

AHEAD

EXECUTIVE PLAYBOOK:

Liquid Cooled Infrastructure for AI

How to design, deploy, and scale liquid-cooled infrastructure
for high-density AI workloads



Modern AI workloads

are driving unprecedented rack densities and power draws, pushing traditional air-cooled data centers to their limits. For example, standard racks have historically had an average installed power of 10-20kW, but the current power requirements for a full rack of NVIDIA's latest Blackwell NVL72 reaches 120kW. The next generation of AI compute could push the power demand even further to 400kW and beyond.

As AI and high performance computing (HPC) continue to increase in scale and complexity, liquid cooling has become a necessity. The direct-to-chip liquid cooling approach in particular is becoming a mainstream solution for modern data centers because it enables highly localized and efficient heat dissipation for high-density AI infrastructure.

In this playbook, we'll discuss how to navigate the key decisions around AI infrastructure, including power, cooling, facility constraints, architecture designs, and more. Read on to learn our best practices for adapting data centers for AI workloads with liquid cooling.

Clarify AI Strategy and Use Cases



Start by anchoring infrastructure decisions to specific business use cases because different AI initiatives will require very different hardware and architecture designs. For example, large language model (LLM) training workloads typically run on GPU-dense clusters with massive power draws, while real-time inferencing requires efficiency and low latency using high-capacity CPUs and GPUs.

This means it's important to understand which AI workloads are in-scope, and what their requirements are in terms of compute, latency, and throughput. Then you can set a realistic target for compute density over a 3-5 year horizon to avoid over-or under-building infrastructure.



Get a Baseline of Current Data Center Estate

Before designing new AI infrastructure, you'll want to get an understanding of your existing environment and constraints. You should quantify your actual power envelope at the facility level, including available capacity and future expansion potential. Similarly, you should evaluate your existing cooling plant and your maximum achievable rack density on air.

Space, floor loading, and water availability may introduce additional constraints on upgrading power capacity and the feasibility of deploying liquid cooling systems. By understanding the limitations of your existing data center facility, you can identify where upgrades or redesigns may be required before introducing high-density AI infrastructure.

Define AI Factory Design Principles



Now you can translate your strategy and constraints into clear principles centered around modularity, scalability, efficiency, and risk tolerance. We suggest designing vendor-agnostic AI pods that can be easily adapted to different use cases and scaled incrementally as AI technologies mature.

We also recommend defining an AI Factory approach for moving workloads or use cases along a standardized production line that includes infrastructure, data, AI platforms, and an operating model. This ensures AI initiatives can move more quickly and reliably from exploration to production. While building your AI Factory roadmap, it's helpful to decide whether you're optimizing for a few ultra-dense AI pods or moderate density across your broader data center estate.

Select a Liquid Cooling Strategy and Reference Architecture

Now you can choose which liquid cooling approach fits your business strategy and data center constraints:

- **Direct-to-chip cooling:** Uses liquid coolant applied directly to cold plates attached to CPUs and GPUs. This can support extremely dense server configurations and is ideal for advanced AI processing, machine learning, and large-scale analytics workloads.
- **Rear-door heat exchangers:** Replaces standard rack doors with liquid-cooled heat exchangers. This is a transitional solution that can enhance the efficiency of air-cooled environments by removing heat at the rack level.
- **Immersion cooling:** Submerges entire servers into non-conductive fluid for more uniform heat dissipation at the system level. This requires greater infrastructure investment because it's more complex to retrofit into traditional data centers.

In addition, a hybrid approach combining liquid and air cooling could make sense if it's not possible to fully retrofit your entire data center. This enables you to strategically deploy liquid cooling for high-density AI racks while maintaining existing air-cooled infrastructure for less demanding workloads.

Next you can map vendor and OEM reference designs to your mechanical, thermal, and electrical requirements. Leading vendors like CoolIT, Motivair, and Vertiv offer liquid cooling technologies and supporting infrastructure with different capabilities and price points that are worth evaluating.



Engineer the Power Architecture for AI-Class Racks

The fundamental constraint for AI infrastructure today is power consumption, so it's crucial to ensure there's enough capacity available at both rack and facility levels. This means creating engineering power budgets per rack and planning branch circuit designs with appropriate redundancy tiers.

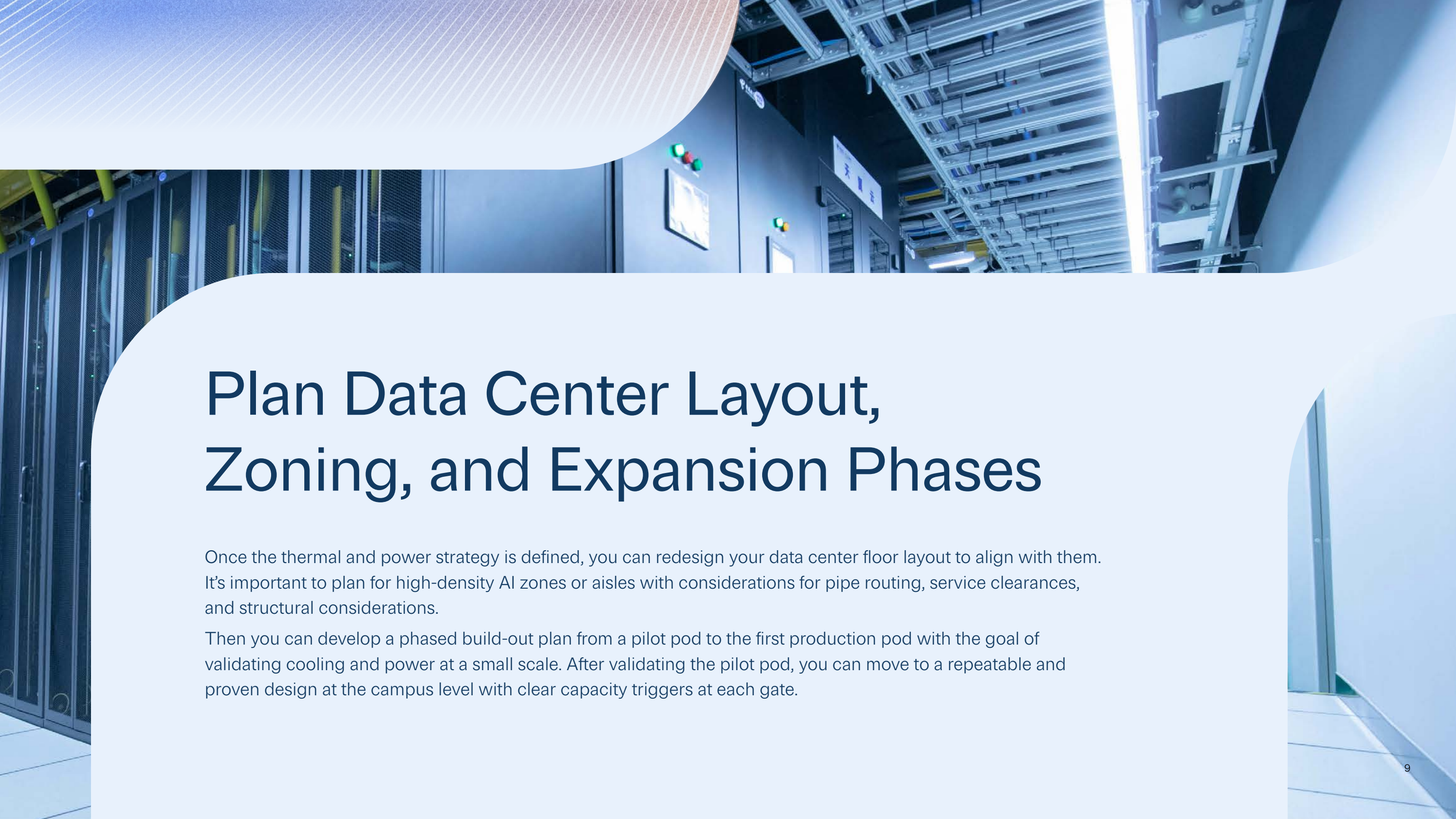
Besides powering the racks themselves, the power budget needs to incorporate the additional load requirements of chillers, pumps, and auxiliary cooling systems for liquid cooling as well. It's also important to consider upstream power distribution, and design redundancy along with a metering strategy to avoid stranded capacity.

Design the Cooling Architecture and Fluid Networks



Now you can translate your cooling strategy into a concrete plant and plumbing design that's reliable and realistic based on your data center constraints. This includes integrating secondary fluid loops for liquid cooling systems with the primary fluid loop of the facility infrastructure like chilled water systems and coolant distribution units (CDUs).

You should also define a redundancy model, controls for leak detection, and failure-mode planning to maintain uptime for the liquid cooling system. For hybrid layouts where liquid-cooled AI pods coexist with legacy air-cooled racks, carefully plan the boundary between thermal zones to ensure existing cooling infrastructure does not degrade.



Plan Data Center Layout, Zoning, and Expansion Phases

Once the thermal and power strategy is defined, you can redesign your data center floor layout to align with them. It's important to plan for high-density AI zones or aisles with considerations for pipe routing, service clearances, and structural considerations.

Then you can develop a phased build-out plan from a pilot pod to the first production pod with the goal of validating cooling and power at a small scale. After validating the pilot pod, you can move to a repeatable and proven design at the campus level with clear capacity triggers at each gate.



Define Rack Integration and Supply Chain Strategy

Next you'll want to consider your deployment strategy and whether to leverage [rack integrations](#). Rather than shipping individual components and assembling them at deployment sites, racks can be pre-integrated at the factory with the necessary cables and other equipment, and pre-validated before deployment. This reduces on-site complexity and accelerates deployment timelines.

In addition, it's important to define a comprehensive supply chain strategy. This should include a standardized bill of materials and configurations to simplify hardware investments and asset lifecycle management. By setting realistic expectations for lead times, staging, and deployment timelines, you can plan capacity better and more effectively scale your AI infrastructure rollout.

Engineer the Networking Fabric and Storage Architecture



Throughput is just as critical as cooling for AI infrastructure performance, so the network architecture needs to be designed to support your specific workloads. You'll want to carefully choose between the two primary fabric options for high-density AI infrastructure — RDMA over Converged Ethernet (RoCE) or InfiniBand — depending on performance requirements. Additional networking considerations include segmentation and traffic patterns between east-west and north-south flows.

You'll also need to align the storage architecture with your rack designs and anticipated patterns for AI training and inferencing. This requires a thorough understanding of throughput, locality, and resilience requirements for your specific workloads.

A man in a checkered shirt is standing and pointing at a large screen displaying a diagram titled "The Clean Architecture". He is in a modern office with large windows in the background. Two people are seated at a table in the foreground, looking at the screen. A laptop and a glass of beer are on the table.

Build the Operational Model and Lifecycle Management

After building the AI infrastructure, the focus shifts to running and maintaining it over time. You should consider creating runbooks for operations teams that include drain/refill procedures, planned maintenance, and emergency response for liquid-cooled racks. You can also integrate monitoring and telemetry for power, thermal, GPU utilization, fabric health, and more with data center infrastructure management (DCIM) and observability platforms to gain real-time visibility into system performance and health.

In addition, it's crucial to implement a lifecycle management approach that ties rack builds, changes, and decommissions into platforms like ServiceNow and lifecycle tools so that AI infrastructure doesn't become an opaque black box over time. For example, AHEAD's clients can leverage our Hatch IT Lifecycle Management platform to track thousands of hardware components from procurement through to deployment and ongoing operations within ServiceNow.



Address Risk, Compliance, and Sustainability

Before deploying at scale, you should proactively work to de-risk your AI infrastructure investments and ensure they're aligned with board-level priorities. This includes addressing regulatory and safety requirements related to liquid cooling systems. You should also plan for business continuity and disaster recovery for new AI workloads tied to new, denser facilities.

In terms of sustainability, it's helpful to consider how liquid-cooling helps meet ESG goals by improving energy efficiency and reducing cooling overhead compared to traditional air cooling. This means tracking water usage, power usage effectiveness (PUE), and other sustainability metrics for your AI facilities to achieve your ESG targets.



Establish the Commercial and Sourcing Model

Technology choices for liquid-cooled AI infrastructure also have significant financial implications, with each architectural decision influencing operating costs, footprint efficiency, and scalability. It's important to evaluate the [tradeoffs between capital expenses and operating costs](#) across different deployments like on-premises, colocation, and integrator-hosted models.

You'll also want to define a sourcing model that evaluates multiple vendors and ecosystem partners across GPUs, cooling systems, integration services, and other technologies related to AI infrastructure. Rather than committing to a single proprietary technology or vendor, it's crucial to prioritize a flexible and modular architecture that allows for future upgrades without overhauling the entire design.

The Complete Roadmap: From Pilot Pod to Global AI Factory

As you can see, building out liquid cooled AI infrastructure requires aligning your business strategy, facility constraints, infrastructure requirements, and operations into a cohesive plan. A phased and modular approach helps reduce risk while accelerating deployment timelines and enabling scalability as AI hardware evolves. We hope this playbook has provided a practical starting point as you seek to plan and scale your liquid-cooled AI infrastructure.

If you're looking to accelerate your AI infrastructure rollout with greater confidence, consider partnering with AHEAD. We're an enterprise solutions leader with deep experience in planning, building, and deploying high-density AI infrastructure. Our liquid-cooled rack integration facility is purpose-built to design, configure, and validate fully-integrated racks at scale. This ensures enterprise can scale AI infrastructure faster with less risk and greater efficiency.



Contact AHEAD

[Contact AHEAD](#) to learn more about our comprehensive services for liquid cooling rack integration, AI factory design, and data center modernization.

AHEAD

Engineering the Future

© 2026 AHEAD, LLC. All rights reserved.