (GPU or TPU generation and rack form factor), cooling system design (air, direct liquid, or hybrid), network fabric selection (InfiniBand, Ethernet, or a hybrid of the two), electrical and grid interconnection (utility power, on-site generation, or both), and facility resiliency tier (Uptime Institute Tier I through Tier IV). Because transformer lead times can now stretch to four years and GPU allocations are negotiated years in advance, these five workstreams must be planned concurrently rather than sequentially (Source: pv-magazine-usa.com). The remainder of this report treats each workstream in turn, then examines the quantitative data on cost, density, and power, followed by named case studies and a forward-looking assessment.

Rack and Compute Architecture: The Building Block of an AI Datacenter

Understanding the Modern AI Rack

The fundamental unit of AI datacenter design has shifted from the individual server to the **rack as a single computer**. NVIDIA describes its GB200 NVL72 as "an exascale computer in a single rack," combining 36 Grace CPUs and 72 Blackwell GPUs into one NVLink domain that behaves as a single massive GPU rather than 72 discrete accelerators (Source: nvidia.com). Physically, the rack consists of 18 compute trays (each holding two Grace-Blackwell "Bianca" boards) and 9 NVSwitch trays, weighing about 1.36 metric tons and drawing roughly 120 to 132 kW depending on the exact

configuration (Source: [newsletter.semianalysis.com](https://www.semianalysis.com/newsletter)) (Source: [journal.uptimeinstitute.com](https://www.uptimeinstitute.com/journal)). As Uptime Institute puts it, just a few small rows of these Blackwell rack systems can consume as much power capacity as a sizable data hall filled with hundreds of ordinary racks, which is precisely why the concentration of capacity puts extra demands on power distribution equipment and floor loading tolerances (Source: [journal.uptimeinstitute.com](https://www.uptimeinstitute.com/journal)).

Rack-scale performance figures illustrate why density matters. The GB200 NVL72 delivers 1,440 petaflops (PFLOPS) of FP4 (four-bit floating point) tensor performance and 130 TB/s of NVLink bandwidth per rack, alongside 13.4 terabytes (TB) of high-bandwidth memory (HBM3E) shared across the domain (Source: [nvidia.com](https://www.nvidia.com)). Its successor, the **GB300 NVL72**, integrates 72 Blackwell Ultra GPUs and 36 Grace CPUs with 37 TB of shared fast memory across the domain, delivering 1.5 times more dense FP4 FLOPS and 2 times higher attention-layer performance, with an overall AI factory output improvement of up to 50 times compared with Hopper-generation platforms (Source: [nvidia.com](https://www.nvidia.com)) (Source: [nvidia.com](https://www.nvidia.com)).

Form Factor Tradeoffs

Not every deployment uses the full NVL72 form factor. SemiAnalysis documents four principal rack-scale variants of Blackwell hardware:

- **NVL72** (single rack, 120 kW): the highest-density configuration, used by only a small number of hyperscalers as a primary deployment target because most datacenter infrastructure cannot yet support the density even with direct-to-chip liquid cooling (Source: [newsletter.semianalysis.com](https://www.semianalysis.com/newsletter)).
- **NVL36x2** (two racks, 66 kW each, 132 kW combined): the more commonly deployed form factor, offering the same non-blocking 72-GPU domain across two physically separate racks, at the cost of roughly 10 kW more total power due to additional NVSwitch application-specific integrated circuits (ASICs) (Source: [newsletter.semianalysis.com](https://www.semianalysis.com/newsletter)).
- **Ariel-board variant**: a Meta-specific configuration with a 1:1 Grace-to-Blackwell ratio (versus the standard 1:2) to support recommendation-system workloads that need more CPU memory for embedding tables (Source: [newsletter.semianalysis.com](https://www.semianalysis.com/newsletter)).
- **Miranda (B200 with x86 CPUs)**: a lower-cost variant that swaps NVIDIA's Grace CPU for standard x86 processors, trading lower upfront capital cost for reduced CPU-to-GPU bandwidth (versus up to 900 gigabytes per second, or GB/s, bidirectional on Grace) (Source: [newsletter.semianalysis.com](https://www.semianalysis.com/newsletter)).

Cost of the Compute Layer

Compute hardware, not the building shell, dominates the capital budget of a modern AI datacenter. Tenant technology fit-out can reach \$25 million per MW against a shell-and-core cost of roughly \$11.3 million per MW (Source: [jll.com](https://www.jll.com)), and fleet-level GPU purchases now run into the tens of billions of dollars per site, as the Colossus case study below illustrates in detail. Goldman Sachs Research's own sensitivity analysis shows how quickly this compounds: modeling the effect of moving data center capital expenditure from \$11 million per MW to \$19 million per MW shows the cost-per-MW assumption alone can swing total build-out cost by billions of dollars across a multi-gigawatt campus (Source: [goldmansachs.com](https://www.goldmansachs.com)).

GPU Cluster Architecture and Supercomputer Design

From Server to Supercomputer

Designing a GPU supercomputer means deciding how individual racks are wired together into pods, and how pods are wired into a full cluster. NVIDIA's DGX SuperPOD reference architecture is the most widely cited blueprint: a three-tier fat-tree network topology connecting up to 32 DGX systems using Quantum-2 InfiniBand switches at 400 Gb/s per port, delivering full bisection bandwidth so that the aggregate bandwidth between any two halves of the cluster equals the total bandwidth into either half (Source: [introl.com](https://www.introl.com)). Full bisection bandwidth matters because distributed training workloads generate communication patterns that shift dynamically as different layers of a model synchronize gradients, and a fat-tree design provides predictable performance regardless of which GPU pairs are communicating at any given moment.

Rail-Optimized Networking

A refinement on the basic fat-tree is **rail-optimized topology**, which recognizes that a GPU server contains multiple accelerators sharing high-speed internal interconnects, and aligns network connections with each GPU's physical placement rather than treating every GPU as an independent endpoint. NVIDIA's own NVL72 AI Factory reference architecture uses exactly this approach: the compute fabric is "a full non-blocking fat tree topology," with GPUs connected "using a rail-optimized network topology through their respective ConnectX-8 NICs," designed to provide the shortest

possible hop count for both intra-tray and inter-tray communication (Source: docs.nvidia.com). In NVIDIA's dual-plane implementation, each 800 Gb/s GPU interface is split into two 400 Gb/s interfaces, each connected to a physically distinct leaf switch on an independent fabric plane, so that a failure or degradation on one plane only reduces available bandwidth linearly rather than taking the GPU offline entirely (Source: docs.nvidia.com).

Reliability at Scale

Network reliability becomes a first-order design concern once clusters cross tens of thousands of GPUs, because a single misconfiguration can stall an entire training run. Meta's own infrastructure research found that network configuration errors caused 10.7% of significant GPU job failures across its fleet, underscoring why fat-tree and rail-optimized designs emphasize redundant paths and automated validation over raw peak bandwidth (Source: introl.com). Alternative topologies exist for specific workloads: Google's TPU pods historically use a three-dimensional torus, and Amazon Web Services (AWS) uses workload-optimized topologies for its Trainium accelerators, but the fat-tree and rail-optimized family remains dominant for GPU-based training clusters.

Designing a GPU Supercomputer Step by Step

Practically, engineering teams work through the following sequence when designing a GPU supercomputer:

- **Choose the accelerator generation and rack form factor** (for example, GB200 NVL72 versus NVL36x2) based on available floor loading, cooling infrastructure, and power delivery at the target site.
- **Select the scale-up fabric** that connects GPUs within a rack (NVLink, at up to 1.8 TB/s of GPU-to-GPU bandwidth per rack in Microsoft's Fairwater design) (Source: blogs.microsoft.com).
- **Select the scale-out fabric** that connects racks and pods into the full cluster, choosing between InfiniBand, Ethernet (Spectrum-X or a Broadcom Tomahawk-based design), or a hybrid.
- **Size the oversubscription ratio** at each tier of the fat tree, trading network cost against worst-case bandwidth contention.
- **Plan physical cable runs**, since latency-sensitive workloads make cable length a first-order constraint; Microsoft's two-story Fairwater datacenter design exists specifically to shorten cable runs by placing racks in three dimensions (Source: blogs.microsoft.com).
- **Validate the fabric** under synthetic collective-communication workloads before committing production training jobs, since a single misrouted plane can silently degrade an entire job's throughput.

InfiniBand vs Ethernet: Choosing the GPU Cluster Networking Fabric

The Core Tradeoff

The single most consequential and most debated decision in GPU cluster design is whether to build the scale-out fabric on InfiniBand or Ethernet. InfiniBand originated in supercomputing and uses credit-based flow control to guarantee lossless transmission, with a centralized Subnet Manager computing deterministic routes; NVIDIA's Quantum-2 switches deliver 400 Gb/s per port, 64 ports (or 128 ports at 200 Gb/s) per switch, an aggregate 51.2 Tb/s of bidirectional throughput, and more than 66.5 billion packets per second of switching capacity (Source: nvidia.com) (Source: nvidia.com). At its largest scale-out configuration, NVIDIA reports Quantum-2 can support over one million 400 Gb/s nodes in a four-switch-tier DragonFly+ network, 6.5 times higher than the previous generation (Source: nvidia.com). Ethernet, by contrast, evolved for local-area networking with best-effort, lossy delivery, relying on congestion control algorithms and, for AI workloads, RDMA over Converged Ethernet (RoCE) to approximate InfiniBand's lossless behavior.

Closing the Gap: Spectrum-X and Tomahawk 6

The performance gap has narrowed substantially since 2024. NVIDIA's own production data from xAI's Colossus cluster shows Spectrum-X Ethernet sustaining 95% data throughput with zero application latency degradation or packet loss due to flow collisions across a 100,000-GPU Hopper cluster, versus only 60% data throughput on standard Ethernet at that scale (Source: nvidianews.nvidia.com) (Source: nvidianews.nvidia.com). Spectrum-X achieves this through fine-grain, per-packet adaptive routing and real-time congestion detection built into Spectrum-4 switches, which support up to 128 ports of 400 Gigabit Ethernet (GbE) at sub-microsecond latency, paired with BlueField-3 SuperNICs (Source: weka.io) (Source: weka.io). Broadcom's Tomahawk 6, shipping since June 2025, pushes the open-Ethernet case further: a single chip delivers 102.4 Tbps of switching capacity,

double the prior generation, supports a scale-up cluster size of up to 512 XPUs, and supports scale-out clusters of 100,000 to one million XPUs using only two switch tiers instead of three, which Broadcom states reduces latency, simplifies load balancing, and requires fewer optics (Source: investors.broadcom.com) (Source: investors.broadcom.com) (Source: investors.broadcom.com).

The Ultra Ethernet Consortium

A parallel open-standards effort, the **Ultra Ethernet Consortium (UEC)**, was founded in July 2023 as a Joint Development Foundation project hosted by the Linux Foundation, with founding members including AMD, Arista, Broadcom, Cisco, Eviden, Hewlett Packard Enterprise (HPE), Intel, Meta, and Microsoft, with the explicit goal of building a complete Ethernet-based communication stack architecture for AI and high-performance computing (HPC) workloads (Source: linuxfoundation.org). As UEC chair Dr. J Metz put it at the group's founding, the effort "isn't about overhauling Ethernet," but about "tuning Ethernet to improve efficiency for workloads with specific performance requirements" at every layer from the physical up through the software stack (Source: linuxfoundation.org).

Microsoft's Approach: Ethernet at Hyperscale

Microsoft's Fairwater AI superfactory design illustrates a large hyperscaler betting on open Ethernet rather than InfiniBand for scale-out networking. Fairwater uses "a two-tier, ethernet-based backend network that supports massive cluster sizes with 800 Gbps GPU-to-GPU connectivity," running on Microsoft's own SONiC (Software for Open Networking in the Cloud) operating system, explicitly to avoid vendor lock-in and to use commodity switching hardware instead of proprietary solutions (Source: blogs.microsoft.com) (Source: blogs.microsoft.com).

Table 1 below summarizes the primary quantitative differences between the two fabrics based on vendor-published specifications gathered in this research.

ATTRIBUTE	INFINIBAND (NVIDIA QUANTUM-2)	ETHERNET (SPECTRUM-X / TOMAHAWK 6)
Per-port speed	400 Gb/s (Quantum-2) (Source: nvidia.com)	Up to 800 Gb/s (Spectrum-X on Fairwater) (Source: blogs.microsoft.com)
Single-chip switching capacity	51.2 Tb/s aggregate (Quantum-2 switch) (Source: nvidia.com)	102.4 Tbps (Broadcom Tomahawk 6) (Source: investors.broadcom.com)
Measured throughput at 100k-GPU scale	Not separately disclosed by NVIDIA for Colossus	95% sustained (Spectrum-X on Colossus) vs. 60% on standard Ethernet (Source: nvidianews.nvidia.com)
Governance model	Proprietary NVIDIA ecosystem	Open standard; UEC founded 2023 with founding members including AMD, Broadcom, Cisco, Meta, and Microsoft (Source: linuxfoundation.org)
Reference topology	Three-tier fat tree, up to 32 DGX systems, full bisection bandwidth (NVIDIA DGX SuperPOD reference architecture, cited above)	Two-tier scale-out possible at up to 128,000 XPUs with Tomahawk 6 (Source: investors.broadcom.com)

As the table shows, the two fabrics have converged on similar aggregate capacity, but they differ fundamentally in governance and multi-vendor flexibility rather than raw throughput. Organizations planning a build should treat this as a build-versus-buy decision about ecosystem lock-in as much as a pure networking benchmark, since NVIDIA's own GB300 platform is explicitly designed to run on either Quantum-X800 InfiniBand or Spectrum-X Ethernet interchangeably (Source: nvidia.com). As a general rule of thumb, smaller GPU deployments can typically be managed on Ethernet without heavy tuning, while very large deployments increasingly benefit from either InfiniBand's automated management or from a Spectrum-X- or Tomahawk-class Ethernet deployment engineered specifically for AI traffic patterns.

Power, Cooling, and Facility Design

Why Air Cooling No Longer Works

Air cooling becomes physically inadequate above approximately 30 to 40 kW per rack, because the air volume and velocity required to remove that much heat from a standard rack form factor exceeds what data center airflow systems can deliver without creating hot spots and thermal runaway (Source: computeforecast.com). At the GB200 NVL72's roughly 120 kW rating, direct liquid cooling (DLC) is not optional: the cold plates handle around 100 kW of thermal load while air exhaust manages the remaining 25 kW or so (Source: journal.uptimeinstitute.com). Direct-to-chip cooling eliminates most of the thermal resistance between GPU dies and the cooling system by removing several of the interfaces heat must otherwise cross in an air-cooled path, which is why one infrastructure engineering firm reports it drops measured facility PUE from 1.58 down to 1.15 in production deployments (Source: introl.com).

Density De-Escalation Options

Not every organization can or should deploy the flagship 120 kW rack. Uptime Institute's Daniel Bizo lays out a "density de-escalation" ladder for operators constrained by existing facility power or floor loading:

- **132 kW per rack (NVL72):** requires DLC combined with high-density air cooling, 1.4 metric tons of floor loading per rack, and mandates the highest tier of power distribution (Source: journal.uptimeinstitute.com).
- **80 to 90 kW per rack (64-GPU, fiber-only):** roughly a third lower in power and weight than the flagship, but requires about 50% more racks to match NVL72's training performance (Source: journal.uptimeinstitute.com) (Source: journal.uptimeinstitute.com).
- **Below 50 kW per rack (four 8-GPU frames per rack):** makes air cooling technically viable again with containment and rear-door heat exchangers, at the cost of fan power reaching up to 10% of total load (Source: journal.uptimeinstitute.com).
- **26 kW and below:** manageable on standard 400V/32A or 415V/40A circuits under North American derating codes, with dynamic power capping software throttling chips to stay within circuit limits (Source: journal.uptimeinstitute.com).

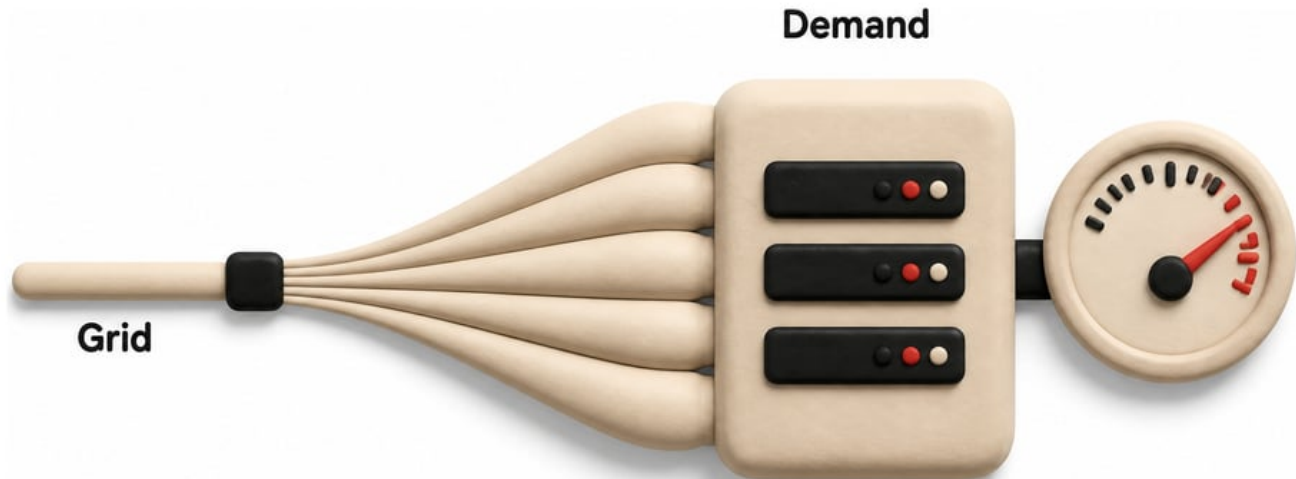
Facility-Level Resiliency Standards

Beyond individual racks, facility owners must choose a resiliency tier under the Uptime Institute's classification system, the international standard used by data center operators for over 30 years. **Tier I** provides basic capacity with an uninterruptible power supply (UPS) and engine generator but requires full shutdown for maintenance. **Tier II** adds redundant capacity components. **Tier III** is concurrently maintainable, meaning any component can be taken offline for maintenance without affecting information technology (IT) operations, achieved through redundant distribution paths. **Tier IV** adds full fault tolerance through physically isolated, independent systems so that a single equipment failure does not affect operations at all (Source: uptimeinstitute.com) (Source: uptimeinstitute.com). Uptime Institute has issued over 4,300 Tier Certifications across more than 120 countries, and remains the sole licensed body able to issue this specific certification (Source: uptimeinstitute.com).

Energy Efficiency and Regulation

Energy efficiency, measured by PUE, has plateaued industry-wide at 1.55 to 1.59 since 2020, though facilities under 15 years old average around 1.48 and those built in the past five years average about 1.45 globally, with the newest North American and European sites achieving better than 1.4 (Source: journal.uptimeinstitute.com) (Source: journal.uptimeinstitute.com). Germany's Energieeffizienzgesetz (EnEfG) is the strictest data center performance law in the European Union: for new facilities commissioned from July 2026 onward, PUE must reach 1.2 within two years, while existing facilities face a more graduated schedule of a PUE of 1.5 or better by July 2027 and 1.3 or better by July 2030 (Source: moduledge.com). New German data centers must also meet escalating waste-heat reuse targets and renewable energy requirements that ramp from 50% unsubsidized in 2024 to 100% from January 2027 (Source: moduledge.com). Individual hyperscalers already operate well inside the German threshold: reporting compiled by one EU compliance analysis puts Google's fleet-average PUE at about 1.09 and Amazon Web Services at about 1.15, showing that sub-1.2 PUE is achievable at scale even though the wider industry average remains far higher (Source: moduledge.com). This German requirement, not a directly binding European-Union-wide mandate, is what most operators mean when they refer to an "EU PUE 1.2 rule," since the broader EU Energy Efficiency Directive itself only mandates reporting rather than a universal PUE ceiling (Source: moduledge.com).

Grid Interconnection and the Power Bottleneck



Electricity Demand Is Now the Binding Constraint

Data centers consumed approximately 415 TWh of electricity globally in 2024, about 1.5% of world electricity consumption, having grown roughly 12% per year since 2017, more than four times the growth rate of total electricity consumption (Source: [iea.org](https://www.iea.org) (Source: [iea.org](https://www.iea.org))). The United States accounted for 45% of that consumption, followed by China at 25% and Europe at 15% (Source: [iea.org](https://www.iea.org)). Goldman Sachs Research forecasts data center power demand will grow 50% by 2027 and by as much as 165% by 2030 relative to 2023 (Source: [goldmansachs.com](https://www.goldmansachs.com)), and estimates roughly \$720 billion of additional grid spending will be needed through 2030 to accommodate the growth, with transmission projects taking several years to permit and several more to build (Source: [goldmansachs.com](https://www.goldmansachs.com)). Outside the United States, Goldman Sachs Research reports the same pressure building: a survey of European utilities found that the number of connection requests received by power distribution operators, a leading indicator of future demand, has risen exponentially over the past few years, driven mostly by data centers (Source: [goldmansachs.com](https://www.goldmansachs.com)).

Grid Queues and Delays

More than 2,500 GW of projects, spanning renewables, storage, and large new loads such as data centers, currently sit in grid connection queues worldwide, and the IEA estimates nearly 20% of planned data center projects globally could face significant delays as a result (Source: [iea.org](https://www.iea.org)). Building new transmission lines can take four to eight years in advanced economies, and wait times for critical grid components such as transformers and cables have doubled over the past three years (Source: [iea.org](https://www.iea.org)). Roughly half of data centers currently under development in the United States are located in these same pre-existing large clusters, which the IEA warns raises the risk of concentrated local bottlenecks rather than a more geographically distributed strain on the grid (Source: [iea.org](https://www.iea.org)). Texas illustrates the scale of the pressure regionally: ERCOT's large-load interconnection queue reached over 410 GW by April 2026, driven by a wave of requests submitted through Oncor Electric Delivery, and data center projects specifically account for the large majority of the queue (Source: [rtoinsider.com](https://www.rtoinsider.com)).

Equipment Shortages Compound the Delay

Even where grid capacity exists, the physical equipment needed to connect a datacenter is scarce. Large power transformers that took two years or less to deliver before 2020 now face lead times as long as four years according to analysts cited by Reuters Events and PwC, driven by a 274% increase in demand for generator step-up transformers and a 116% increase in demand for substation transformers between 2019 and 2025, according to Wood Mackenzie data (Source: [pv-magazine-usa.com](https://www.pv-magazine-usa.com)) (Source: [pv-magazine-usa.com](https://www.pv-magazine-usa.com)). Prices for the same equipment have risen roughly 80% over five years, and manufacturers including Hitachi Energy (over \$1 billion in new U.S. capacity) and Siemens (\$421 million for a Charlotte, North Carolina plant) are investing to close the gap, though the shortage is expected to persist for years (Source: [pv-magazine-usa.com](https://www.pv-magazine-usa.com)) (Source: [pv-magazine-usa.com](https://www.pv-magazine-usa.com)). This is precisely why on-site power generation, discussed in the case studies below, has become a competitive advantage rather than an exception.

Data Analysis and Evidence

The quantitative picture of AI datacenter construction is dominated by four figures: rack power density, construction cost per MW, electricity demand growth, and equipment lead times. On density, the AFCOM State of the Data Center Report found average rack density reached 27 kW per rack in 2026, up from 16 kW in 2025, the largest year-over-year jump recorded in the report's decade-long history, and up from just 7 kW in 2021 (Source: datacenterknowledge.com) (Source: datacenterknowledge.com). Nearly 70% of surveyed operators expect further density increases within 12 to 36 months (Source: datacenterknowledge.com).

On cost, JLL's shell-and-core figure of \$11.3 million per MW for 2026 sits well below Goldman Sachs Research's all-in estimate of \$15 million to \$20 million per MW for next-generation AI facilities once tenant fit-out, redundancy, and specialized cooling are included, illustrating how sensitive total capital expenditure is to the definitional boundary of "cost per MW" (Source: jll.com) (Source: goldmansachs.com). JLL separately estimates \$3 trillion of investment will be required to deliver 100 GW of new data center supply by 2030, a figure consistent with the scale of capital commitments already announced by OpenAI's Stargate consortium and by Meta (Source: jll.com).

Table 2 below compares announced power capacity and disclosed investment for the largest publicly documented AI datacenter campuses as of mid-2026, drawn from company statements, utility filings, and trade press coverage gathered in this research; full sourcing for the xAI Colossus figures appears in the Case Studies section immediately below the table.

FACILITY / OWNER	DISCLOSED OR PLANNED CAPACITY	DISCLOSED COST OR FINANCING	NOTABLE DETAIL
xAI Colossus (Memphis, Tennessee)	~2 GW across three buildings	~\$18 billion for 555,000 GPUs	Original 100,000-GPU phase built in 122 days (Source: nvidianews.nvidia.com)
Stargate (Abilene, TX and 5 additional sites)	~7 GW planned (Source: openai.com)	\$500 billion over 4 years; \$100 billion immediate (Source: openai.com)	Backed by SoftBank, OpenAI, Oracle, and MGX
Meta Hyperion (Richland Parish, Louisiana)	2 GW by decade end, scaling toward 5 GW (Source: datacenterfrontier.com)	~\$10 billion buildout; \$29 billion financing (\$3B equity / \$26B debt) (Source: datacenterfrontier.com)	2,250 acres; 4 million sq ft of buildings (Source: datacenterfrontier.com)
Microsoft Fairwater (Wisconsin + Atlanta)	Hundreds of thousands of GB200/GB300 GPUs across linked sites (Source: blogs.microsoft.com)	Not separately disclosed	140 kW per rack; 1,360 kW per row (Source: blogs.microsoft.com)
CoreWeave (multi-site, North America and Europe)	3.1 GW+ contracted; targeting 5 GW+ by 2030	Not separately disclosed for infrastructure alone	250,000+ GPUs across 43 datacenters (Source: nextplatform.com)

The wide variance in disclosure format between these five projects (per-GPU spend for xAI, total multi-year investment for Stargate, buildout-plus-financing for Meta, per-rack density for Microsoft, and fleet size for CoreWeave) itself illustrates why standardized cost-per-MW benchmarks from JLL and Goldman Sachs remain the most useful cross-comparison tool available to planners, even though they average over very different site-level realities.

On water, U.S. data centers directly consumed 17.4 billion gallons of water in 2023, with a 2024 academic estimate suggesting that training a single large language model on the scale of GPT-4 consumed approximately 700,000 liters (about 185,000 gallons) for cooling, while typical inference workloads consume far less per query (roughly 500 milliliters, about one bottle, per 10 to 50 responses). Larger data centers can require up to 5 million gallons of water per day, comparable to the demand of a city of 50,000 people. Microsoft's Fairwater design addresses this directly with a closed-loop cooling system that reuses liquid continuously after an initial fill (equivalent to what 20 homes consume in a year) with no evaporative loss and a design life of six-plus years before replacement is needed (Source: blogs.microsoft.com).

Case Studies and Real-World Examples

xAI Colossus (Memphis, Tennessee)

xAI's Colossus facility is the clearest public demonstration of how fast an AI datacenter can be built when power is generated on-site rather than sourced from the utility grid. The original 100,000-GPU phase, using NVIDIA Hopper GPUs and Spectrum-X Ethernet networking, was built in 122 days total, with only 19 days elapsing between the first rack rolling onto the floor and the start of training, a timeline NVIDIA CEO Jensen Huang called "superhuman" (Source: nvidianews.nvidia.com). By January 2026, Musk announced a third building expanding total site capacity to nearly 2 GW and 555,000-plus GPUs, purchased for approximately \$18 billion, implying an average cost of roughly \$32,400 per GPU across the fleet and making the site the world's largest single-site AI training installation (Source: introl.com) (Source: introl.com). The site bypasses conventional utility interconnection entirely by building its own gas-fired power generation adjacent to the datacenter, a 2 GW load equivalent to powering approximately 1.5 million homes.

Meta's Hyperion and Prometheus Superclusters

Meta's Hyperion campus in Richland Parish, Louisiana, illustrates the scale of grid infrastructure now required to feed a single AI datacenter. Utility Entergy has proposed a 100-mile, 500-kilovolt (kV) transmission project costing approximately \$1.2 billion, alongside three new combined-cycle gas plants generating about 2.25 GW at a cost of just under \$4 billion, to serve the site while broader transmission is built out (Source: datacenterfrontier.com) (Source: datacenterfrontier.com). Local community groups, the Union of Concerned Scientists, and the Louisiana-based Alliance for Affordable Energy have raised concerns about ratepayer cost exposure and what happens once initial 15-year power contracts expire (Source: datacenterfrontier.com). Meta's companion Prometheus site in New Albany, Ohio, occupies a 740-acre industrial tract and is described by Mark Zuckerberg as Meta's first "titan" AI supercluster, due online in 2026 (Source: datacenterfrontier.com). To finance Hyperion, Meta secured \$29 billion (\$3 billion in equity from Blue Owl Capital and \$26 billion in debt from PIMCO), reflecting a broader industry shift toward direct ownership rather than pure cloud leasing (Source: datacenterfrontier.com).

Microsoft's Fairwater AI Superfactory (Wisconsin and Atlanta)

Microsoft's Fairwater sites demonstrate a networked, multi-site design philosophy rather than a single-campus approach. The Atlanta Fairwater 2 facility, a two-story datacenter that began operation in October 2025, is directly linked to the original Wisconsin Fairwater site via a dedicated AI Wide Area Network (WAN) built on more than 120,000 miles of fiber, allowing the two sites to function as a single "AI superfactory" capable of training a model across geographically separated locations (Source: blogs.microsoft.com). The design achieves 140 kW per rack and 1,360 kW per row through closed-loop liquid cooling, and Microsoft states that every GPU is connected to every other GPU across the network, requiring the two-story design specifically to shorten cable runs (Source: blogs.microsoft.com).

Google's Ironwood TPU Pods

Google's seventh-generation TPU, Ironwood, demonstrates an alternative accelerator architecture built around a custom interconnect rather than NVIDIA's NVLink and InfiniBand/Ethernet stack. Ironwood scales up to 9,216 liquid-cooled chips per pod, delivering 42.5 exaflops of aggregate compute when scaled to that size, more than 24 times the compute power of El Capitan, the world's largest supercomputer at the time of Google's announcement (Source: blog.google). Each individual chip offers 192 gigabytes (GB) of high-bandwidth memory (HBM) with 7.37 TB/s of bandwidth, connected through an Inter-Chip Interconnect (ICI) network that spans nearly 10 MW at full pod scale and, in Ironwood, is enhanced to 1.2 terabytes per second (TBps) of bidirectional bandwidth, 1.5 times that of the prior generation (Source: blog.google) (Source: blog.google) (Source: blog.google). Google also reports that Ironwood delivers roughly twice the performance per watt of Trillium, its prior-generation TPU, at a time when available power is one of the central constraints on delivering AI capacity (Source: blog.google). Ironwood scales beyond a single pod into clusters of hundreds of thousands of chips using Google's Pathways software stack and Jupiter datacenter network, illustrating that fat-tree- and Ethernet-centric scale-out designs are not unique to NVIDIA-based clusters (Source: blog.google).

CoreWeave: A Neo-Cloud Built Purely for AI

CoreWeave illustrates how a company that started outside the traditional hyperscaler tier can compete on infrastructure specialization. As of its public disclosures, CoreWeave operates more than 250,000 GPUs across 43 data centers with over 3.1 GW of contracted power capacity, and is targeting more than 5 GW of AI datacenter capacity by 2030 (Source: nextplatform.com). NVIDIA has invested \$2 billion directly in CoreWeave, becoming its second-largest shareholder, reflecting how deeply chip suppliers are now embedding themselves in the economics of the datacenters that consume their hardware.

Implications and Future Directions

The trajectory of AI datacenter design over the next several years is set by three forces already visible in the data gathered above. First, rack density will keep climbing faster than facility power infrastructure can be upgraded to match: Uptime Institute notes that rack-scale systems of 300 kW and above are slated to begin shipping as early as 2026, which will push direct liquid cooling from a high-end option to an absolute baseline requirement across the industry. Second, the networking fabric debate between InfiniBand and Ethernet is very unlikely to resolve into a single winner. Instead, expect continued bifurcation, with InfiniBand or InfiniBand-class deterministic fabrics retained for the largest, most latency-sensitive training runs, and open Ethernet fabrics such as Spectrum-X and Tomahawk 6 handling inference and increasingly large fractions of training as the Ultra Ethernet Consortium's specification matures through 2026.

Third, and most consequentially, electrical grid capacity rather than chip supply or even capital availability will most likely be the binding constraint on how quickly new AI datacenter capacity can come online through the rest of this decade. With more than 2,500 GW of global interconnection requests queued and transformer lead times stretching to four years, expect on-site power generation, the model pioneered at xAI's Colossus and now being replicated by Meta at Hyperion, to become the default strategy for the largest training clusters rather than a workaround reserved for the most impatient builders. This shift carries its own second-order implications: reliance on gas-fired on-site generation raises emissions and permitting questions that grid-connected renewables-backed projects do not face, and the IEA's own projection that natural gas will expand by 175 TWh specifically to meet data center demand suggests this tension will intensify rather than resolve on its own (Source: iea.org). Organizations planning a build in this environment should treat power procurement, not chip procurement, as the long-lead-time item that determines a project's realistic in-service date.

Frequently Asked Questions (FAQs)

How much does it cost to build an AI datacenter? Shell-and-core construction alone runs approximately \$11.3 million per MW in 2026 according to JLL, while Goldman Sachs Research estimates a fully equipped next-generation AI facility, including tenant technology fit-out, costs \$15 million to \$20 million per MW, roughly double the \$10 million per MW typical of a conventional hyperscale build (Source: jll.com) (Source: goldmansachs.com).

Should organizations choose InfiniBand or Ethernet for a GPU cluster? For the largest, most latency-sensitive training deployments, InfiniBand's deterministic, centrally managed fabric remains the safer default, while Ethernet fabrics built specifically for AI, such as Spectrum-X or Tomahawk 6-based designs, now close most of the performance gap while offering open, multi-vendor economics and are increasingly favored for inference and mixed workloads; NVIDIA's own GB300 platform is designed to run on either fabric interchangeably (Source: nvidia.com).

How much power does an AI rack need? Current flagship racks (NVIDIA GB200 NVL72) draw approximately 120 to 132 kW, roughly ten times the density of a conventional enterprise server rack at about 12 kW, and industry-wide average rack density reached 27 kW in 2026, up from 16 kW the year before (Source: newsletter.semianalysis.com) (Source: datacenterknowledge.com).

Why is grid interconnection the biggest bottleneck? Transmission lines can take four to eight years to build, transformer lead times now reach four years, and roughly 20% of planned data center projects worldwide could face delays because of grid connection constraints, which is why some of the fastest-moving operators now build gas-fired generation on-site to bypass the utility queue entirely, as detailed in the Colossus case study above (Source: iea.org).

How fast can an AI datacenter actually be built? xAI's Colossus facility went from an empty factory shell to a training-ready 100,000-GPU cluster in 122 days, with only 19 days between the first rack arriving and training commencing, though this timeline relied on on-site power generation and is not representative of grid-connected projects, which more typically take two to four years (Source: nvidianews.nvidia.com).

What Tier classification should an AI datacenter target? Most hyperscale AI training facilities target Tier III (concurrently maintainable, allowing component-level maintenance without downtime) given the extreme cost of an unplanned outage during a multi-week training run, while facilities running latency-critical inference for paying customers increasingly consider Tier IV's full fault tolerance despite its higher construction cost (Source: uptimeinstitute.com).

Conclusion

Building an AI datacenter in 2026 is no longer primarily a real estate or construction problem; it is a systems engineering problem in which power, cooling, and networking must be co-designed with the compute hardware from day one. The reference architectures published by NVIDIA, the operational data disclosed by Microsoft, Meta, Google, and xAI, and the market data compiled by JLL, Goldman Sachs Research, and the IEA all point to the same conclusion: rack density has outpaced facility infrastructure, grid capacity has outpaced transmission and transformer supply, and networking has bifurcated into a deterministic InfiniBand tier and an increasingly capable open Ethernet tier rather than converging on one winner. Organizations planning a build should treat electrical interconnection and equipment lead times, not GPU allocation, as the schedule-critical path, and should size their cooling and power distribution architecture for the density tier they can realistically support today rather than the flagship rack NVIDIA markets, since the density de-escalation options documented by Uptime Institute show credible, lower-risk paths to comparable training performance at meaningfully lower infrastructure risk. The organizations moving fastest, from xAI's 122-day Colossus buildout to Meta's dual-track ownership and leasing strategy, share a common thread: they have stopped treating power, cooling, and networking as downstream implementation details and instead treat them as the primary design constraint around which the entire facility is built.

Tags: ai datacenter, gpu cluster architecture, infiniband vs ethernet, data center cooling, nvidia gb200 nvl72, data center power density, gpu supercomputer, ai datacenter cost per megawatt, rack density, data center networking topology

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 [GPUSmith](#) 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.

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

© 2026 [GPUSmith](#). All rights reserved.