



开放计算标准
工作委员会
Open Compute Technology
Committee

OCTC AB03—2026

OCP China Community

GW-scale Open AIDC

Blueprint White Paper

An Open, Efficient, Green, and Intelligent Architecture for GW-Scale AI Data Centers

Framework Technical Report

Prepared by: GW-scale Open AIDC Working Group (OCP China Community / OCTC)

Guided by: OCP China Community · Open Compute Technology Committee (OCTC)

Release Date: July 2026



Table of Contents

Contributors and Preparation Information	6
Executive Summary.....	8
Introduction.....	10
1. Industry Background	11
2. Chinese Characteristics of AIDC Development.....	11
3. Paradigm Shift	12
4. Vision and Objectives.....	12
Chapter 1 Infrastructure Layer	14
1.1 Building and Space Planning	14
1.1.1 Core Planning Principles.....	14
1.1.1.1 Policy Alignment Principle	14
1.1.1.2 Safety-First Principle	14
1.1.1.3 High Power Density Adaptation Principle.....	14
1.1.1.4 Network Topology Adaptation and Evolution Principle.....	15
1.1.1.5 Elastic, Flexible Design and Expansion Principle.....	15
1.1.1.6 Green, Intensive, and Efficient Principle.....	15
1.1.2 Site Selection and Site Planning	15
1.1.2.1 Basic Site Selection Requirements.....	16
1.1.2.2 Basic Site Planning Requirements	16
1.1.3 Core Building Unit Design.....	17
1.1.3.1 Functional Zoning	17
1.1.3.2 Building Design	18
1.1.3.3 Structural Design	19
1.2 Power Supply and Backup System	22
1.2.1 Power Demand Characteristics of GW-Scale AIDCs.....	22
1.2.1.1 Power Demand of AI Training and High-Density Power Racks	22
1.2.1.2 Overall Planning of the Power Supply and Backup System	22
1.2.2 Campus-Level High-Voltage Power Supply System	23
1.2.3 Campus Medium-Voltage Distribution System	24
1.2.3.1 35 kV vs 10 kV.....	24
1.2.3.2 Medium- and Low-Voltage Supply Architecture.....	24
1.2.3.3 Energy Storage Replacing Part of Diesel Generation	24
1.2.4 In-Building Power Distribution System.....	24
1.2.4.1 Evolution Path of the Supply Architecture	24
1.2.4.2 All-DC Supply Architecture	25
1.2.4.3 HVDC Energy Storage Configuration.....	25
1.2.5 Operation Protection and Control System	26
1.2.5.1 Multi-Level Protection Strategy.....	26
1.2.5.2 AIDC Intelligent Control System.....	26

1.2.6 Conclusion and Outlook	26
1.3 GW-Scale AIDC Cooling System	28
1.3.1 Cooling System Architecture and Core Equipment	28
1.3.1.1 System Architecture	29
1.3.1.2 Control System	34
1.3.1.3 Heat Exchange Equipment	37
1.3.2 Hall-Level Cooling System Engineering Validation	40
1.3.3 Future Outlook	40
Chapter 2 IT Facility Layer	42
2.1 AI Supernodes	42
2.1.1 Full-Rack Supernodes	42
2.1.1.1 Compute Nodes	42
2.1.1.2 Racks	44
2.1.1.3 Full-Rack Power Supply	45
2.1.1.4 Full-Rack Cooling	47
2.1.1.5 Scale-Up Interconnect	52
2.1.2 Distributed Supernodes	55
2.1.2.1 Modules	55
2.1.2.2 Topology	58
2.1.2.3 Power Supply	61
2.1.2.4 Cooling	61
2.1.2.5 I/O Expansion Cards	61
2.2 High-Speed Interconnect Physical Layer and Connectors	62
2.2.1 High-Speed Interconnect Networks in GW-Scale AIDCs	62
2.2.1.1 Supernode High-Speed Interconnect Hardware Architecture	62
2.2.1.2 Compute Node High-Speed Interconnect	62
2.2.1.3 Switch Node High-Speed Interconnect	63
2.2.1.4 Intra-Rack Node-to-Node High-Speed Interconnect	64
2.2.1.5 Rack-to-Rack Interconnect	66
2.2.2 Interconnect Technology Implementation Paths	66
2.2.2.1 Termination Forms: Press-fit/SMT/BGA/LGA	66
2.2.2.2 High-Speed Connectors	67
2.2.2.3 Internal High-Speed I/O	67
2.2.2.4 External High-Speed I/O	67
2.2.2.5 High-Speed PCB Laminates	68
2.2.2.6 High-Speed Raw Cable	68
2.2.3 Outlook for Next-Generation Interconnect Technology Paths	69
2.3 AI Network and Switching System	70
2.3.1 Scaling	70
2.3.1.1 Rapidly Evolving Switch Chip Bandwidth Drives Generational Leaps in AIDC Networks	70
2.3.1.2 Elastic AIDC Network Architecture Supporting Hyperscale Expansion	71
2.3.2 Communication Efficiency	73
2.3.2.1 Pure Network-Scheduled Load Balancing Technologies	74

2.3.2.2 End-Network Converged Scheduling Load Balancing Technologies	74
2.3.3 Reliability and Maintenance	75
2.3.3.1 Switch Node-Level Reliability and Maintenance	75
2.3.3.2 Network Architecture-Level Reliability and Maintenance	75
2.3.3.3 Link-Level Reliability and Maintenance	75
2.3.3.4 Service-Level Reliability and Maintenance	76
2.3.4 Future Outlook	76
2.3.4.1 Network Management and Control Software Supporting AIDC Services	76
2.3.4.2 A Healthy Ecosystem of Openness and Decoupling	77
2.4 AI Storage Technology and Systems	78
2.4.1 Overview	78
2.4.2 AI Storage Technology	78
2.4.2.1 High-Speed Inference Storage Architecture	78
2.4.3 AI Storage Nodes	79
2.4.3.1 CPU	79
2.4.3.2 Memory	80
2.4.3.3 NVMe SSD	80
2.4.3.4 RDMA NICs	82
2.4.3.5 DPU NICs	82
2.4.3.6 Storage Drive Expansion	83
Chapter 3 Network and High-Speed Interconnect Layer	84
3.1 Scale-Up Technology Analysis	84
3.1.1 Research Background and Significance of Scale-Up	84
3.1.2 Inter-Card High-Speed Interconnect Protocols and Technical Standards	85
3.1.2.1 NVLink Protocol	85
3.1.2.2 Omni-directional Intelligent Sensing Express Architecture (OISA)	85
3.1.2.3 SUE Technology	87
3.1.2.4 UALink Protocol	87
3.1.2.5 Other Inter-Card Interconnect Protocols	88
3.1.3 Future Development Trends and Application Potential	90
3.2 Scale-Out Network Technology	91
3.2.1 Scale-Out Network Requirements and Goals	91
3.2.1.1 Scale-Out Requirements	91
3.2.1.2 Goals of the Scale-Out Network	91
3.2.2 Network Operating System SONiC	91
3.2.3 Scale-Out Protocols and Software Stack	93
3.2.3.1 Communication Protocols	93
3.2.3.2 Network Congestion Control	96
3.2.4 Performance Evaluation and Tuning	96
3.2.4.1 Scale-Out Network Performance Metrics	96
3.2.4.2 Performance Tuning Technologies	97
3.2.5 Industry Frontier Cases and Trends	97
3.2.5.1 AIDC Scale-Out Network Practices	97

3.2.5.2 Future Trend Outlook	98
3.3 Scale-Across Network Technology	98
3.3.1 AIDC Interconnection Requirements	98
3.3.1.1 Compute Expansion	98
3.3.1.2 Physical Constraints	99
3.3.2 Lossless RDMA	99
3.3.2.1 Distance-Adaptive Load Balancing	99
3.3.2.2 End-to-End Congestion Control	99
3.3.3 Wide-Area Deterministic Interconnection	100
3.3.3.1 Precise Latency Management	100
3.3.3.2 Network State Awareness	100
3.3.4 Industry Frontier Solutions and Trends	101
Chapter 4 Unified Operations Management	104
4.1 Current Status and Analysis of AIDC Operations	104
4.2 Full-Rack Supernode Management Architecture	104
4.3 Supernode Rack-Level Management	105
4.3.1 Power System Management	106
4.3.2 Liquid Cooling System Management	106
4.3.3 Rack Management Protocols	107
4.4 Supernode Node-Level Management	107
4.4.1 Core Component Management	107
4.4.2 Power Management	109
4.4.3 Thermal Management	110
4.4.4 Log Diagnostics	110
4.4.5 Alarm Settings	111
4.4.6 Remote Access	111
4.4.7 Users and Security	111
4.4.8 BMC Network Management	112
4.5 Development Trend Outlook	113
Conclusion	114
Glossary	115
References and Further Reading	117
About OCTC and the OCP China Community	118

Contributors and Preparation Information

Preparing Organizations

This technical report was prepared by the GW-scale Open AIDC Working Group under the joint guidance of the OCP China Community and the Open Compute Technology Committee (OCTC). Tailored to China's national context and engineering practice, the working group systematically reviews and constructs an open, efficient, green, intelligent, and mass-replicable gigawatt (GW)-scale AIDC reference architecture and practical framework.

Principal Authors (Organizations)

OCP China Community Open Compute Technology Committee (OCTC) of the China Electronics Standardization Association Baidu Online Network Technology (Beijing) Co., Ltd.

IEIT, Inc.

AVIC Jonhon Optron Technology Co., Ltd.

VNET Group, Inc.

China Mobile Research Institute China Electronics Standardization Institute New H3C Technologies Co., Ltd.

Beijing Geek Tiancheng Technology Co., Ltd.

Jiangsu Zhiwang Technology Co., Ltd.

Guangdong Connector Association China Institute of Building Standard Design & Research Co., Ltd.

Huizhou EVE Energy Co., Ltd.

Dongguan Luxshare Technology Co., Ltd.

Qinghong Electronics (Suzhou) Co., Ltd.

Jiangsu Anlan Wanjin Electronics Co., Ltd.

Xingyin Information Technology (Shanghai) Co., Ltd.

Principal Authors and Other Contributors

The chapters of this report were jointly written, reviewed, and proofread by experts from the working group's member organizations, covering AIDC, supernodes, connectors, power supply and distribution, cooling, networking, storage, interconnect protocols, management and operations, and architectural design. The principal authors include:

Ye Yurui, Wu Zhenghui, Zhang Bin, Zhang Qun, Liang Yutong, Yuan Junfeng, Chen Xuanhao, Gao Xiaoqi, Liu Shilei, Jiang Zhengshun, Li Hongzhe, Zou Shengliang, Lei Xiong, Xie Shuang, Qi Taoran, Wang Jiasen, Si Changliang, Wu Xiaohui, Xi Xiaoguang, Ren Xiaopan, Luo Zhen, Zhang Junping, Li Hengyu, and others.

Thanks are also extended to He Yongzhan, Chen Hai, Yu Xiongjie, Li Dianlin, Bao Yi, Qin Shengyu, Qin Shijuan, Fang Siyuan, Lu Xin, Bu Yingqiong, Gao Jianguo, Wang Xinglong, Liu Shengfu, Zhang Xiaoguang, Wang Shuo, Wang Zhiqiang, Zhao Wei, Zhang Zheng, Pang Bin, Hao Yutao, Wang Chao, Chen Yue, Zhao Jingtao, Deng Haokun, Li Botao, Yang Haizhong, Miao Rui, Zhang Yang, Zhang Xianguo, Guan Zhengchao, and others for their contributions to this technical report.

Technical Review

Expert panels from the OCP China Community, the Open Compute Technology Committee (OCTC), and related organizations.

Special Acknowledgements

We thank the Open Compute Project (OCP) Open Systems for AI strategic initiative for its methodological foundation and open collaboration framework, which provided an important reference for the system-level design perspective of this technical report.

Executive Summary

A System-Level Vision and Chinese Practice for Next-Generation AI Infrastructure Artificial intelligence is evolving at an unprecedented pace. Emerging applications such as large models, AI-generated content (AIGC), and agentic AI continue to push compute demand toward ever greater scale, density, and energy efficiency. Compute infrastructure is accelerating from the traditional data center form centered on servers and clusters toward a new generation of AI data centers (AIDCs) characterized by system-level co-design. In this process, data center power is leaping from several megawatts (MW) to hundreds of megawatts and even to the gigawatt (GW) scale, and traditional architectures can no longer meet the energy-efficiency and scalability requirements of future large models, artificial general intelligence, and the agent economy.

The Open Systems for AI initiative proposed by the Open Compute Project (OCP) community emphasizes open standards, a systems perspective, and cross-layer co-design to drive an overall leap in the performance, energy efficiency, scalability, and sustainability of AI infrastructure, providing an important methodological foundation and open collaboration framework for its global development. At the same time, China is at a critical stage in the scaled development of compute infrastructure: under the "East Data, West Computing" project, it has developed unique strengths in new energy supply, high-voltage direct-current (HVDC) transmission, cross-regional compute deployment, manufacturing supply chains, and engineering organization, creating fertile ground for exploring the new form of GW-scale AIDCs.

A GW-scale AIDC is not a linear scale-up of the traditional data center, but a fundamental shift in system paradigm:

the design scale rises from the server level and rack level to the campus level and even the regional level; energy, computing, and networking are no longer independent subsystems, but must be co-designed and jointly optimized at the planning stage; at hyperscale, reliability, maintainability, and total lifecycle cost are often more decisive than single-point performance; and rising system complexity makes open interfaces and standards-based collaboration a key means of reducing systemic risk. Building a GW-scale AIDC is therefore, in essence, a systems engineering problem.

This technical report is not a specification for any specific product or single technical solution, but a system-level reference framework. The document is organized into four interdependent layers: the infrastructure layer (building and space planning, power supply and backup systems, cooling systems), the IT facility layer (AI supernodes, high-speed interconnect physical layer and connectors, AI networking and switching systems, AI storage), the network and high-speed interconnect layer (Scale-Up, Scale-Out, and Scale-Across), and the system software and resource management layer (AI system software, power and environment monitoring, intelligent autonomous operations, and unified operations management). Constrained by available capacity and time, the current version of this report focuses on these layers.

Each layer is organized around the four main themes of openness, efficiency, greenness, and intelligence, and defines system-level design principles, layered architectures, and key interface directions.

Through this technical report, the GW-scale Open AIDC Working Group hopes to distill the common problems of GW-scale AIDCs in cross-domain coordination across energy, computing, storage, and networking, and to provide an alignable, co-buildable reference foundation for the industry, operators, equipment vendors, and standards organizations, promoting two-way exchange and co-evolution between Chinese practice and the OCP global ecosystem — both "bringing in" and "going global" — thereby building industry consensus, reducing the cost of system-level innovation, fostering an open ecosystem, and providing solid support for the sustainable evolution of next-generation AI infrastructure.

Introduction

Artificial intelligence (AI) technology is evolving at an unprecedented pace, and new application forms such as large models, AI-generated content (AIGC), and agentic AI continue to push compute demand toward greater scale, higher density, and higher energy efficiency. Globally, computing infrastructure is accelerating from the traditional data center form centered on servers and clusters toward a new generation of AI data centers (AIDCs) characterized by system-level co-design.

Against this backdrop, the Open Compute Project (OCP) community launched the Open Systems for AI initiative, which emphasizes open standards, a systems perspective, and cross-layer co-design to drive an overall leap in AI infrastructure performance, energy efficiency, scalability, and sustainability. This initiative provides an important methodological foundation and an open collaboration framework for the development of AI infrastructure worldwide.

At the same time, China is at a critical stage in the scaled development of compute infrastructure. On the one hand, new demand from large-model training and inference, industrial intelligent upgrading, and agentic applications is driving the continuous construction of hyperscale AIDCs; on the other hand, China has developed unique strengths in new energy supply, HVDC transmission, cross-regional compute deployment, manufacturing supply chains, and engineering organization, creating fertile ground for exploring the new form of gigawatt (GW)-scale AIDCs.

In this context, under the joint guidance of the OCP China Community and the Open Compute Technology Committee (OCTC), the GW-scale Open AIDC Working Group was formally established. Tailored to China's national conditions and engineering practice, it aims to systematically review and construct an open, efficient, green, intelligent, and mass-replicable GW-scale AIDC reference architecture and practical framework.

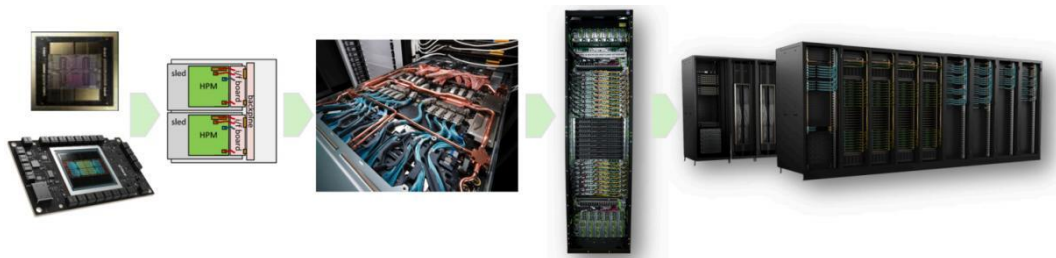
This technical report is not a specification for any specific product or single technical solution, but a system-level reference framework whose objectives are to:

distill the common problems of GW-scale AIDCs in cross-domain coordination across energy, computing, storage, and networking; clarify system-level design principles, layered architectures, and key interface directions; provide an alignable, co-buildable reference foundation for the industry, operators, equipment vendors, and standards organizations; and promote two-way exchange and co-evolution between Chinese practice and the OCP global ecosystem. Through this technical report, the GW-scale Open AIDC Working Group hopes to build industry consensus, reduce the cost of system-level innovation, foster an open ecosystem, and provide solid support for the sustainable evolution of next-generation AI infrastructure.

1. Industry Background

The training scale and computing demand of AI models are growing exponentially, driving simultaneous leaps in compute density, network bandwidth, energy efficiency, and green compliance requirements:

(1) Training cluster size: from ten thousand GPUs to a hundred thousand to a million. (2) High-speed interconnect: from 400G to 800G to 1.6T and beyond. (3) Cooling: from air cooling to cold plate to immersion. (4) Data center power: from several MW to hundreds of MW and toward the gigawatt scale. Traditional data center architectures will be unable to meet the energy-efficiency and scalability demands of future large models, AGI, and the agent economy, and new forms of data infrastructure centered on "energy-computing integration" are emerging worldwide. AI computing architecture design must evolve from the chip level and server level to the rack level and system level, as shown in the figure below.



Introduction – Figure 1

2. Chinese Characteristics of AIDC Development

China enjoys a set of unique conditions for exploring hyperscale AIDC forms: (1) Coordinated advancement of energy structure and compute deployment. Under the "East Data, West Computing" project, China has accumulated system-level engineering experience in cross-regional compute deployment, power dispatching, and network coordination. The northwest and north China regions possess concentrated, contiguous wind and photovoltaic resources, providing a sustainable, low-carbon energy base for hyperscale AIDCs; meanwhile, HVDC transmission and grid dispatching systems create realistic conditions for cross-domain coordination of computing and energy.

(2) Advantages in manufacturing and key component supply chains. China has built relatively complete and internationally competitive industrial chains in key areas such as complete servers, liquid cooling systems, connectors, optical modules, power supplies, and power devices. This advantage provides important support for high power density, high-bandwidth interconnect, and rapid scaled delivery.

(3) Strong application-driven demand. AI applications in China are highly diverse, spanning internet and cloud computing scenarios as well as industrial intelligence needs across manufacturing, energy, transportation, and finance. This pattern of parallel multi-scenario development requires AIDC infrastructure to be highly general-purpose, scalable, and capable of long-term evolution.

(4) Policy and demonstration-project leadership. Governments at all levels regard AI and compute infrastructure as important strategic directions, accelerating the construction of new-type AIDCs through policy guidance, demonstration projects, and industrial coordination, thereby providing a realistic testing ground for system-level innovation. Beginning in 2026, China introduced policies related to "computing-power and electricity coordination."

Taken together, these factors mean that in advancing the evolution of AI infrastructure, China needs not only "bigger data centers" but also a systematic architectural methodology for the gigawatt scale.

3. Paradigm Shift

A GW-scale AIDC is not a linear scale-up of the traditional data center, but a fundamental shift in system paradigm, whose core characteristics are:

- (1) An expanded design scale. Architecture design rises from the server level and rack level to the campus level and even the regional level.
- (2) Cross-domain coordination as a precondition. The energy, computing, and networking systems are no longer independent subsystems, but must be co-designed and jointly optimized at the planning stage.
- (3) Priority on system stability and operability. At hyperscale, reliability, maintainability, and lifecycle cost are often more decisive than single-point performance.
- (4) The growing importance of openness and standardization.

As system complexity rises, closed, customized solutions become difficult to sustain over the long term, and open interfaces and standards-based collaboration become key means of reducing systemic risk.

Building a GW-scale AIDC is therefore, in essence, a systems engineering problem that requires holistic planning across technology, energy, networking, operations, and ecosystem at the architectural level.

4. Vision and Objectives

Based on the above background, the GW-scale Open AIDC Working Group proposes the following overall vision: (1) Build an open reference architecture for the gigawatt scale. Through systematic layered design, clarify key capability boundaries and interface directions, and reduce the cost of cross-entity collaboration. Promote the standardization and normalization of China's AIDC infrastructure through an open system, and drive the re-globalization of the supply chain.

(2) Promote a green development model that coordinates energy and computing: make energy efficiency and sustainability core constraints of system-level design, achieve the coordinated evolution of compute growth and green goals, and advance the sustainable development of intelligent computing power.

(3) Support the long-term evolution of diverse AI application scenarios: meet the differentiated demands of large models, agents, and industrial intelligence on computing systems. (4) Promote the co-building of Chinese practice and the global standards system: on the basis of aligning with OCP's global methodology, summarize Chinese engineering experience and foster systemic solutions that can be understood and adopted internationally — that is, both "bringing in" and "going global."

Chapter 1 Infrastructure Layer

1.1 Building and Space Planning

Building and space planning for a gigawatt-scale AI data center (AIDC) is campus-level civil engineering design and coordination for gigawatt-scale total power, high-power-density clusters, and energy-computing integration. It focuses on five dimensions — site layout, individual building space, modular adaptation, fire protection and security, and long-term evolution — providing a benchmark for the civil and spatial design of cross-regional, large-scale, standardized AIDC infrastructure.

1.1.1 Core Planning Principles

Building and space planning for GW-scale AIDCs breaks from the design logic of traditional MW-scale data centers. Centered on system design and measured at the campus scale, it establishes six core principles, all focused on building spatial form, layout, scale, and compliance.

1.1.1.1 Policy Alignment Principle

Fully align with national new-type data center construction, the "East Data, West Computing" national computing hub layout, "computing-power and electricity coordination" policy requirements, dual-carbon goals, and local special policies for compute infrastructure. Building space planning must meet hard policy requirements such as green power consumption, HVDC power access, liquid cooling deployment, and energy conservation. Core indicators such as building classification, fire protection, seismic resistance, and evacuation must strictly follow the national Class A data center standards, ensuring that a GW-scale campus complies from the planning stage with national and regional compute infrastructure requirements and serves as the spatial carrier for the national computing network strategy.

1.1.1.2 Safety-First Principle

With Class A data center reliability as the baseline for building space design, build an all-dimensional safe spatial system across site layout, individual structures, functional zoning, and fire evacuation. For the four high-risk zones of a GW-scale campus — the 220 kV high-voltage power supply area, the large-capacity energy storage area, the liquid cooling piping area, and the high-density compute cluster area — use civil engineering measures such as physical isolation, spatial separation, load strengthening, and reserved space for leak/flood prevention to solidify the spatial safety foundation of the hyperscale campus.

1.1.1.3 High Power Density Adaptation Principle

Adapt building space in all dimensions to the high-power-density characteristics of GW-scale AIDCs — over 140 kW per rack and total campus power of at least 1000 MW — focusing on four civil engineering

indicators: structural load, interior clear height, column grid dimensions, and piping space. By enlarging building spatial scales, strengthening structural bearing capacity, and reserving vertical piping space, meet the installation and O&M space requirements of wide-body racks, liquid-cooling CDU units, and rack-level backup power equipment.

1.1.1.4 Network Topology Adaptation and Evolution Principle

The site plan and individual building layouts must synchronously match the AIDC's high-speed network topology. Core compute zones, network switching zones, and storage zones should be physically concentrated at close range to reduce network interconnect loss by optimizing spatial distance. Inside buildings, reserve dedicated space and vertical piping channels for fiber distribution, high-speed interconnect racks, and cross-cluster links, supporting the spatial evolution of network bandwidth across generations and accommodating network architecture iteration without later civil modification.

1.1.1.5 Elastic, Flexible Design and Expansion Principle

Adopt a two-dimensional elastic design approach combining campus-level master planning with modular subdivision of individual buildings. Building space is divided by Pod compute units, supporting phased deployment and on-demand expansion. Building structures, power distribution piping, cooling pipe networks, and fire protection systems all use pooled, reserved design. The spatial layout can flexibly adapt to dynamic adjustment of training/inference clusters and to spatial switching between air-cooled and liquid-cooled hybrid deployments, avoiding the space waste of one-time construction and improving the space utilization and operational flexibility of a GW-scale campus.

1.1.1.6 Green, Intensive, and Efficient Principle

With reducing PUE and WUE as the core goals of building space planning, the site layout prioritizes the use of natural cooling sources and optimizes the heat transfer spatial paths between the cooling system and compute zones. Building envelopes use high-efficiency thermal insulation to reduce cooling energy loss. Functional zoning follows the principle of intensive compactness, strictly controlling the area of auxiliary and supporting zones to improve land and building space efficiency. Differentiated spatial design is applied according to the geographic characteristics of waterless western regions and water-rich southern regions, achieving the dual goals of green low-carbon operation and spatial efficiency.

1.1.2 Site Selection and Site Planning

This planning is the prerequisite foundation for a GW-scale AIDC's building space, addressing four civil engineering issues: campus siting conditions, site layout, vertical design, and pipe network coordination.

1.1.2.1 Basic Site Selection Requirements

1. Core resource alignment: prioritize sites within the national computing hub nodes of the "East Data, West Computing" project. The site must have campus-level power access conditions — see Section 1.2 for specific access capacity and voltage levels — and be close to regions rich in new energy. Trunk communication connectivity must be guaranteed to meet the network transmission needs of cross-regional compute scheduling. Water-rich areas must have stable industrial water quotas; waterless or water-scarce areas must reserve installation space for water-free cooling systems.
2. Environmental safety isolation: the site must be far from dust, harmful gases, corrosive and flammable/explosive premises, and must avoid areas with geological hazard risks, ensuring the structural safety of the building site from the source.
3. Interference avoidance: the site must be far from strong vibration sources, strong noise sources, and strong electromagnetic interference zones, and must maintain compliant safety distances from residential buildings and public facilities, meeting the spatial environment requirements for quiet AIDC operation.

1.1.2.2 Basic Site Planning Requirements

1. Building layout. The building layout may adopt the main equipment hall building plus office/support building model (hereinafter the traditional model) or the building cluster model. The building cluster model splits the equipment rooms and power rooms of the traditional main building into separate buildings laid out as a cluster. When this model is used, the site plan should adopt a centripetal building cluster with physical isolation — the equipment hall building sits at the center of the campus with power buildings arranged around the IT buildings — forming a standardized civil spatial structure in two typical layout modes:

- 1) High-rise single-building centripetal layout: a single high-rise equipment hall building serves as the geometric center of the campus, with power buildings symmetrically arranged on all sides, forming a "one core, four auxiliaries" symmetric structure. In this layout, the piping paths between power buildings and the core equipment hall building are equal and the spatial distances shortest, greatly reducing the civil installation length of power, cooling, and communication pipelines.
- 2) Multi-story group centripetal layout: two to three multi-story equipment hall buildings form the core group of the campus, surrounded by power buildings, forming a "multi-core, multi-auxiliary" expandable structure. This layout supports phased construction: new equipment hall and power building groups can be added as compute demand grows, without adjusting the campus's main pipeline routing. It fits the elastic expansion principle of GW-scale campuses while keeping the spatial distances between equipment hall buildings and power buildings consistent, enabling standardized O&M and pipeline layouts.

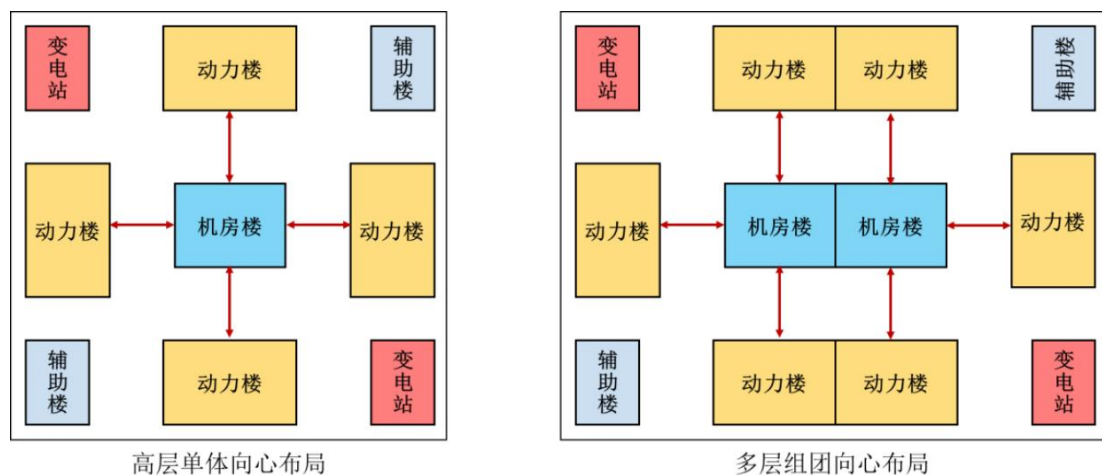


Figure 1.1 Schematic of the centripetal building cluster layout

2. Entrances and circulation. The campus sets independent pedestrian, vehicle, and logistics entrances for three-way separation. The main entrance reserves turning and maneuvering space for transporting large prefabricated modules and heavy racks. Fire lanes must be at least 4 m wide with a turning radius of at least 15 m accounting for freight, forming a continuous fire loop around the campus.

3. Site vertical design. Site drainage slope must be at least 0.5% to eliminate standing water. The indoor-outdoor height difference is determined rationally from terrain and 100-year flood level information; in typhoon regions, appropriately increase the indoor-outdoor height difference and add anti-backflow measures.

4. Integrated pipe network. Coordinate HV power supply, liquid cooling/fluorine circulation, communication optical cables, fire protection, and waste-heat recovery pipelines, using a centralized utility tunnel. The pipe network routing matches the building layout and network topology, with expansion interfaces reserved to avoid later modification that would damage civil structures.

1.1.3 Core Building Unit Design

Core building units are the central carriers of compute, energy assurance, and cooling operation in a GW-scale AIDC; their design directly determines the campus's operational efficiency, safety and reliability, and long-term evolution capability. This section focuses on five core civil engineering dimensions — functional zoning, spatial scale, structure, envelope, and interior fit-out — combining national standards with GW-scale campus practice to detail design requirements and clarify their basis, providing dedicated, compliant, and implementable building space support for GW-scale computing, energy, cooling, and O&M systems.

1.1.3.1 Functional Zoning

The main equipment hall building should adopt a centripetal single-building layout with the main equipment halls at the center and supporting functions around the perimeter. It is divided into four core zones by function, each spatially independent with clear boundaries. The detailed design requirements and basis for each zone are as follows:

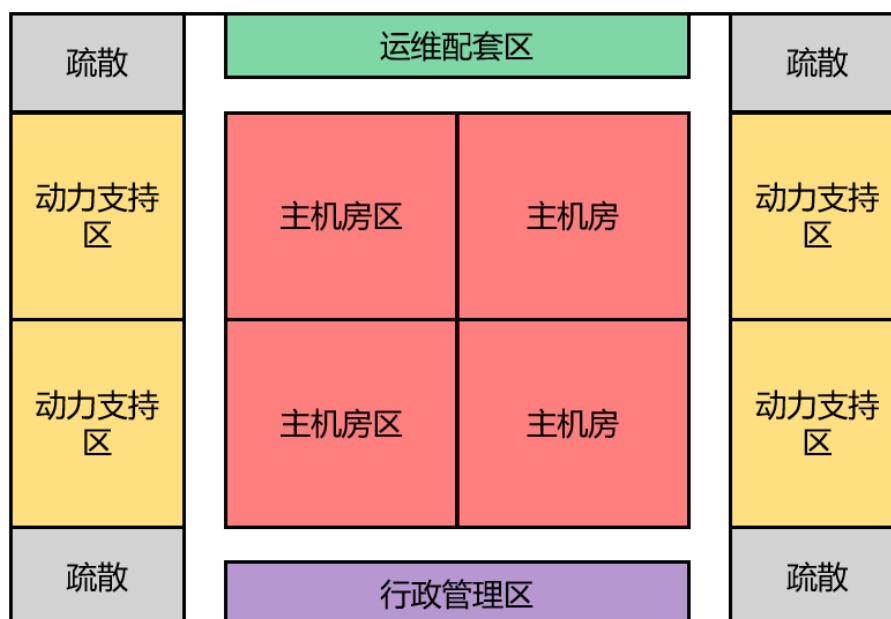


Figure 1.2 Schematic of the centripetal single-building layout

1. Core main equipment hall zone: as the core space of the compute building, it occupies 20%-30% of the total area of each equipment hall building. The main equipment hall carries high-density rack clusters and should be laid out close to the power support zone to shorten pipeline transmission paths. The space should be regular in shape, suit batch deployment of wide-body racks and full-rack supernodes, and reserve rack row spacing and O&M operating space.
2. Power support zone: this may be an independent building or an independent zone within an equipment hall building; an independent power building must maintain fire separation distance from the compute building. When located inside the building, the support zone occupies 40%-50% of the total building area, growing as per-rack power increases. Diesel generator sets may use outdoor containerized form.
3. O&M support zone: occupying 8%-10% of the total building area, it meets the space needs of daily O&M, monitoring, equipment commissioning and fault handling, and spare parts storage.
4. Administration zone: set as needed, occupying 5%-10% of the total building area. Office areas must be physically isolated from core compute zones and energy assurance zones to avoid interference.

1.1.3.2 Building Design

1. Unit scale: the floor area of each core equipment hall building in a GW-scale campus should be at least 20,000 m², modularly subdivided by Pod compute units, with each Pod's space suited to 50-100 MW power load deployment. Buildings should be 3-5 stories: more stories increase vertical transport costs and cooling pipe network losses, while fewer waste land resources. Building length should not exceed 120 m and width should not exceed 60 m, facilitating fire evacuation and equipment O&M, suiting the transport and installation of modular prefabricated components, and conforming to the green, intensive, and efficient principle.

2. Plan layout: when the core compute zone uses non-immersion cooling, it should adopt a row layout with enclosed hot/cold aisles, racks arranged face-to-face and back-to-back to form clear hot and cold aisles, with cold aisle width at least 1.5 m and hot aisle width at least 1.8 m to facilitate cooling circulation and heat exhaust. Liquid-cooled racks and CDU units should be placed close together to shorten secondary-side piping distance and reduce cooling loss and pipe resistance. Liquid cooling rooms and CDU equipment areas use locally raised, sloped floors for drainage, with anti-flood drainage space. Dedicated piping shafts should be reserved inside buildings to centrally route power, cooling, and network pipelines, avoiding exposed pipelines that impede O&M and space utilization. Shaft dimensions should be at least 1.2 m × 1.2 m for pipeline installation and maintenance.

3. Story height and column grid: a story height of 7.2 m per floor is recommended, reserving vertical installation space for liquid-cooled racks (about 2.5 m high), overhead cable trays (about 2.9 m high), and CDU equipment (about 2 m high), plus a recommended 10%-15% vertical redundancy for later equipment upgrades. Column grid spacing should be chosen in conjunction with rack arrangement, ensuring a reasonable number of racks per row and sufficient O&M space.

4. Building envelope: exterior walls should use Class A fireproof insulation materials, preferably high-efficiency insulation, to reduce cooling energy consumption. Liquid cooling areas should add a secondary anti-leakage structural layer to prevent cooling medium leaks from damaging the building structure. Doors and windows should use high-airtightness, dustproof, and soundproof designs to reduce heat conduction and external interference and meet equipment room cleanliness requirements.

5. Interior fit-out: all areas should use Class A non-combustible, anti-static, dustproof, and moisture-proof finishing materials to eliminate fire hazards and dust pollution.

Floors should use anti-static flooring; liquid cooling areas add anti-slip surfaces and leak drainage channels connected to collection sumps for collecting and handling leaked medium. Walls should use fireproof, moisture-proof, dustproof latex paint or color steel panels, preferably in light colors to enhance indoor lighting.

1.1.3.3 Structural Design

1. Structural design standards: for new GW-scale campus building units, the structural safety class should be Class I; the seismic fortification category should be no lower than Category B (key fortification) for new buildings and no lower than Category C (standard fortification) for renovated buildings; the structural design standards of supporting buildings (structures) should be consistent with the main equipment halls;

2. Rational structural layout: 1) Neither new GW-scale AIDCs nor their supporting building units should adopt single-span structures or structural forms dominated by partially single-span frames;

2) Given the large load and clear-height requirements of GW-scale AIDCs, the structural column grid can use different spans in the X and Y directions according to rack arrangement: maximize the span in one direction to fit the most rack rows, and appropriately reduce it in the other. In the long-span direction,

secondary beam spacing can be densified; in the short-span direction, the main load-bearing beams should have increased frame beam section depth. Optimal span modules include 9.6 m × 6.0 m (short direction) and 12.0 m × 7.2 m (short direction);

3) Structural joint layout: structural joints are strictly prohibited from crossing main equipment halls; the overall building should avoid structural joints wherever possible through appropriate measures;

4) Given the characteristics of GW-scale equipment halls — heavy loads, large spans, tall story heights, and high investment — the structure should adopt seismic isolation/energy dissipation design. Isolation applies to the whole building, and vibration isolation pads should also be used for heavy individual racks, or overall isolation implemented at the raised floor level; at least one of these isolation/energy dissipation measures should be adopted;

5) Relatively heavy equipment should be placed on lower floors and lighter equipment on the top floor. 6) For roof structures, engineering experience shows that structural sloping is advantageous, freeing up load margin that can be allocated to large rooftop equipment;

3. Load values: 1) When performing structural calculations for new GW-scale AIDCs, the standard values of floor live load for main functional rooms should not be less than the following:

Liquid-cooling CDU (main equipment hall; 2-ton rack) $\geq 16 \text{ kN/m}^2$; liquid-cooling CDU (main equipment hall; 3-ton rack) $\geq 20 \text{ kN/m}^2$; diesel generator sets (generally built separately outdoors) floor live load $\geq 24 \text{ kN/m}^2$; UPS room 10 kN/m^2 ; gas cylinder room 8 kN/m^2 ; battery rooms and energy storage equipment areas floor live load 24 kN/m^2 ; power distribution rooms 16 kN/m^2 ; air-conditioning areas 8 kN/m^2 ; equipment transport corridors 12 kN/m^2 . For equipment hall areas and diesel generator areas, a dynamic factor of 1.05 should be considered;

Suspended loads under slabs in equipment hall areas: $2.0\text{-}4.0 \text{ kN/m}^2$. Large rooftop equipment (dry coolers, fluorine pumps, cooling towers, etc.) should participate in structural calculations at actual weight.

2) For new GW-scale AIDCs, the combination value, frequent value, and quasi-permanent value coefficients of live loads may be taken per Appendix 1 of the Code for Design of Data Centers. When calculating seismic action, the combination value coefficients of the above variable loads are as follows:

Main equipment hall areas / air-conditioning areas / substations and distribution: 0.9; battery rooms / UPS rooms / pipeline suspension / diesel generators: 0.9-1.0; equipment transport corridors: 0.7. 3) Given the high load levels of GW-scale AIDCs but the spacing requirements in the plan layout of racks and equipment, when calculating members at different structural positions, refer to Table 8.2.2 of the Code for Design of Communication Building Engineering: for rooms with loads greater than 8 kN/m^2 , the live load reduction factor may be 0.75 for main frame beams, columns, walls, and foundation members, and 0.85 for secondary beams; floor slab calculations are not further reduced;

4) For renovation projects of existing buildings, the above live load values need not apply. In design, calculations may be entered as local slab loads based on actual equipment weight and specific positioning;

however, considering maintenance, installation, and liquid leakage risks, a small amount of uniformly distributed live load should be additionally superimposed;

5) Since the design working life of equipment hall buildings generally exceeds the lifecycle of the equipment, and as data industry technology updates iterate ever faster, it is recommended — where feasible — to appropriately relax floor live load values during design and use the upper-limit standard wherever possible, so that layouts can be flexibly rearranged within the building's working life;

4. Material selection: given the high load levels of GW-scale AIDCs, when reinforced concrete structures are used, the main load-bearing members should use HRB500/HRBF500 grade rebar; when steel structures are used, Q355 or higher high-strength steel should be used;

5. Structural design and detailing requirements: 1) When equipment halls require raised (overhead) structures, the design standards and load values of the raised structure should be consistent with the main structure; raised (steel frame) structures must also undergo ultimate limit state and serviceability limit state verification;

2) When heavy equipment acts on the main structural floor slab through raised structures (concentrated forces), the slab's punching shear capacity should be additionally checked;

3) Structural calculations must fully consider the influence of the raised floor construction on the main structure: the mass point for seismic action is actually located at the top surface of the raised floor, so the definition of the stirrup densification zone of the main frame columns must include the raised height. When necessary, model the raised structure and main structure together for an integrated calculation;

4) Seismic supports and hangers for pipelines should have member sizes and spacing determined by calculation. 5) Given the tall story heights of GW-scale AIDC structures, stairs generally have at least 3 or 4 flights; in structural calculations, take special care that stair self-weight and stair live loads are counted as double or triple according to the number of flights;

6) If a GW-scale AIDC requires full-rack delivery, lifts must be arranged in the transport corridor areas of the equipment halls, and structural members related to lifting equipment should consider a dynamic factor of 1.05-1.2;

7) Per substation and distribution room design requirements, areas with transformers should have lifting openings and hooks. Hooks should be mounted on structural beams, and hook beams should account for lifting-condition loads. Lifting opening dimensions should fully consider operating buffer space;

8) When equipment halls use steel structures, fireproofing and anti-corrosion design should be carried out per the corresponding fire protection standards.

1.2 Power Supply and Backup System

1.2.1 Power Demand Characteristics of GW-Scale AIDCs

1.2.1.1 Power Demand of AI Training and High-Density Power Racks

GW-scale AIDCs have become the core construction form of hyperscale compute infrastructure. Their power demand differs essentially from that of traditional megawatt-scale data centers, centered on ultra-high per-rack power, dynamic load characteristics, gigawatt-scale supply, and refined redundancy requirements — the core origin and underlying basis of power supply and distribution architecture design.

The load rate of AIDCs fluctuates much more than that of traditional data centers. When a training task starts, load can climb rapidly from trough to peak; when the task ends, load falls back quickly. When large numbers of racks start simultaneously or training tasks enter peak phase in sync, the power supply system faces large-scale, high-amplitude transient power shocks. The supply and distribution system must therefore have extremely strong dynamic load adaptation, fast voltage regulation, and load balancing and dispatching capabilities, effectively damping transient shocks at the architectural level.

1.2.1.2 Overall Planning of the Power Supply and Backup System

The power capacity planning of a GW-scale AIDC must adapt to the rapid growth of AI compute, following the core principles of "forward planning, phased construction, flexible expansion, and reliable redundancy." The campus's external power access capacity should meet the full supply demand of core loads, with expansion space reserved to accommodate the stepwise growth of compute demand over the next 3-5 years.

As shown in Figure 1.3, the power supply and backup system of a GW-scale AIDC is organized by voltage level from high to low, and the chapters are divided according to deployment location, complete power supply chain, and system functional positioning, with each section describing system architecture, configuration characteristics, and operating logic. The campus trunk interconnect layer is the main artery of power transmission in a GW-scale AIDC, handling power transformation from 220 kV to 10 kV or 35 kV. The building branch interconnect layer is the branch transmission network connecting the campus trunk ring network to each equipment hall building, handling distribution from the campus central substation to the distribution rooms in each building at 10 kV or 35 kV, then converting AC to 800 V DC via energy routers or solid-state transformers (SSTs). The in-building distribution interconnect layer is the trunk distribution network inside each building, distributing power from the building's distribution room to each rack zone at 800 V DC. In-rack power supply is described in detail in later chapters.

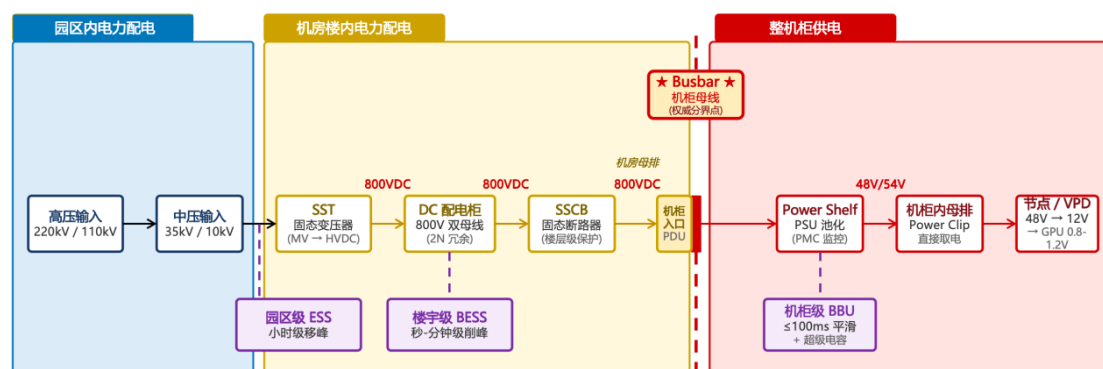


Figure 1.3 Schematic of the power supply and backup system segments

1.2.2 Campus-Level High-Voltage Power Supply System

A GW-scale AIDC campus achieves reliable power supply, efficient transmission, flexible dispatch, and low-carbon consumption by building an integrated overall architecture with source-grid-load-storage coordination. Based on a multi-source coordinated supply system, it gradually increases the share of new energy sources such as wind, solar, hydro, and nuclear, and deploys energy storage to raise the green power consumption ratio and reduce carbon emissions.

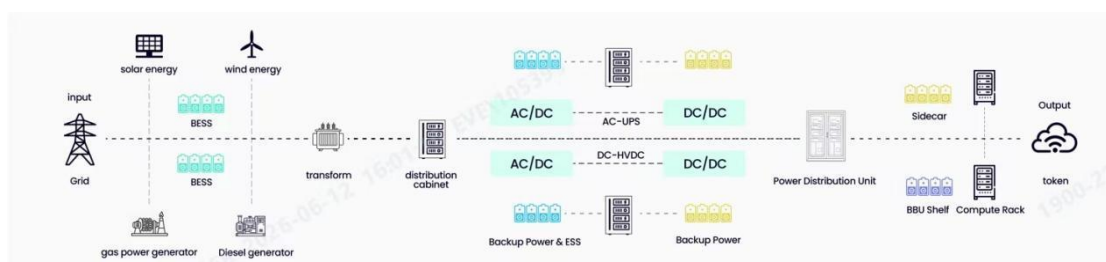


Figure 1.4 Schematic of the multi-level energy storage coordinated layout

As shown in Figure 1.4, the storage side builds a three-level "campus-level, building-level, rack-level" coordinated energy storage layout. Campus-level long-duration storage handles grid peak shaving, long-duration backup, and virtual power plant operation. Building-level short-duration storage handles peak-valley arbitrage, power quality management, and short-duration backup. Rack-level BBUs provide short-term energy buffering, transient power support, and backup power transition support, forming an all-time, all-scenario energy storage assurance and dispatch system.

Through the campus-level energy management system, data interconnection and coordinated dispatch achieve optimal power distribution and maximum energy efficiency under normal conditions, uninterrupted load assurance under grid fault conditions, and maximum renewable energy consumption and power trading operations in low-carbon scenarios — fully adapting to the large-capacity, high-reliability, low-carbon operating needs of GW-scale AIDCs.

1.2.3 Campus Medium-Voltage Distribution System

1.2.3.1 35 kV vs 10 kV

GW-scale AIDCs cover large areas, have dispersed equipment hall buildings, and carry large loads per building. The medium-voltage distribution system must balance three core needs: low-loss, large-capacity long-distance transmission; flexible dispatch across multiple load zones; and rapid isolation in fault scenarios. Although 10 kV distribution still dominates today, 35 kV — the mainstream medium-voltage level in China — offers differentiated advantages concentrated in large capacity, low loss, and long-distance transmission. At the same current-carrying capacity, 35 kV transmits more than three times the capacity of 10 kV, greatly saving investment in cables, pipe galleries, and switchgear.

1.2.3.2 Medium- and Low-Voltage Supply Architecture

A GW-scale AIDC campus adopts a two-level layout of "campus centralized main substation + distributed attached distribution rooms in equipment hall buildings." The medium-voltage side may use a double-bus sectionalized scheme, with the two buses coming from different sections of the campus 35 kV/10 kV system and backing each other up. The low-voltage DC side uses a double-bus scheme, each DC bus section powered by an independent rectifier group, forming a dual-redundant DC supply architecture that directly powers the equipment hall buildings.

1.2.3.3 Energy Storage Replacing Part of Diesel Generation

For the backup power requirements of GW-scale AIDCs, adopt a backup architecture of "long-duration storage as primary, retaining some diesel generation for emergency backup": configure 2-8 hour energy storage for the campus's core IT loads to serve as the priority backup during external power interruptions, and configure diesel generators at 50% of core load capacity, started only in extremely low-probability scenarios such as prolonged grid failures or extreme energy storage failures.

Energy storage uses prefabricated cabin modular form. Large-capacity batteries in the market are evolving toward cell capacities above 600 Ah; on the basis of high energy density, they enable minimalist integration with improved system reliability, safety, and economics — an ideal energy storage route for GW-scale AIDCs.

1.2.4 In-Building Power Distribution System

1.2.4.1 Evolution Path of the Supply Architecture

The traditional 380 V low-voltage AC supply architecture has a long supply chain with many energy conversion stages, facing efficiency bottlenecks in high-power scenarios. The industry has gradually introduced 240 V/336 V DC schemes, powering servers directly from a DC bus with batteries directly attached.

By further raising the output voltage level — such as 800 VDC or ± 400 VDC — data centers effectively reduce transmission current and line losses while meeting higher power demand, improving system efficiency. High-voltage DC architectures also have a natural advantage in direct coupling with DC energy systems such as photovoltaics and storage, restructuring traditional supply modes at the system level and becoming an important evolution direction for data centers with high compute density, high efficiency, and green sustainable development.

1.2.4.2 All-DC Supply Architecture

The currently promising 800 VDC or ± 400 VDC supply architecture uses energy routers or solid-state transformers (SSTs) to convert 10 kV medium-voltage AC directly into stable, controllable 800 VDC or ± 400 VDC, then powers servers at 48 V or lower via high-efficiency DC-DC modules. This architecture significantly shortens the supply chain and reduces AC-DC conversion stages. In efficiency, traditional UPS and HVDC overall efficiency is typically 95%-97.5%, while SST-based 800 VDC or ± 400 VDC architectures achieve measured efficiency above 98%, with significant annualized energy savings and an end-to-end system efficiency improvement of about 3%-5%.

In materials, the high-voltage, low-current all-DC scheme uses about one-third the copper of traditional AC schemes. In power density and space utilization, reducing the number of supply devices cuts footprint by 40%-50%. In construction difficulty, modularized, prefabricated systems greatly shorten construction cycles. In energy terms, new energy and storage can connect in DC form, reducing unnecessary AC-DC conversion stages.

1.2.4.3 HVDC Energy Storage Configuration

HVDC-based energy storage must meet 5-15 minute (6C-10C discharge rate) short-duration backup needs, supporting backup power switchover transitions and load transient power fluctuations. As campus scale and per-building power density continue to rise, storage systems must provide higher backup capacity within limited deployment space, placing higher demands on high energy density, high rate capability, high reliability, high safety, and intelligent monitoring.

Lead-acid batteries can no longer meet current needs. The lithium iron phosphate (LFP) route — leveraging its large-scale deployment base in the energy storage industry and especially the development of stacking technology — offers high-rate discharge, low temperature rise, high safety, and real-time status monitoring, making it the main technical direction for large-scale short-duration backup in data centers.

1.2.5 Operation Protection and Control System

1.2.5.1 Multi-Level Protection Strategy

The core protection challenge of an all-DC power supply system is that DC current has no natural zero crossing, making arcs hard to extinguish during short circuits, and fault current rises extremely fast. With its microsecond-level breaking speed and arc-free breaking advantage, the solid-state circuit breaker has become the core protection device of all-DC supply systems. Combined with solid-state circuit breakers, a multi-level, full-chain protection strategy spanning campus level, building level, rack level, and device level achieves rapid fault isolation with minimal impact.

From device level to campus level, protection operating times are set in a progressively increasing time gradient, ensuring that lower-level protection acts first, with upper-level protection acting as backup only when the lower level fails to operate. Each protection level has fault direction discrimination, acting only on faults within its protection zone and not on through faults outside it. Through the campus SCADA system, all protection devices share data and coordinate. When a fault occurs, the fault point can be precisely located with synchronized commands, achieving precise fault isolation and automatic load transfer, further improving the protection system's reliability.

1.2.5.2 AIDC Intelligent Control System

Through a full-lifecycle digital management system for the power supply and distribution system, combined with the campus-level energy management system, the transition from "passive O&M" to "active and predictive O&M" is achieved. Precisely reproducing the full-chain supply architecture enables real-time mapping and two-way interaction between the physical system and its digital model. Core functions and applications include full-system real-time visualization, multi-scenario simulation and solution verification, remote control and unattended O&M, and integrated energy management and power trading operations — comprehensively improving system reliability, O&M efficiency, and energy efficiency, while enabling additional economic functions such as maximum renewable energy consumption and power trading in low-carbon scenarios.

1.2.6 Conclusion and Outlook

Combining current industry technology trends, development mainlines and technical gaps, and engineering implementation needs, the core technology trends of GW-scale AIDC power supply and backup systems can be summarized in six directions, covering architecture upgrades, safety protection, energy storage backup, standards development, intelligent O&M, and low-carbon transformation.

1. High-voltage all-DC high-efficiency power supply. 800 V all-DC supply comprehensively replaces traditional AC UPS supply. Centered on the energy router SST, it simplifies the supply chain and improves system efficiency and power density, adapting to megawatt-scale single-rack supply needs. Future work will explore wide-bandgap semiconductor devices in supply rectification and protection, validate at scale

35 kV AC direct-to-800 VDC conversion, and continuously optimize the campus-rack-chip multi-level full-chain supply architecture to improve power transmission efficiency in all dimensions.

2. Refined safety protection for high-voltage DC. DC protection technology is iterating toward microsecond-level, full-chain, intelligent operation. Solid-state circuit breakers will deliver full-chain refined protection, DC short-circuit arc extinguishing, and multi-level protection coordination. Future work will focus on 800 V and above HVDC scenarios, emphasizing fast short-circuit detection, arc-free breaking, and precise ground-fault location, optimizing multi-level intelligent protection coordination, and improving the HVDC supply safety protection system and industry technical standards.

3. Multifunctional energy storage backup integration. Energy storage has iterated from single emergency backup to a multifunctional integrated system of "backup + peak shaving + power quality management + virtual power plant + green power consumption." The "long-duration storage + backup power" model makes backup systems low-carbon and highly economical. Future work will research large-scale application of hybrid storage systems, optimize capacity configuration and coordinated dispatch strategies, build multi-level storage safety protection systems, and improve the intrinsic safety of backup power systems.

4. Standardization of power supply and distribution interconnect. For all-DC supply architectures, the industry is accelerating unified interconnect standards, standardizing the parameters and design of core components such as device interfaces, transmission lines, and connectors to achieve interconnection and compatibility across the whole industry chain and reduce engineering application costs. Standardization research on all-DC supply and distribution systems is underway; future work will develop unified technical parameters, test standards, and design specifications for 800 V DC power supplies, busways, and connectors, solving the pain points of inconsistent interfaces and poor compatibility and driving large-scale adoption of all-DC technology.

5. All-scenario intelligent O&M and control. Leveraging digital twins, AI large models, and IoT, the industry is gradually building full-lifecycle digital management systems for power supply and distribution, enabling visual monitoring, predictive O&M, and intelligent coordinated dispatch. Future work will focus on real-time closed-loop control between digital twins and physical systems, developing AI large-model-driven intelligent dispatch algorithms to maximize system operating efficiency and economic benefit, and deepening intelligent fault diagnosis and predictive maintenance to help supply and distribution systems achieve zero-fault operation.

6. Building a zero-carbon power supply and distribution system.

Deep integration of the supply and distribution system with renewable energy is a core development trend. Through source-grid-load-storage integration, direct green power connection, energy-efficient equipment, and clean backup power, system losses and carbon emissions are continuously reduced, building a zero-carbon supply and distribution system. Future work will continuously optimize the deep coordination of new energy with AIDCs, improve green power consumption and coordinated control,

build a full-chain zero-carbon supply and distribution system, and drive the comprehensive transformation of GW-scale AIDCs into zero-carbon computing facilities.

1.3 GW-Scale AIDC Cooling System

The AIDC cooling system is key infrastructure for ensuring stable IT equipment operation, controlling energy costs, and enhancing core competitiveness. Its core goal is to efficiently transfer the large amount of heat generated during server operation to the outside, keeping server core components (CPU, GPU, etc.) within a safe and stable operating range while achieving optimal energy efficiency and controllable cost. With the explosion of AI compute demand, per-rack power density has soared from the traditional 8-15 kW to 30-60 kW; with the rollout of architectures such as NVIDIA GB300, per-rack heat dissipation has exceeded 150 kW, and full liquid-cooled architectures such as Rubin Ultra push per-rack power to 600 kW. AIDC energy consumption is gradually moving from megawatt (MW) scale toward gigawatt (GW) scale. Against this backdrop, cooling system architectures for GW-scale AIDCs have iterated from traditional air cooling toward liquid-cooling-led, multi-technology fusion.

1.3.1 Cooling System Architecture and Core Equipment

The mainstream heat dissipation architectures in the AIDC industry fall into three technical paths: a) Forced air cooling: fans dissipate heat from servers, suiting 12-50 kW per-rack power density; advantages are low cost and simple O&M, disadvantages are high noise and high energy consumption. b) Cold-plate liquid cooling: the core is a cold plate attached to heat-generating components such as CPUs and GPUs; circulating coolant directly absorbs the heat, exchanges it through a CDU, and finally conducts it to the outside atmosphere via cold-source equipment. This method has good compatibility, balancing heat dissipation efficiency and O&M difficulty, and is currently the preferred solution for high-power-density data centers. c) Immersion liquid cooling: servers are fully immersed in dielectric coolant; heat dissipation efficiency is highest, there are no local hot spots, and noise is low. As domestic coolant technology matures and system processes, supply chains, and commercial applications improve, the economics of single-phase immersion cooling are gradually emerging. Commercial deployment of immersion cooling is advancing rapidly and is on the eve of scaled deployment. Comparing these architectures, cold-plate liquid cooling solves the pain points of insufficient heat dissipation efficiency and high energy consumption of traditional air cooling, while immersion cooling has not yet reached scaled deployment and its industrialization has not yet formed. Based on the mainstream position and broad application value of cold-plate liquid cooling, the following sections fully dissect its system architecture, operating principles, core equipment, and engineering validation to provide a reference for practical application.

1.3.1.1 System Architecture

1. Primary-side system architecture. The primary-side liquid cooling system architecture should be selected according to the project's local meteorological conditions (including dry-bulb and wet-bulb temperatures), water resource availability, and the inlet temperature of the liquid-cooled servers. See the table below for the selection method.

Table 1.1 Primary-side liquid cooling system architecture

Climate conditions	Cold season present: low ambient temperatures enable year-round free cooling		Warm season present: high ambient temperatures require supplemental mechanical cooling	
Water availability	Water-rich areas	Water-scarce areas	Water-rich areas	Water-scarce areas
System architecture	Open cooling tower + plate heat exchanger, or closed cooling tower (Figure 1.5)	Dry cooler (Figure 1.6), air-cooled condenser (Figure 1.7)	Chiller + cooling tower + plate heat exchanger (Figure 1.8)	Air-cooled chiller (Figure 1.9)
Economics comparison				
Cost	Low	Relatively low	Relatively high	High
Energy consumption	Low	Relatively low	Relatively high	High
O&M difficulty	Relatively low	Low	High	Relatively high



Figure 1.5 Cold-source solution: open cooling tower + plate heat exchanger or closed cooling tower

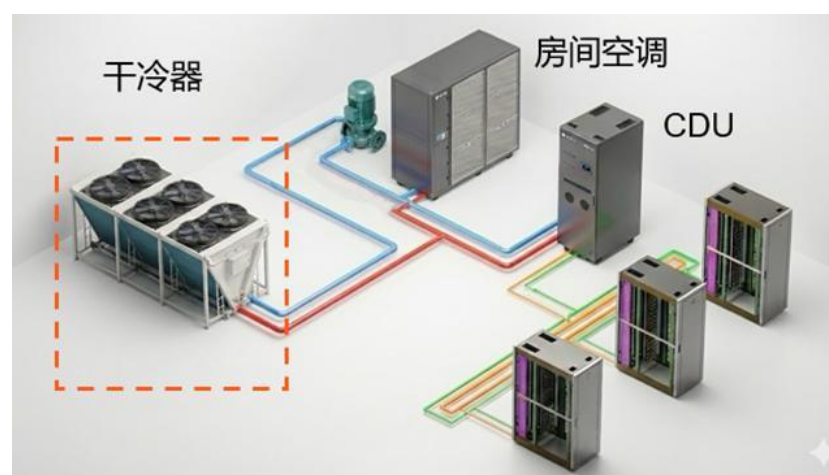


Figure 1.6 Cold-source solution: dry cooler



Figure 1.7 Cold-source solution: air-cooled condenser



Figure 1.8 Cold-source solution: water-cooled chiller

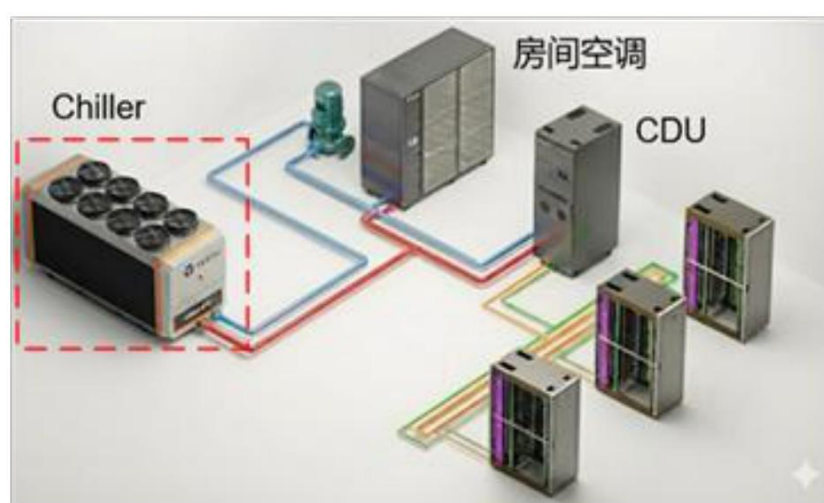


Figure 1.9 Cold-source solution: air-cooled chiller

2. Secondary-side system architecture 1) Centralized liquid cooling system architecture. The centralized liquid cooling system is characterized by centralized CDU management and unified pipe network distribution. Mainstream designs commonly use 2N or N+X redundancy, adaptable to the reliability level requirements of different compute scenarios and providing redundancy assurance for the stable operation of high-density compute equipment.

2N architecture (figure below): a fully redundant, dual-bus architecture whose core is the deployment of two completely independent, equivalent CDUs each with 100% load capacity. The two systems are fully independent in structure, electricals, and control, yet back each other up. Under normal conditions, the two CDUs operate in parallel, jointly carrying 100% of the system load. If either CDU fails, the other switches automatically without interruption per preset logic to carry the full system load, ensuring stable operation of compute equipment. The 2N architecture is convenient to maintain: one CDU can be powered down for maintenance without interrupting the other's normal operation.

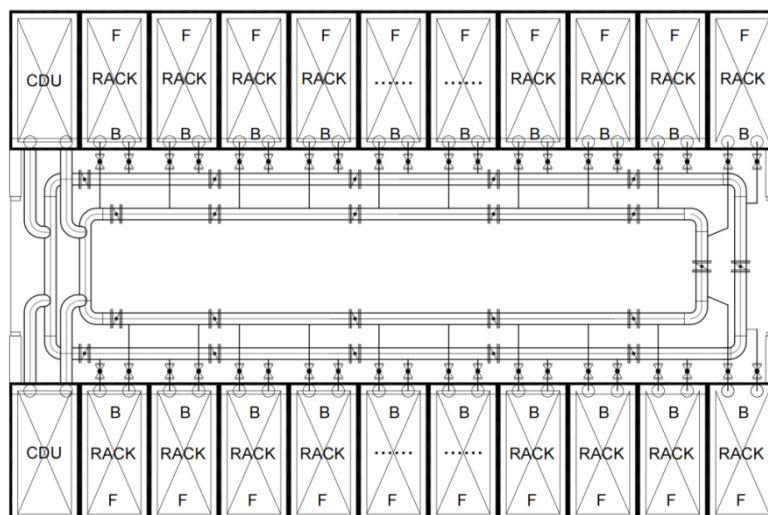


Figure 1.10 Centralized liquid cooling system 2N architecture

N+X architecture (figure below): a modular redundancy architecture that achieves load balancing and automatic failover through intelligent group control. N is the minimum number of CDUs needed for the total system load; X is the number of redundant standby CDUs. All CDUs connect in parallel to the same supply/return pipe network and jointly participate in system cooling distribution and intelligent control. Under normal conditions, N CDUs run at full load and X CDUs are in hot standby; when a CDU fails, a standby CDU cuts in quickly and automatically to keep compute equipment running safely. Compared with the 2N architecture, N+X offers better cost and is friendly to linear expansion.

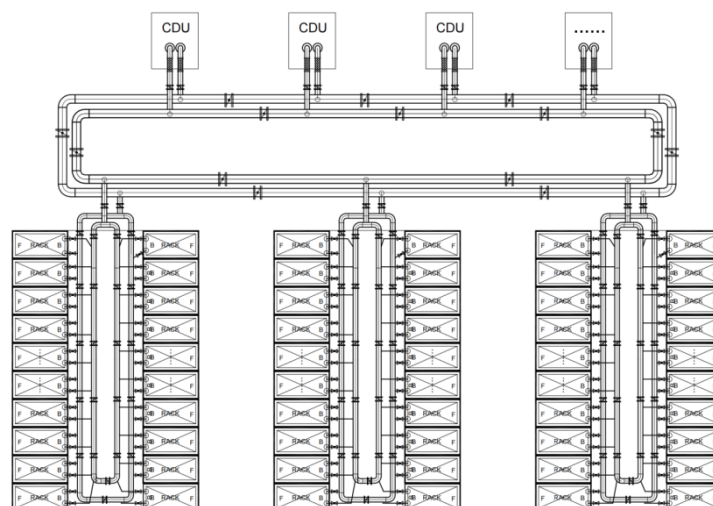


Figure 1.11 Centralized liquid cooling system N+X architecture

The 2N architecture focuses on zero-interruption requirements of core compute scenarios, providing the highest level of reliability assurance; the N+X architecture balances reliability and cost-effectiveness. Both architectures achieve standardized, modular, highly reliable design, providing solid technical support for green, low-carbon, efficient, and stable operation of high-density AIDCs.

2) Distributed liquid cooling system architecture. The distributed liquid cooling system is characterized by CDUs placed inside racks, achieving rack-level precision heat dissipation through independent closed

loops. It features flexible deployment, on-demand expansion, shipment with full-rack server deployment, and good isolation — a single rack failure does not affect the whole. It suits retrofit equipment rooms of high-density traditional data centers and edge computing nodes with demanding space requirements. For more details, see Section 2.1, AI Supernodes.

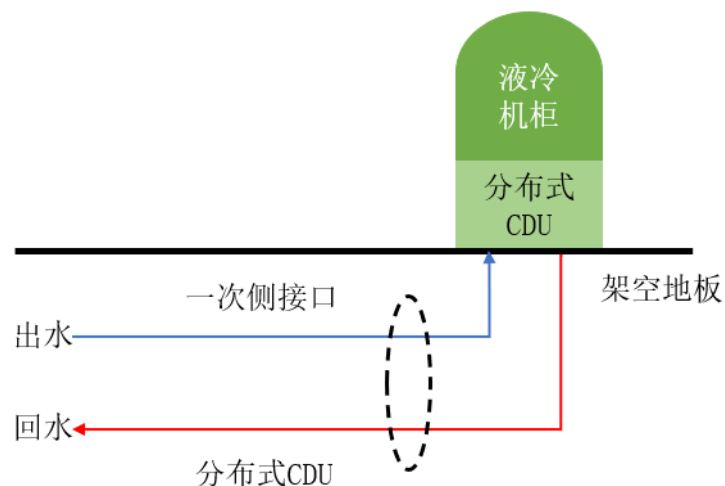


Figure 1.12 Distributed liquid cooling system

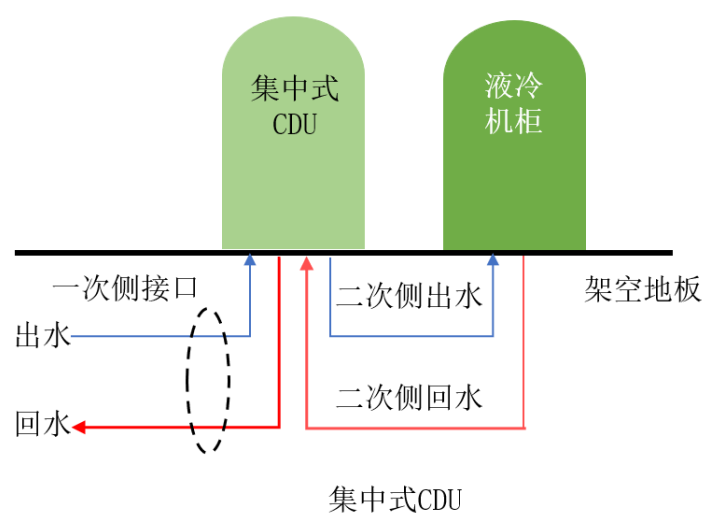


Figure 1.13 Centralized liquid cooling system

A typical liquid-cooled full-rack layout is shown below, including compute nodes, switch nodes, power shelves, blind-mate manifolds, cable cartridges, and more.



Figure 1.14 Full-rack architecture

1.3.1.2 Control System

1. Composition and architecture of the cold-source automatic control system. The cold-source automatic control system is mainly responsible for monitoring primary-side circulation equipment and cooling towers. It consists of a cold-source automatic controller, pump flow meters, pressure sensors, temperature sensors, motorized valves, cooling tower sensors, CDU monitoring, other sensors, and a display. The system architecture is as follows:

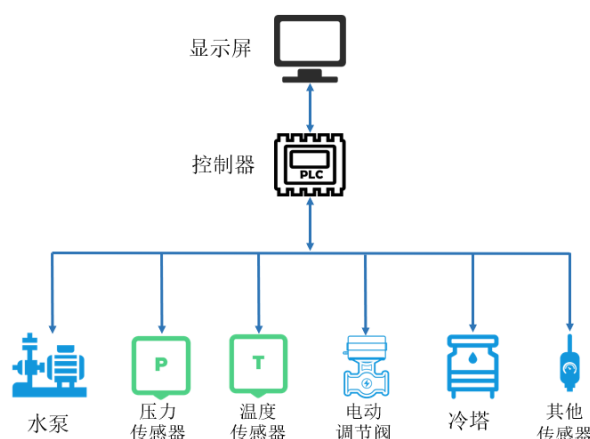


Figure 1.15 Cold-source automatic control system

Two common solutions for cold-source automatic control systems are PLC and DDC. PLCs are stable and reliable, flexible in function development (e.g., hot standby, timed rotation), available in many models, easy to expand, and industrial-grade for complex environments. DDCs are better suited to building HVAC control and are easier to develop and commission.

2. Secondary-side system control architecture. 1) Control architecture of 2N systems. The 2N system architecture has 2 CDUs in the same ring network, with N ring networks in total, as shown below:

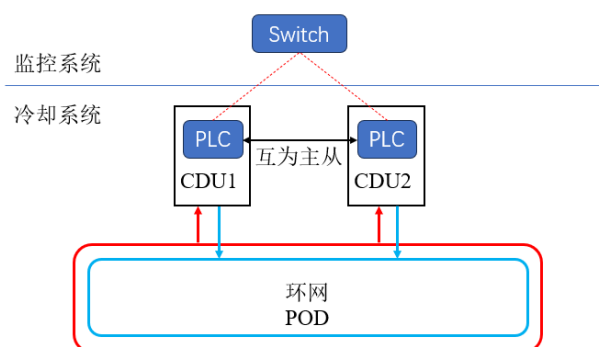


Figure 1.16 2N control architecture

Each CDU in a 2N system typically has an independent controller (such as a PLC or embedded microcontroller). CDUs are mutual master-slave peers. Unit operating modes include group control mode and standalone mode: in group control mode, a CDU receives command parameters from the group control master; a CDU can also leave group control for standalone mode, with its own PLC executing the control logic. Each CDU's controller communicates with the supervisory monitoring system and uploads operating data.

2) Control architecture of N+X systems. The N+X system architecture suits systems with 3 or more CDUs in the same ring network: N CDUs run and X CDUs stand by in cold reserve, as shown below:

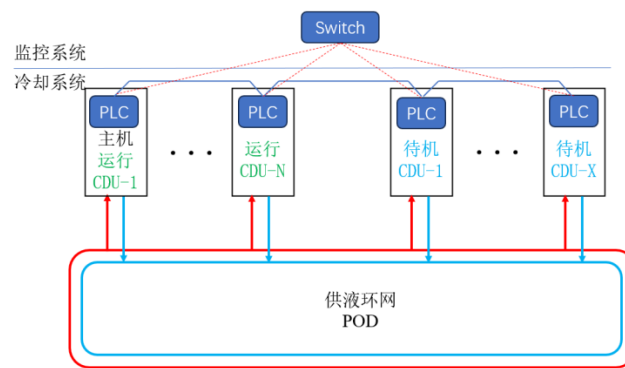


Figure 1.17 N+X control architecture

In an N+X system, one CDU must serve as the group control master. When the system is in group control mode, CDUs execute the master's commands; a CDU can also leave group control for standalone mode, with its own PLC executing the control logic.

Each CDU in an N+X system has an independent PLC controller. The controllers are peers with no separate one-to-many master controller; any CDU controller can serve as the system group control master. CDU controllers communicate over the control network to synchronize key data, enabling fault cutover and manual master/standby assignment. Any two PLCs in the automation network can communicate with each other. Each CDU's controller communicates with the monitoring system and uploads operating data.

Master assignment and switching in N+X systems: the master can be designated manually, and there is exactly one master in the system at any time. In special cases — such as the master PLC losing communication with all other PLCs, or the master PLC losing power or failing — the remaining online CDUs can automatically elect a master to perform group control.

3) Group control logic. The group control master uniformly schedules start/stop, fault cutover, unit rotation, and adjustment of secondary-side pumps and secondary-side two-way valves inside CDU units across all CDU units in the system.

Cold standby/hot standby: in group-controlled cold standby mode, standby CDUs are stopped and run in turn at fixed rotation intervals to balance unit operating hours and extend service life; in group-controlled hot standby mode, all CDUs in the system run together and share the heat exchange load equally.

Data synchronization: the control modes and parameters of CDUs in the same system remain consistent; modifying parameters on any CDU synchronizes the data to the entire system.

Fault cutover: CDU fault levels are managed hierarchically. When a fault affecting unit operation occurs in the system, the standby unit must be started and the faulty unit stopped, ensuring stable flow during the switchover.

Reliability design: important sensors are backed up inside the system to prevent sensor failures from causing system failures. In addition, to ensure a CDU does not lose heat exchange capacity due to controller failure, the CDU pump keeps running in its pre-failure state after the controller fails.

1.3.1.3 Heat Exchange Equipment

1. Key points and trends of centralized CDUs. Centralized CDUs are divided into in-row and non-in-row types according to their physical layout in the AIDC equipment hall. In-row CDUs are arranged alongside server racks, generally matching their height and depth, and provide heat dissipation for the server racks in that row. Non-in-row CDUs are placed in specific areas of the equipment hall outside server rack rows; their dimensions are generally not strictly specified, and multiple CDUs typically jointly serve all server racks in the hall.



Figure 1.18 Centralized CDU

1) Design points. a) Heat exchanger: when the cold plate is copper, use a copper-brazed heat exchanger; when the cold plate is aluminum, use an all-stainless-steel brazed or detachable heat exchanger; b) Pumps: pump mechanical seals should be silicon carbide vs silicon carbide or silicon carbide vs tungsten carbide, not graphite mechanical seals; wetted materials should be 304 stainless steel or better. To ensure system reliability, the pump group needs 2N redundant backup and online replacement capability. c) Piping flow velocity: internal piping flow velocity on the CDU primary and secondary sides should be ≤ 3 m/s. d) Filtration requirements: primary-side filtration precision should be ≥ 50 mesh, secondary-side ≥ 270 mesh; primary- and secondary-side filters should support online cleaning. e) Automatic venting: automatic air vents should be installed at the CDU's highest point and local high points. f) Pressure relief protection: the CDU must have an automatic pressure relief valve that automatically relieves pressure when system pressure exceeds 5 bar. g) Anti-condensation: the CDU must have anti-condensation control to avoid condensation risk in the secondary-side system. h) Wetted materials: all wetted materials must meet compatibility requirements. Piping should be 304 or better stainless steel, solder should be the corresponding stainless steel solder, and gaskets and hoses should use EPDM, PTFE, or PVDF. Bronze, carbon steel, brass with zinc content $> 15\%$, high-lead brass, aluminum and aluminum alloys, PTFE tape, and threadlocker are prohibited. 2) Development trends. a) High heat flux: the power of a single CDU keeps increasing; single CDUs with 2 MW cooling capacity have appeared, and 3 MW or

higher-capacity CDUs will emerge to meet the high-density needs of AIDCs and supercomputing clusters. Multi-module parallel connection supports linear expansion, suiting phased equipment hall construction.

b) Full-chain intelligence: hall-level AI scheduling can dynamically optimize supply temperature and flow based on outdoor temperature/humidity and IT load, saving 20%-30% energy; CDU + piping + rack full-chain simulation provides early warning of blockage, leakage, and hydraulic imbalance; and trends in pump life, heat exchanger scaling, and pipe corrosion can be predicted. c) Higher temperatures + waste heat recovery: raising supply temperature to 40-55°C enables year-round free cooling, reducing PUE to 1.05-1.1; using coolant waste heat for campus heating, domestic hot water, and dehumidification brings comprehensive energy utilization to $\geq 80\%$. d) Low energy consumption and long life: high-efficiency magnetic levitation pumps + variable-frequency control cut transport energy by 40%, and corrosion-resistant new materials (titanium alloys, specialty plastics) extend service life to over 15 years.

2. Key points and trends of the secondary-side pipe network. Secondary-side pipe networks are usually prefabricated in factories: pipes are grouped, and most pipe module welding, cleaning, and tightness testing are completed in the factory, then shipped under nitrogen pressure to the project site for assembly. Only a small amount of piping is installed on site, commonly using "hot-work-free" non-welded connections such as clamps and flanges.

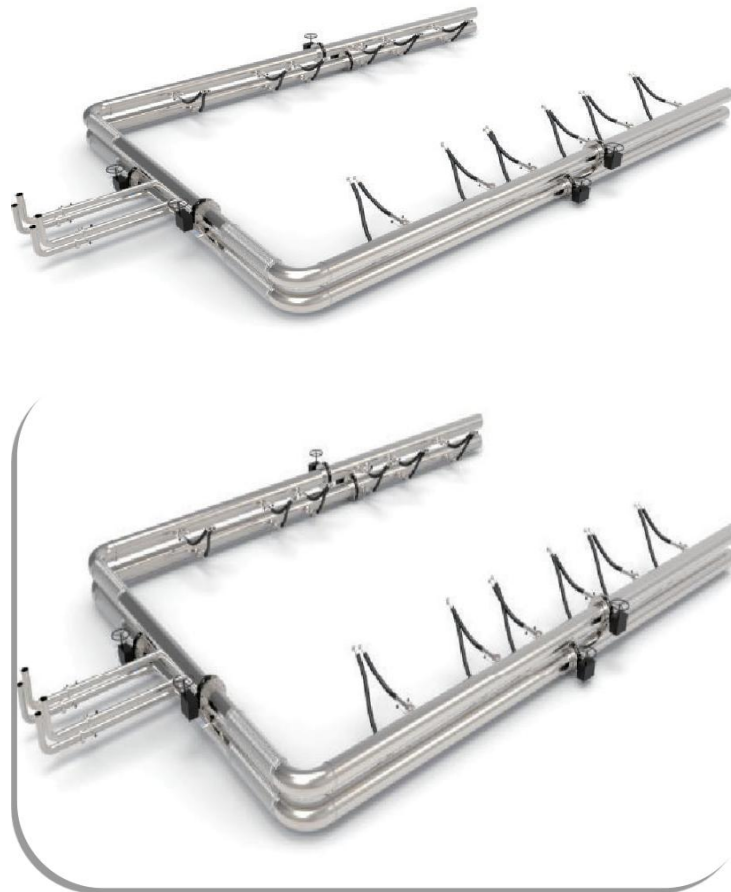


Figure 1.19 Secondary-side engineering pipe network

1) Design points. a) Material selection: liquid-cooling CDUs typically operate between 15°C and 45°C, and the working medium is mostly pure water or inhibited ethylene glycol. Piping materials must have corrosion resistance, mechanical strength, workability, and chemical compatibility with the cooling medium. In practice, 304L and 316L are the most widely used materials; the main differences are nickel (Ni) content and molybdenum (Mo) addition. Considering overall performance and cost, 304L is preferred as the secondary-side piping material. b) Flexible connections: piping design and installation must account for thermal expansion and contraction, allowing free movement without stressing connected pipes or equipment, and must pass pressure tightness tests under both high- and low-temperature conditions. The tightness test pressure must exceed the design pressure (e.g., 1.5 times design pressure). c) Maintainability: the pipe network should be designed as a detachable, flushable structure with inspection ports reserved at key nodes. d) Cleanliness requirements: before commissioning, the entire system must be circulated and flushed, and samples tested for coolant dielectric constant, volume resistivity, water content, and other indicators, with the liquid free of impurities and the surface free of oil. e) Drip pans: drip pans must be installed at the bottom of the secondary-side pipe network, usually with full welding around the perimeter. Matching different pipe sections, drip pans are also prefabricated in sections and designed with a slope for directional drainage. 2) Development trends. a) As AIDC liquid cooling systems demand higher reliability and stability, liquid cooling pipe networks are evolving from "passive heat dissipation channels" to "intelligent thermal management infrastructure," for example: b) Distributed

sensor networks: deploy pressure, temperature, flow, and leak detection sensors on each branch for real-time sampling and monitoring, promptly discovering faults and risks in system operation. c) Digital twin integration: feed sensor data from the liquid cooling pipe network into the AIDC digital twin platform for fault prediction and heat flow simulation optimization.

1.3.2 Hall-Level Cooling System Engineering Validation

1. Tightness test. Before a hall-level liquid cooling system enters operation, an overall secondary-side tightness test is performed to ensure no seepage or leakage in the secondary-side system, typically using gas or liquid detection.
2. Supply characteristics test. This assesses whether the relationship between secondary-side flow resistance and flow rate meets technical specifications. Simulation tests are run based on the liquid cooling system's flow rate and water pressure differential, and the upper and lower flow limits of the overall secondary-side system are also mapped, providing data support for later operation.
3. CDU fault cutover and switching time test. Simulate the time for unit switching to stabilize when a CDU raises a fault alarm, and record flow fluctuation values during switching, providing data support for later operation.
4. Heat exchange capacity test. Simulate and test whether the CDU's heat exchange capacity and the cooling tower's heat dissipation capacity reach design values.

1.3.3 Future Outlook

As AI large models iterate at an astonishing pace, compute demand climbs exponentially, and server power consumption and heat dissipation needs rise accordingly. Squeezed by both compute density and energy policy, liquid cooling has gradually changed from an "option" to a "necessity," becoming a standard rigid requirement. Looking ahead, GW-scale AIDC cooling systems will evolve around high-density heat dissipation, dual-track cold-plate + immersion development, extreme PUE, and full-stack intelligence.

1. High-density heat dissipation: per-chip power has jumped from 200 W to a future 3,000 W (NVIDIA GB300/Rubin), and per-rack power has jumped accordingly to over 150 kW, even to the megawatt level. Demand for heat dissipation power density in liquid cooling systems keeps rising, as do heat flux requirements for heat exchange equipment directly involved in server cooling, such as CDUs and cold plates. Rack-mount CDU heat flux of 4U-150 kW has become mainstream, centralized CDU heat exchange power has exceeded the 2 MW level, and liquid-cooling cold plates achieve heat flux above 150 W/cm² through structures such as manifold microchannels. For GW-scale AIDCs, demand for various new technologies, new materials, and new architectures is emerging.

a) New technologies: two-phase liquid cooling systems based on the physical characteristics of working-medium phase change (boiling/condensation) latent heat of vaporization further improve heat dissipation efficiency to 1.3-1.5 times that of current single-phase liquid cooling. b) New materials: cold plates combining ultra-high-thermal-conductivity materials such as diamond with manifold microchannel and other technologies achieve heat flux exceeding 1000 W/cm^2 , but due to processing and other issues this remains in the pre-research stage.

2. Dual-track cold-plate + immersion development: cold-plate liquid cooling currently dominates thanks to mature technology, stability, reliability, and controllable cost. But as heat flux rises, immersion liquid cooling is gradually developing. For some time, the two mainstream architectures — cold-plate and immersion liquid cooling — will coexist. As immersion technology continues to iterate and optimize, standards systems gradually improve, equipment interoperability becomes fully compatible, hardware architectures are deeply decoupled, and commercial scale keeps expanding, immersion liquid cooling will become an indispensable mainstream cooling solution for high-density AIDCs, supporting the scaled deployment of GW-scale AIDCs, reducing the cost per unit of compute per token, and helping realize the goal of inclusive computing power.

3. Full-stack intelligence: to improve the stability and reliability of liquid cooling systems and enhance system PUE, intelligent and refined management of future liquid cooling systems is imperative. From full-parameter sensing of CDUs to building CDU health management systems that predict CDU faults and maintenance needs, to introducing AI group control systems that adaptively regulate liquid cooling temperatures with load prediction and fault self-healing, liquid cooling systems will gain a smarter "brain," achieving true "predictive heat dissipation."

4. Extreme PUE: combining the liquid cooling system with waste heat recovery — using coolant waste heat for campus heating, domestic hot water, dehumidification, and other scenarios — achieves comprehensive energy utilization and reduces the overall PUE of the entire AIDC campus.

Chapter 2 IT Facility Layer

2.1 AI Supernodes

This chapter focuses on the AI supernodes that carry the main compute load in an AIDC. General-purpose computing (management, login, data preprocessing, etc.) serves as necessary support; its architecture can reuse mature IDC practices and is not separately elaborated in this version of the report.

An AI supernode is essentially an integrated super-server that vertically scales AI acceleration modules such as GPUs over high-speed buses such as UALink, Ethernet, and CLink. Leveraging cutting-edge technologies such as architecture decoupling and liquid cooling, it achieves extreme high-density integration of AI acceleration modules. By shortening high-speed signal transmission paths between modules, it effectively increases interconnect bandwidth and reduces communication latency, fully suiting application scenarios such as MoE, PD disaggregation, agents, long-sequence input/output, and multi-turn interaction.

The development of AI large models has driven continuous, substantial increases in the performance and power consumption of AI acceleration modules, while the power supply and cooling infrastructure conditions of existing data centers vary by region. To adapt to the carrying capacity of existing equipment hall infrastructure, the distributed supernode form has emerged, complementing the full-rack supernode and jointly ensuring efficient, stable operation of upper-layer AI services.

The full-rack supernode uses the rack as its unit and suits data centers with high power density and liquid cooling. The distributed supernode uses the chassis as its unit and can be mounted directly on already-deployed racks, suiting data center environments with constrained power and cooling conditions or those being built out incrementally.

2.1.1 Full-Rack Supernodes

The new full-rack supernode architecture achieves high-density integration, high-speed interconnect, liquid cooling, and centralized power supply. GPU chips are interconnected via high-speed buses and switch chips, enabling GPU interconnect at scales of 32, 64, or 128 GPUs within a single rack or across multiple racks, breaking through technical challenges in in-rack computing, memory, storage, power supply, and heat dissipation, and solving the problems of resource pooling and integrated design.

2.1.1.1 Compute Nodes

The compute node is the core carrier unit for CPUs, GPUs, memory, and in-rack high-speed Scale-Up interconnect interfaces, and is the core compute carrier of a full-rack supernode. The extreme pursuit of performance and compute density brings the dual challenges of high-power supply and heat dissipation.

In today's mainstream industry supernode compute node designs, centralized power supply and full liquid cold-plate cooling are widely adopted.

Take the supernode of Baidu, a leading Chinese internet company, as an example: its compute node integrates three functional modules — a CPU compute board, a PCIe Switch board, and a GPU board — in a 1:2:4 resource ratio, with four SmartNICs for expansion. The compute node is 1U high, 21 inches wide and 1.2 m deep; NICs, NVMe SSDs, M.2 and other components are front-mounted, while Scale-Up high-density connectors are rear-mounted.



Figure 2.1 Baidu full-rack supernode compute node design example

Different AI application scenarios have significantly different CPU-to-GPU resource ratio requirements, leading to very different supernode architecture designs.

In AI training scenarios, CPUs mainly handle data loading, task distribution, and GPU compute coordination scheduling, emphasizing low-latency, high-bandwidth interconnect rather than simply stacking core counts for greater compute power. The NVIDIA GB200 NVL72 uses 36 Grace CPUs with 72 cores plus 72 Blackwell GPUs, a 1:2 CPU-to-GPU ratio, and the Vera Rubin supernode architecture continues the 1:2 CPU-to-GPU ratio (with CPU cores increased to 88), representing the current mainstream supernode architecture for training scenarios. In AI inference scenarios — especially complex applications such as AI agents — as task complexity rises and multi-model scheduling increases, CPU load grows significantly in task orchestration, memory management, scheduling, and data transfer to GPU accelerators. Tightly coupled architectures with fixed CPU-to-GPU ratios are prone to resource mismatch and wasted compute, whereas disaggregated, pooled CPU-GPU architectures adapt to a wider range of applications and significantly improve resource utilization, becoming an important direction for inference clusters — Alibaba Cloud's Panjiu AL128 supernode server is a typical example. At GTC 2026, NVIDIA introduced the Groq 3 LPX LPU full-rack solution, further enriching the types of compute resources for inference supernodes with more flexible combination ratios.

2.1.1.2 Racks

The rack is the physical carrier and structural core of a full-rack supernode system. It hosts compute nodes, power shelves, distributed CDUs, network nodes, the Scale-Up high-speed network, and the Scale-Out interconnect network, and provides reliable fixing and structural support for system power busbars and liquid cooling pipelines, enabling precise interconnection and coordination among in-rack functional modules — the core physical framework for compute aggregation in a supernode. The significant improvements of full-rack supernode racks in structural strength, machining precision, and coupling adaptation with internal equipment are the core differences from general ICT equipment racks.

Rack specifications for full-rack supernodes must match the overall architecture design of the supernode system, and several differentiated schemes exist in the industry. Among them, the purpose-optimized rack based on OCP Open Rack V3.0 (ORV3), promoted by NVIDIA, has a relatively mature industrial ecosystem. Wide-body racks with higher scalability have also been introduced; OCP Open Rack Wide (ORW) is a widely used industry example, demonstrated by Meta, AMD, and others at the 2025 OCP Global Summit. However, Meta also noted in its technical sharing that larger racks mean greater design and manufacturing difficulty, higher shipping costs, and greater O&M challenges, so the future direction may be multiple decoupled low-density ORV3 racks achieving higher scalability through optical interconnect.

The supernode rack of Baidu, a leading Chinese internet company, inherits the design philosophy of the Scorpio full-rack server architecture. The rack is 46U high, with overall dimensions of 2300 mm (± 2) $\times$ 600 mm (0, -3) $\times$ 1200 mm (± 2) and a unit RU height of 46.5 mm. The rack's rated load capacity is no less than 1600 kg, and it supports blind-mate interfaces for water, power, and network, enabling blind-mate node servicing. The liquid cooling manifold uses top/bottom supply-return design and supports integrated full-rack delivery, significantly improving delivery and O&M efficiency.

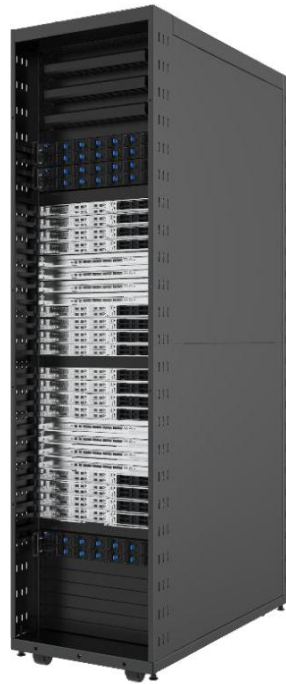


Figure 2.2 Baidu full-rack supernode rack design example

2.1.1.3 Full-Rack Power Supply

For power access, North American data centers traditionally use 480Y/277 VAC and 208 VAC as the main input systems, while Chinese data centers generally use 380Y/220 VAC, 240 VDC, and 336 VDC. As AI supernode power density rises rapidly, higher-voltage DC input schemes such as ± 400 VDC and 800 VDC will mature and proliferate.

1. Power shelf supply. The power shelf is the power conversion hub of a full-rack supernode system. It uses a pooled architecture to centrally deploy in-rack power supply units (PSUs) and is equipped with a power management controller (PMC) that monitors the operating status and health of PSUs and other supply components in real time. The power shelf receives input power through a PDU, which the PSUs step down and convert to 54 VDC or 48 VDC before outputting to the rack bus, supplying all in-rack equipment uniformly. A purpose-optimized power shelf can also provide several standard power interfaces for other rated-voltage input devices.

As per-rack power rises rapidly, improving power density and conversion efficiency has become the core design goal of PSUs. Applying new materials and technologies — third-generation semiconductors (such as GaN devices), high-frequency low-loss magnetic components, and advanced power conversion topologies (such as the flying-capacitor four-level FC4L topology) — effectively shrinks device size, reduces hysteresis and eddy current losses, lowers device voltage stress, and significantly improves the device figure of merit, ultimately achieving higher power density and conversion efficiency.

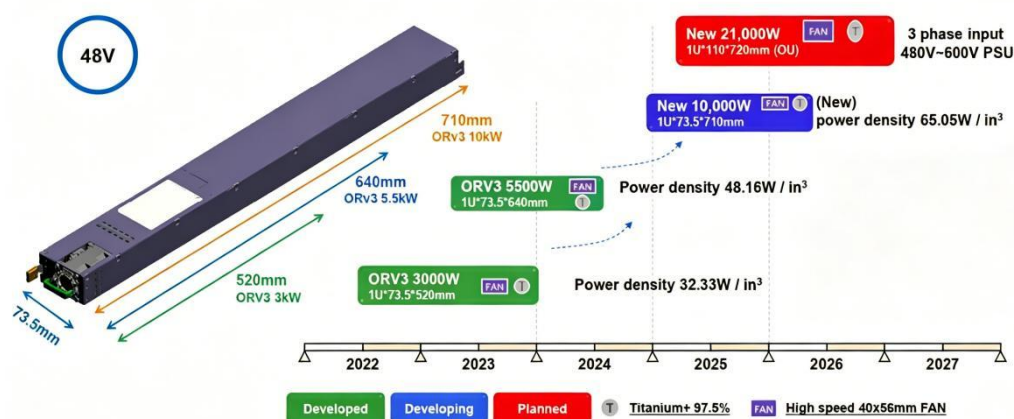


Figure 2.3 PSU power and size development roadmap

An in-rack integrated power shelf occupies deployment space for compute and switch nodes. When per-rack power exceeds 200 kW, continuing with the built-in power shelf architecture significantly reduces compute deployment density and can even affect the operating efficiency of the entire compute cluster. As per-rack power rises further, co-locating power units with compute and network equipment in the same rack becomes unsustainable; power must be separated from the service racks, and the sidecar independent power rack solution will gradually become the mainstream AIDC solution.

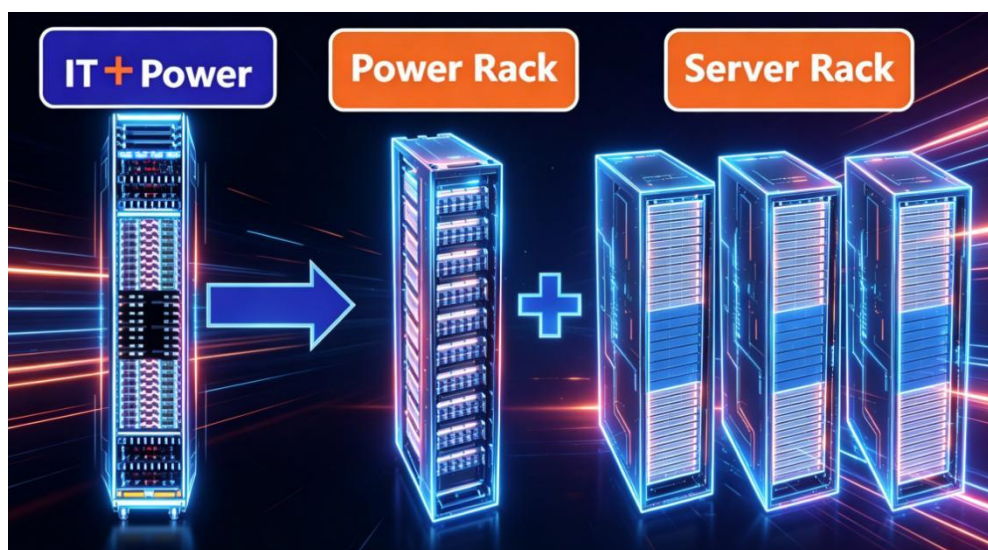


Figure 2.4 Full-rack supernode power architecture shifting from co-rack to sidecar independent power racks

2. Compute node supply. In current industry practice, compute nodes and network nodes typically draw 54 VDC or 48 VDC input directly from the rack busbar via power clips to power GPU modules, while CPUs and other node components require further step-down conversion to 12 VDC. With the adoption of ± 400 VDC and 800 VDC DC input schemes, busbar voltage will also upgrade to ± 400 VDC and 800 VDC, requiring additional DC conversion modules inside compute and network nodes to step down and output 54 VDC, 48 VDC, 12 VDC, and other voltages for GPU modules and other node components. As the power consumption of chips such as GPUs and CPUs keeps climbing (a single GPU now consumes up to 2 kW),

power path design challenges grow, and more efficient new supply architectures such as vertical power delivery have emerged.

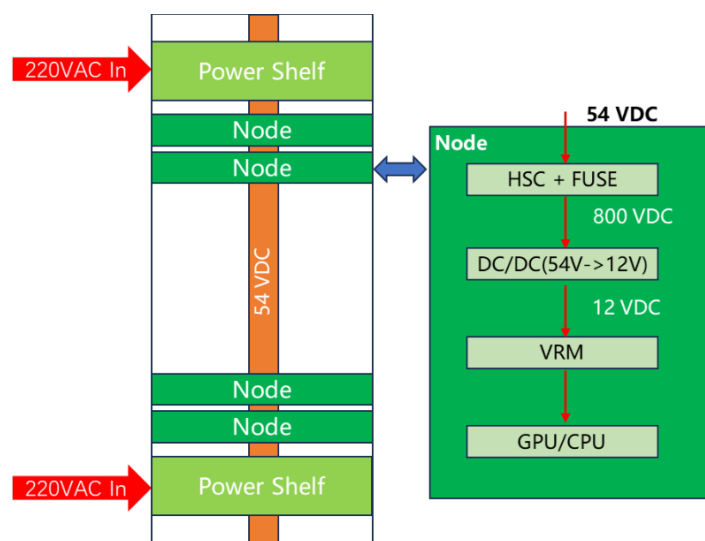


Figure 2.5 Current 220 VAC full-rack supply topology

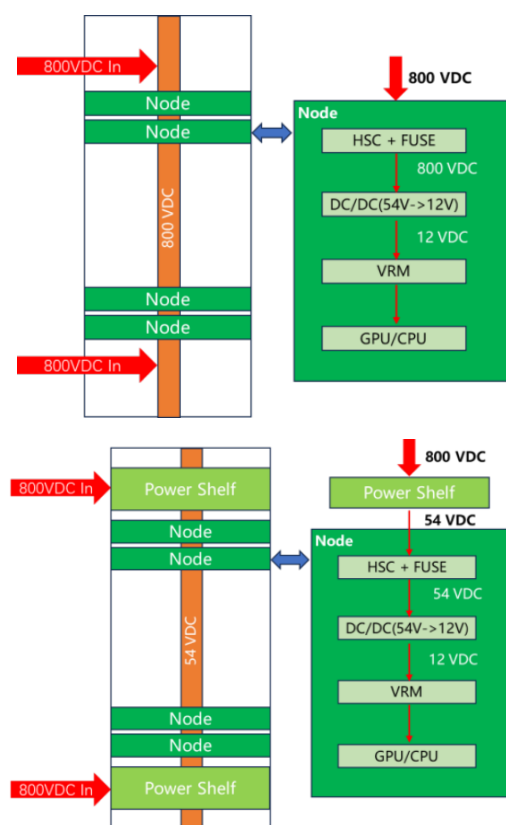


Figure 2.6 Future 800 VDC supply topology

2.1.1.4 Full-Rack Cooling

At the full-rack system level, rack-level cooling is moving from air cooling to hybrid air-liquid cooling and ultimately to full liquid cooling. As compute, network, and storage nodes increase in capability and power consumption, maintaining outstanding performance and system stability is a technical challenge

supernodes must overcome, and liquid cooling provides the better heat dissipation solution for high-power-density scenarios.

To achieve extreme compute density, large cloud service providers (CSPs) and next-generation cloud data center users generally adopt full-rack supernode solutions with liquid cooling architectures. Small and medium-sized users with modest compute density requirements or data centers not yet equipped for liquid cooling can choose air cooling or hybrid air-liquid cooling.

Current liquid-cooled full-rack supernodes generally adopt the cold-plate approach. In-rack liquid cooling components mainly include the manifold, cold plates, quick disconnects, and the in-rack CDU.

The in-rack manifold is the piping system connecting the CDU to node cold plates. According to how quick disconnects connect the manifold to cold plates, manifolds are divided into manual-mate and blind-mate types; blind-mate manifolds usually have anti-splash components and drainage channels. For reliable assembly, the flatness of the manifold connector mounting surface must not exceed 0.1 mm. To ensure balanced cooling of in-rack nodes, manifold branch flow deviation must not exceed 10%.



Figure 2.7 Manifold and anti-splash components

Cold plates are located inside compute and switch nodes. Their efficient design and manufacturing processes are critical to the long-term reliable operation of high-power components such as GPUs, CPUs,

and ASICs, and to system energy savings. At the design level: optimize skived-fin structures to increase heat dissipation area and enhance fluid turbulence; adopt jet chamber designs for direct impingement cooling of hot spots, raising local heat dissipation intensity; and introduce two-phase cold-plate technology with high-efficiency evaporative heat transfer structures, using liquid phase change to absorb large amounts of latent heat and further strengthen heat dissipation. At the manufacturing and inspection level: perform post-weld machining to ensure flatness; conduct 100% ultrasonic non-destructive testing of welds to eliminate defects; and apply nickel plating or passivation for long-term protection.

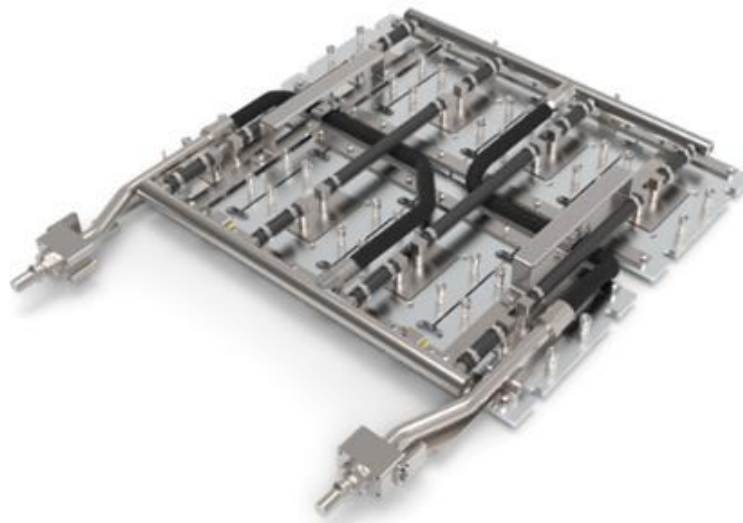


Figure 2.8 GPU cold plate assembly

By installation location, quick disconnects are divided into in-rack mounting (on the manifold side) and in-server mounting (on the cold-plate side). By mating operation, they are divided into manual-mate and blind-mate quick disconnects. Blind-mate quick disconnects are widely used in supernode systems. Reliable blind-mate connection requires strict radial, axial, and angular float tolerances: radial float of no less than ± 1.0 mm, axial float of no less than 1.0 mm, and — with a floating module — angular float of no less than $\pm 2.5^\circ$.



Figure 2.9 Liquid cooling quick disconnects

By overall CDU system layout, full-rack supernode liquid cooling solutions divide into distributed and centralized cooling architectures. Distributed cooling integrates the CDU with IT equipment in the same rack, better suiting small and medium deployment scales that do not pursue extreme compute density. Large CSPs and neo-cloud users generally adopt centralized cooling with independent CDU racks, which enables higher compute density, simplifies liquid cooling piping topology, and greatly improves O&M efficiency.



Figure 2.10 Full-rack supernode with distributed cooling

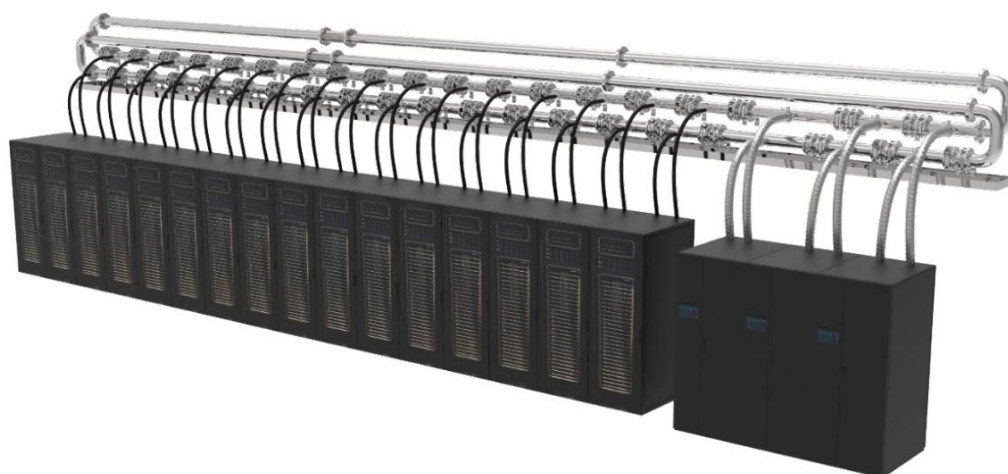


Figure 2.11 Full-rack supernode cluster with centralized cooling

2.1.1.5 Scale-Up Interconnect

Relying on dedicated Scale-Up interconnect protocols, a supernode can build a globally unified address space, abstracting the memory and VRAM of all AI acceleration chips in the system into a unified memory resource pool. Compute units can directly access other chips' memory resources with memory semantics, eliminating complex data copy operations and reducing end-to-end latency to sub-microsecond levels.

Supernode Scale-Up interconnect must focus on core requirements of high bandwidth, low latency, high reliability, and high density, and divides into two technical routes: copper and optical interconnect.

1. The copper interconnect route uses copper conductors as the transmission medium and currently dominates thanks to high reliability, low latency, a mature supply chain, and controllable cost. Cable backplane interconnect and board-to-board orthogonal interconnect are the two core interconnect architectures.

1) Cable tray/backplane interconnect. Through continuous optimization of connectors and copper cable technology, cable backplanes have greatly increased link loss margin, and interconnect length can extend to 2 m. Their blind-mate self-guiding floating mechanism completely eliminates the high-speed signal degradation caused by incomplete mating in traditional hard connections, while supporting flexible conversion between cable-to-board and cable-to-cable forms, offering extremely high design flexibility and reliability. This solution is especially suited to architectures where compute and network nodes are horizontally laid out and vertically stacked, and is currently the most widely used in-rack interconnect solution for full-rack supernodes.

To ensure long-term reliable cable backplane connection, requirements for the floating mechanism and mating force must be quantified. Taking Baidu's supernode as an example, its cable backplane quantitative indicators are as follows:

a) Panel-side X/Y axis float: within the specified bearing pressure range, achieve ± 3.0 mm self-correcting guidance with no sticking, locking, or spring-back failure throughout the mechanism's travel;

b) Panel-side Z-axis travel float (spring compression direction): within the specified bearing pressure range, achieve a maximum 3.5 mm spring compression under node over-pressure conditions, with no sticking, locking, or spring-back failure throughout;

c) Spring performance: springs are the core components enabling float in the X/Y/Z axes and must stably meet elastic performance requirements within different designed bearing pressure ranges.

2) Board-to-board orthogonal interconnect: orthogonal interconnect breaks the structural limits of horizontal layout and vertical stacking of in-rack nodes, enabling front-to-back vertical mating of compute and network nodes (both in vertical orientation), further increasing the compute integration density of the supernode system.

a) Orthogonal midplane: suited to small-scale, short-distance interconnect scenarios, typically applied inside a GPU tray and between adjacent GPU trays and switch trays (interconnect distances usually 0.5-0.7 m). With advantages of simple architecture, high integration, low cost, and easy maintenance, it is widely used in distributed supernode systems.

b) Orthogonal direct: relying on low-loss PCB materials, high-speed connector material upgrades and process iterations, plus architecture optimization, it completely eliminates midplane signal loss and greatly extends high-speed signal transmission distance. Its high-speed channels come in PCB and copper cable forms: copper cable interconnect reaches 1-1.3 m at 112G PAM4 and 0.4-0.8 m at 224G PAM4; PCB

trace interconnect distance is slightly shorter than copper cable. Orthogonal direct interconnect is especially suited to ultra-high-integration-density full-rack supernode architectures and is an important technical route for next-generation Scale-Up high-speed interconnect.

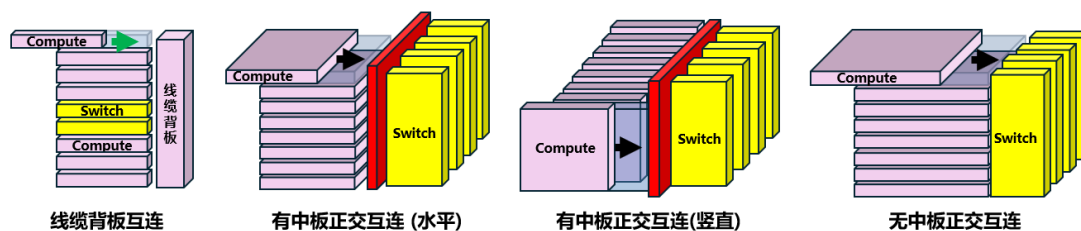


Figure 2.12 Multiple board-to-board orthogonal interconnect solutions

As supernode signal rates keep climbing, solving the interconnect and transmission bottleneck for signals above 112 Gbps within the chassis makes NPC and CPC the optimal solutions:

NPC (Near Package Copper): targeting rack-level long-distance interconnect, centered on loss optimization and distance extension, supporting 112G/224G rates with transmission distances of 1.5-2 m;

CPC (Co-Packaged Copper): targeting package-level ultra-high-density interconnect, centered on rate improvement and power reduction, bringing 224G/448G signals inside the package with transmission distances of 1 m.

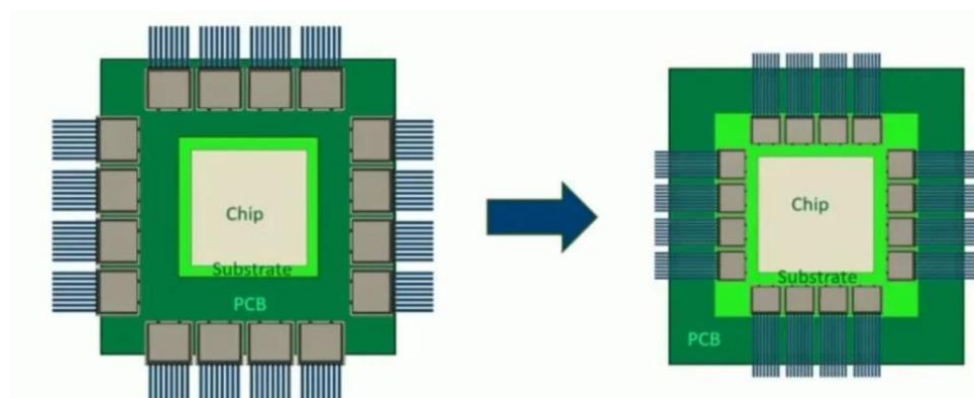


Figure 2.13 NPC and CPC interconnect solutions

2. The optical interconnect route uses optical fiber as the transmission medium, achieving long-distance, ultra-high-bandwidth data transmission via optical signals — an important technical route for breaking the interconnect bottleneck within and between supernode racks.

LPO (Linear Pluggable Optics) retains the traditional pluggable module form; by removing the built-in DSP and moving complex functions such as equalization and clock recovery to the SerDes chip on the ASIC side, it has become an important transitional form for short-distance, low-power pluggable solutions.

The current hot technologies in supernode optical interconnect are NPO (Near Packaged Optics) and CPO (Co-packaged Optics). Both remove the built-in DSP from optical modules and integrate more closely with the ASIC, significantly reducing optical module power (from about 8 W to 3 W per 400G port), cutting single-end latency by about 100 ns, and significantly improving port density and system reliability. The

NPO solution, being more industrially mature and open, currently enjoys broader support from system vendors and users.

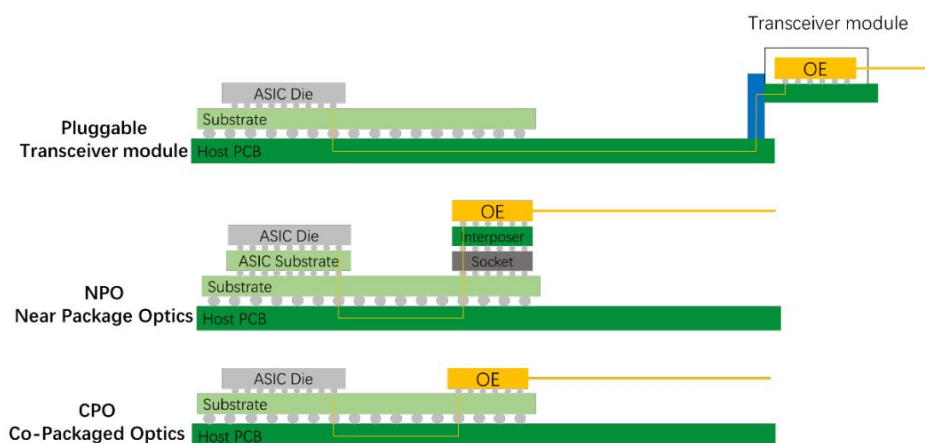


Figure 2.14 Link diagrams of pluggable optics, near-packaged optics, and co-packaged optics

Since the second half of 2025, Chinese vendors have intensively pursued NPO technology exploration and product deployment: Tencent led the launch of the ETX Ultra 512-GPU supernode solution, using NPO modules for all-optical interconnect inside the supernode; Alibaba Cloud debuted NPO modules and a 102.4T NPO Switch concept prototype at the Apsara Conference in September 2025, and was first to apply 3.2T NPO technology in a new-generation domestic four-chip switch.

2.1.2 Distributed Supernodes

A distributed supernode is a system form that efficiently aggregates small-scale compute chips within a single chassis through tight coupling of architecture and high-speed protocols. Its unit scale is smaller and its compute granularity finer, and it can flexibly expand according to business scenario needs, suiting small and medium deployment scales or existing data centers with limited power density.

2.1.2.1 Modules

1. AI acceleration modules. Distributed supernode AI acceleration modules fall into two physical forms: PCIe plug-in cards and board-integrated modules. PCIe plug-in AI acceleration modules typically use PCIe as the inter-card bus, offer high standardization, but have relatively limited heat dissipation and power ceilings, suiting compute scenarios with moderate per-module compute requirements.

Board-integrated AI acceleration modules: using open-standard or vendor-defined interfaces and protocols, they break the interface and size limits of PCIe plug-in AI acceleration modules, support higher power input and system heat dissipation, and deliver stronger AI compute — the currently dominant AI acceleration module form.



Figure 2.15 Board-integrated AI acceleration module design example

With Rubin/Rubin Ultra as the technical reference, GPU integrated modules for training and high-density inference will see power rise further to 3600 W by 2027, with volume, heat flux, and power supply requirements increasing significantly — posing systemic challenges to cold-plate/immersion liquid cooling, power regulation, rack cooling, and structural design.

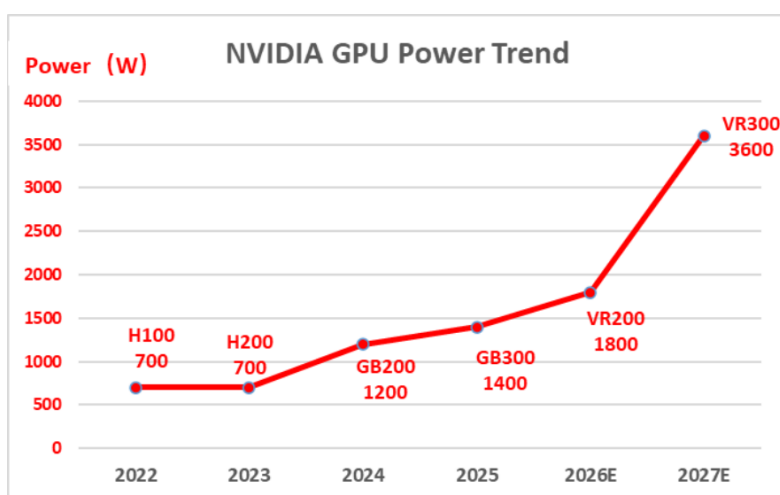


Figure 2.16 NVIDIA GPU power trend

2. AI carrier board modules. The AI carrier board module is the core substrate unit carrying AI acceleration modules such as GPUs. A single carrier board can host multiple AI acceleration modules that interconnect via bus protocols, forming a basic AI acceleration resource pool unit. The carrier board provides high-speed bus and expansion network interfaces, supporting multi-board Scale-Up or Scale-Out expansion.

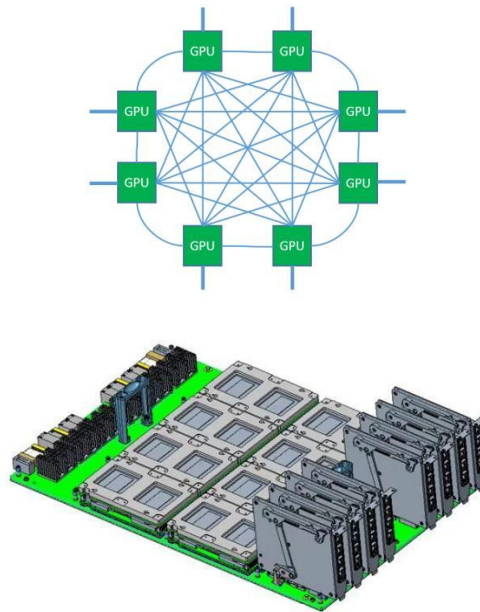


Figure 2.17 AI accelerator carrier board module with expansion cards and its bus topology

As the power, size, and bus rates of AI acceleration modules such as GPUs grow rapidly, AI carrier board module design faces increasingly complex architecture, topology, and engineering implementation challenges.

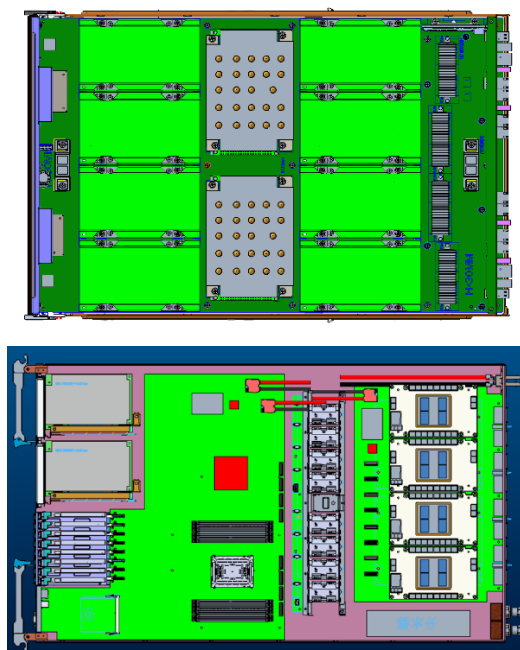


Figure 2.18 Design comparison of 8-GPU and 4-GPU carrier boards

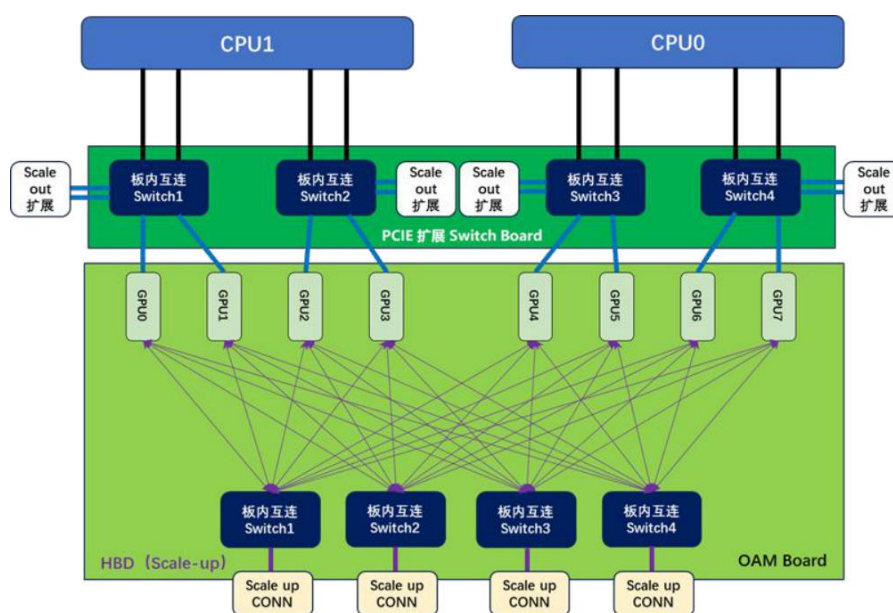


Figure 2.19 8-GPU carrier board system topology

2.1.2.2 Topology

1. Cluster-layer topology. At the cluster layer, distributed supernodes network via IB or Ethernet protocols in fat-tree and other network architectures, with flexible configuration of NICs and data center network topology.

A typical cluster topology is the three-layer CLOS/fat-tree.

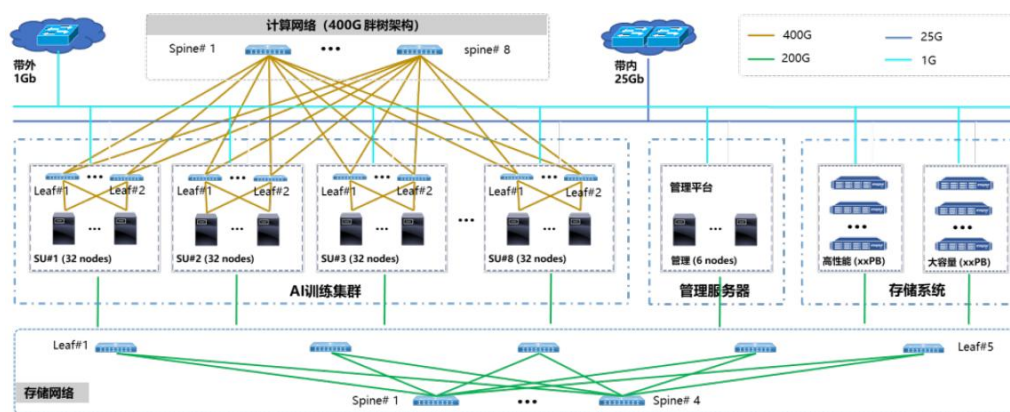


Figure 2.20 Example of a typical AI cluster networking topology

2. Rack-layer topology. The rack layer has two mainstream expansion schemes:

a) Integrating compute, storage, and acceleration nodes in the same rack — more flexible deployment, suiting small and medium clusters; b) integrating only a single node type per rack, then expanding the cluster rack by rack — suiting large-scale clusters. For schemes integrating multiple node types in one rack, interconnection among multiple CPU compute nodes, among GPU nodes (typical scenarios such as PD disaggregation), and among storage nodes should be implemented via bus architectures, preferably using one or more open-standard interconnect protocols including Ethernet, OISA, UALink, and CLink.

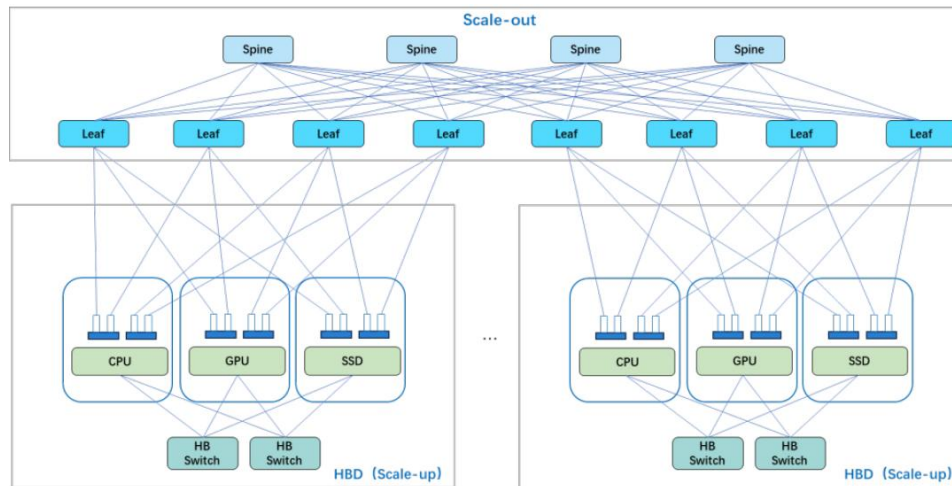


Figure 2.21 Cluster topology integrating compute, storage, and AI acceleration nodes in a single rack

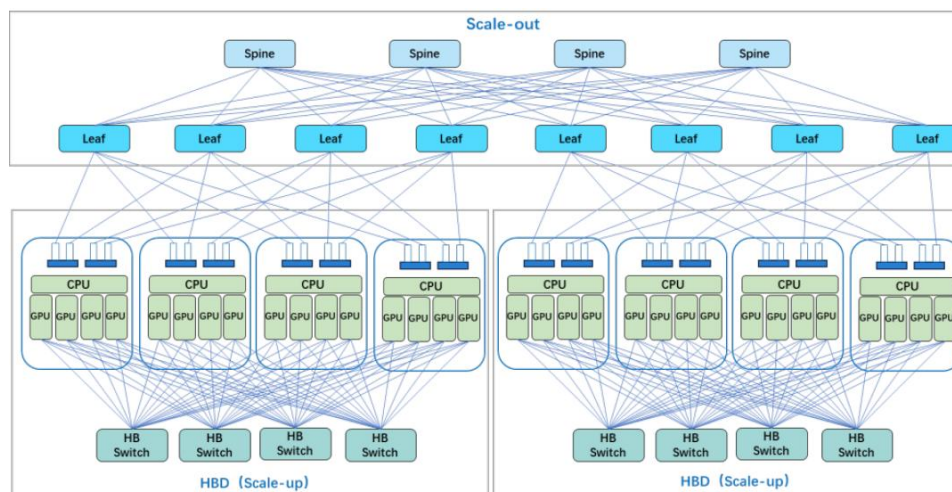


Figure 2.22 Cluster topology integrating only AI acceleration nodes in a single rack

3. Node-layer topology. Includes multiple physical topologies, supporting more flexible CPU, memory, and GPU ratios. a) Traditional CEM standard cards: the CPU connects directly to GPUs via the PCIe bus, and GPUs interconnect within the node via PCIe switches. This topology is simple, highly reliable, and flexible, but has limited intra- and inter-node interconnect bandwidth. b) OAM cards: dominated by direct-connection topologies under the OCP UBB specification (full-mesh, hybrid cubic). Switch chips will be added within carrier boards in the future for programming friendliness and further performance gains. For example, in the NVIDIA BFX NVL16 topology instance below, GPUs connect upstream directly to the CPU and downstream through switches for inter-card interconnect, effectively combining the flexibility of CEM standard cards with NVSwitch's high interconnect bandwidth.

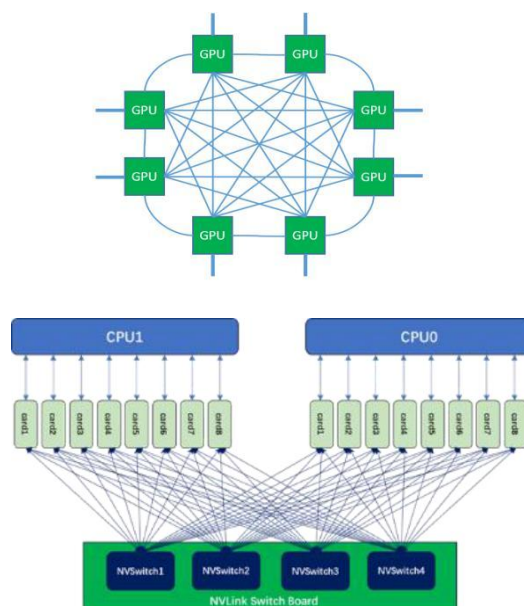


Figure 2.23 Full-mesh topology and NVL16 interconnect topology

2.1.2.3 Power Supply

Distributed supernodes integrate power modules. As GPU module power and deployment density rise, single-system power has reached tens of kW, and single PSU output has exceeded 5 kW. Matching the internal vertical stacking of multiple carrier board modules, systems generally use unified bus power supply.

2.1.2.4 Cooling

Thanks to their small granularity and flexible deployment, a number of distributed supernodes can be integrated with distributed CDUs in the same rack — matching infrastructure conditions — to form a standard rack unit, which then serves as the basis for cluster expansion.

For larger deployments, the independent CDU rack solution can be used for intensive deployment.

2.1.2.5 I/O Expansion Cards

To further improve the expansion flexibility of distributed supernodes and GPU carrier board modules, standard Scale-Up expansion cards can be developed. Such cards often need to integrate signal retiming and amplification to achieve reliable, efficient GPU module interconnection between distributed supernodes or between carrier board modules.

2.2 High-Speed Interconnect Physical Layer and Connectors

2.2.1 High-Speed Interconnect Networks in GW-Scale AIDCs

A supernode is a unified computing unit formed by point-to-point interconnection of multiple GPUs. Its core is high-bandwidth, low-latency intra-node communication via Scale-Up technology, solving the communication bottleneck of AI computing. Scale-Up optimizes multi-GPU coordination within a single node, while Scale-Out supports expansion between supernodes, enabling flexible networking of clusters at the ten-thousand-GPU level. Scale-Up high-speed interconnect hardware applications include compute node high-speed interconnect, switch node high-speed interconnect, intra-rack node-to-node high-speed interconnect, and rack-to-rack high-speed interconnect.

2.2.1.1 Supernode High-Speed Interconnect Hardware Architecture

The larger-scale XPU compute integrated in supernodes, together with the high-bandwidth and low-loss goals driven by exploding compute demand, jointly drive the transformation of high-speed interconnect hardware architecture. Supernodes must comprehensively balance compute capacity, efficiency, heat dissipation, and scalability, and architectures are evolving toward cable tray, orthogonal midplane, orthogonal direct, NPO, and CPO architectures.

Cable tray architecture: nodes slide in horizontally from the front panel and interconnect with the backplane at the rear. It is currently the most widely used industry solution, e.g., NVIDIA NVL72~144, ByteDance Dayu, and Baidu Tianchi (cable architecture).

Orthogonal midplane architecture: compute and switch nodes are loaded vertically or horizontally from the front and rear panels, interconnecting through a midplane, e.g., NVIDIA Rubin Ultra NVL576.

Orthogonal direct architecture: compute and switch nodes mate vertically, directly, at the front and rear panels, e.g., Huawei CloudMatrix 384, Sugon scaleX640, Alibaba Cloud Panjiu AL128 supernode server, and Baidu Tianchi (orthogonal architecture).

2.2.1.2 Compute Node High-Speed Interconnect

The compute node is the core compute carrier of a full-rack supernode, with multiple XPU cards interconnected within a single node (commonly 4 to 16 cards per node in the industry). With the extreme pursuit of performance and compute density in AI large models, copper-based high-speed interconnect faces challenges such as signal attenuation and increased power consumption. The industry has proposed high-density, large-channel cable assemblies, near-package copper (NPC), and near-package optics (NPO), while driving rate upgrades of standard (e.g., PCIe protocol) high-speed cables — all working toward the core goals of signal loss optimization and transmission distance extension.

2.2.1.3 Switch Node High-Speed Interconnect

The switch node is the core interconnect hub of a full-rack supernode. With the explosion of XPU inter-card communication across multiple node groups (16-640 cards per supernode), large numbers of compute units and switch chips must achieve high-density interconnection within limited space.

In GW-scale AIDCs, the extreme pursuit of communication latency and network bandwidth in AI large models has driven the rapid development of Ethernet switch chips: per-lane rates are rising from 56 Gbps to 224 Gbps and 448 Gbps. Under this high-speed network evolution, high-speed interconnect faces multiple challenges including signal attenuation, increased power consumption, and constrained rear-panel space on nodes. The industry has proposed a series of technical paths — high-density large-channel cable assemblies, near-package copper (NPC), near-package optics (NPO), co-packaged copper (CPC), and co-packaged optics (CPO) — and through rapid practice and iteration is driving the co-adaptation of RNICs and new interconnect solutions to jointly ensure stable, efficient operation of large-scale AI networks in high-throughput scenarios.



Figure 2.17 Application schematic of an NPC product for node high-speed interconnect

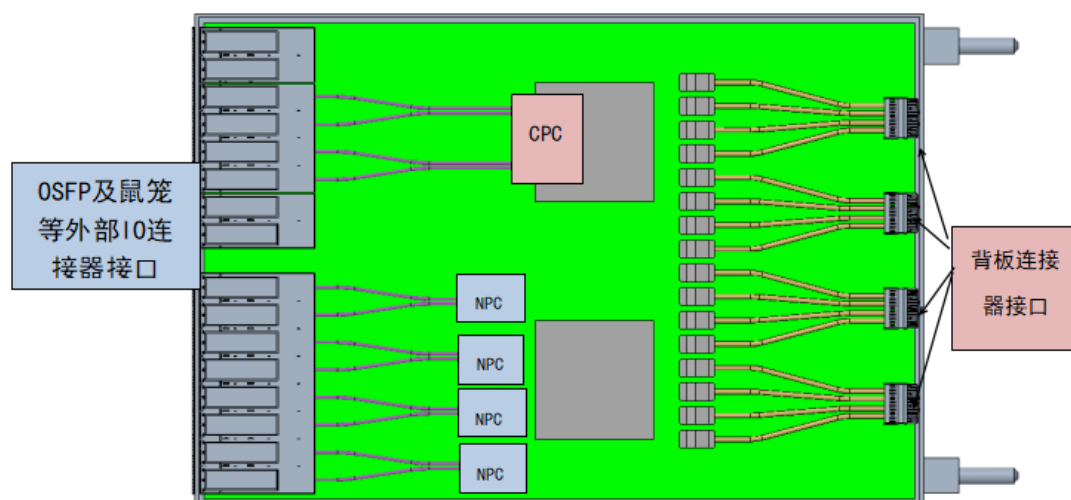


Figure 2.18 Application schematic of switch node high-speed interconnect

2.2.1.4 Intra-Rack Node-to-Node High-Speed Interconnect

In GW-scale AIDC rack applications, the full rack is evolving from the traditional PCB backplane architecture toward cable tray, orthogonal midplane, and orthogonal direct architectures. 1. Cable tray architecture. The cable tray architecture uses high-speed copper cable modules to replace the traditional PCB backplane, building direct connections between compute and switch nodes and forming a flexible mesh topology. It breaks the physical size limits of backplanes, supports high-density, multi-directional expansion, and offers low cost, high reliability assured by floating mechanisms, and multi-XPU compatibility, with cables quickly replaceable for bandwidth upgrades. Its open layout readily accommodates multiple cooling needs — air cooling (optimized airflow) and liquid cooling (with cold plates and manifolds) — enabling efficient cross-node interconnect deployment inside the rack.

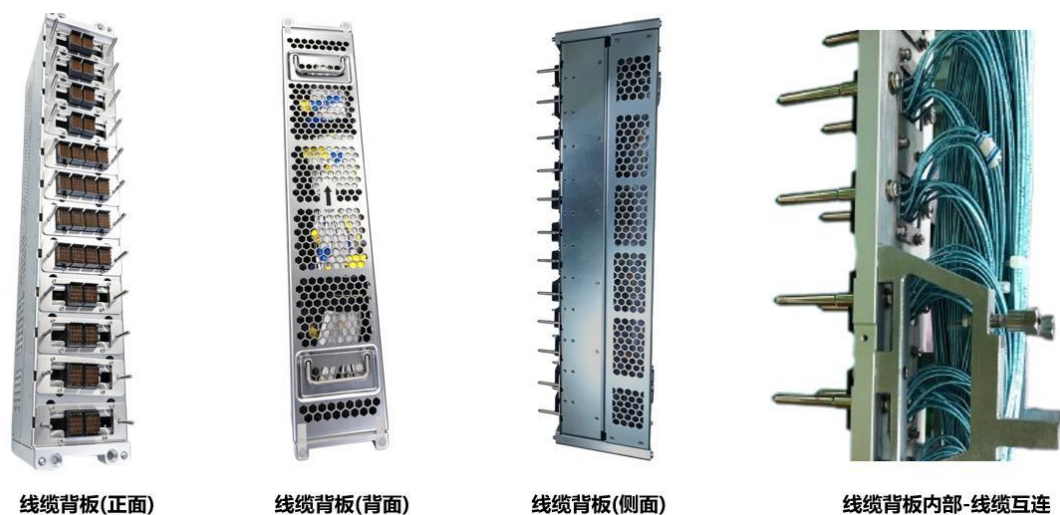


Figure 2.19 A cable tray product schematic

2. Orthogonal midplane architecture. The orthogonal midplane architecture plugs compute and switch nodes at 90° onto a common midplane, building a non-blocking rectangular topology. With extremely short signal links it achieves ultra-low latency and excellent signal integrity; with robust connector design, the connectors support both PCB-level and cabled implementations, front/rear blind-mate maintenance, and efficient direct-through airflow cooling, while offering high density, high reliability, and good cross-generation compatibility — though it also introduces new challenges in architecture and hardware design.

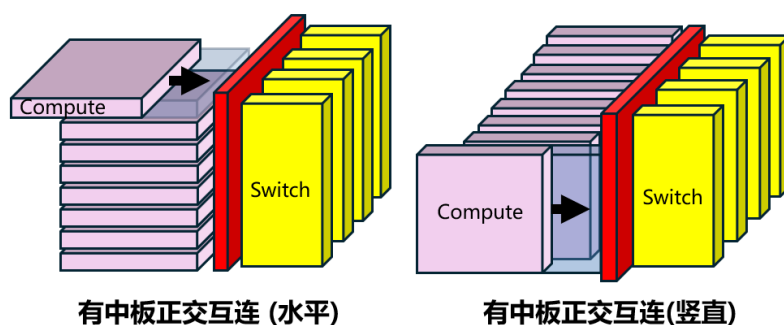


Figure 2.20 Orthogonal midplane architecture schematic

3. Orthogonal direct architecture. The orthogonal direct architecture eliminates the common midplane, using a hardware structure in which compute and switch nodes mate directly at 90°. This design breaks through the line-rate bottleneck of traditional backplanes, achieving ultra-large switching capacity and high scalability while maintaining extremely low signal attenuation. However, this highly integrated direct-mating mode also poses greater engineering challenges for mechanical precision, interconnect stability, liquid cooling planning, power supply, and hardware design.

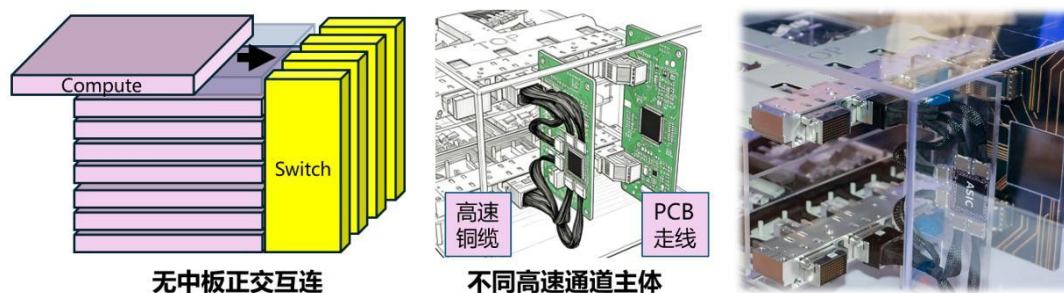


Figure 2.21 Direct orthogonal + PCB architecture interconnect solution

2.2.1.5 Rack-to-Rack Interconnect

Inter-device interconnect solutions fall into two major technical routes: copper and optical. Copper transmission mainly includes DAC (Direct Attach Cable), ACC (Active Copper Cable), AEC (Active Electrical Cable), and direct backplane-connector copper cable assemblies. As data center bandwidth evolves from 100G/400G to 800G/1.6T, copper interconnect is upgrading from traditional DAC to ACC and AEC with stronger signal compensation: ACC suits medium distances such as within a rack, while AEC better suits long-distance interconnect across racks. In optical interconnect, AOC and pluggable optical modules are the core solutions, with evolution focused on system integration achieving higher bandwidth, lower power, smaller size, and better cost.

2.2.2 Interconnect Technology Implementation Paths

For GW-scale AIDCs, high-speed interconnect faces enormous challenges. SerDes technology has evolved from 28 Gbps to 224 Gbps, with 56G/112 Gbps now widely deployed, strongly driving network architecture innovation. A systematic review of interconnect technology paths across all rate segments is therefore key to smooth system upgrades.

2.2.2.1 Termination Forms: Press-fit/SMT/BGA/LGA

To meet rising transmission rates, connector termination forms have evolved: press-fit finished holes keep shrinking, gradually toward 0.25 mm; SMT, BGA, and LGA pitches continue to densify below 0.6 mm, solving impedance consistency problems in high-bandwidth scenarios.

Table 2.1 Comparison of termination form application information

Termination form	Typical specification range	Advantages	Disadvantages
Press-fit	56 Gbps recommended finished hole diameter: 0.31-0.45 mm	High mechanical reliability; easy removal, repair, and rework;	Hole diameter limited by PCB stack-up thickness, stamping, and other processes; limited spatial density;
	112 Gbps recommended finished hole diameter: 0.29-0.36 mm		
	224 Gbps recommended finished hole diameter: 0.25-0.36 mm		

SMT	Terminal pitch: 0.5-0.6 mm	Higher density and miniaturization than Press-fit	Difficult repair and rework;
BGA	Terminal pitch: 0.4-0.6 mm	Higher density and miniaturization	Difficult repair and rework;
LGA	Terminal pitch: 0.4-0.6 mm	Higher density and miniaturization	Mechanical reliability slightly lower than Press-fit;

2.2.2.2 High-Speed Connectors

High-speed backplane connectors implement the interconnect form between a single node's rear panel and its mating node. According to differentiated interconnect application needs, board-level or cabled application types can be selected, extending across different series of 56 Gbps to 224 Gbps high-speed connectors. [For sections 2.2.2.2-2.2.2.4, see the detailed connector selection panorama in the product table of the AIDC High-Speed Interconnect Technology Technical Report (2026).]



Figure 2.22 A high-speed backplane connector and cable assembly product schematic

2.2.2.3 Internal High-Speed I/O

Internal high-speed I/O implements high-speed interconnect within a single node. According to differentiated interconnect protocol needs, standard connectors (MCIO, Gen-Z, etc.) or customized connectors (NPC/CPC) can be selected, adapting to the high-speed characteristics of different protocols such as PCIe 5.0 to PCIe 7.0 (32 GT/s to 128 GT/s per lane), with SerDes rates covering 56 Gbps to 224 Gbps.



高速I/O连接器（标准类型）

高速I/O连接器—NPC（定制类型）

某 LGA 形态的 NPC（定制类型）

Figure 2.23 A high-speed I/O connector and customized connector product schematic

2.2.2.4 External High-Speed I/O

External high-speed I/O implements high-speed interconnect across nodes and racks. According to differentiated interconnect protocol needs, standard cables (QSFP, OSFP-XD, etc.) or customized

backplane connector cables can be selected, adapting to the high-speed characteristics of different protocols. For cable lengths: passive DAC cables are mostly 0.5-2.0 m, active ACC cables mostly 2.0-4.0 m, and active AEC cables mostly 4.0-7.0 m.



Figure 2.24 An external high-speed I/O cable product schematic

2.2.2.5 High-Speed PCB Laminates

Loss requirements for high-speed PCB laminates by signal rate are as follows:

Table 2.2 High-speed PCB laminate application information

Signal rate	Dissipation factor (Df)	Non-domestic laminates	Domestic laminates
56 Gbps	0.003 ~ 0.005	MEGTRON6 (G), EM891, TU-883, IT-968, etc.	NY6300S, H360Z, S6, etc.
112 Gbps	0.002 ~ 0.003	MEGTRON7 (N), EM890K, TU-933, IT988GSE, etc.	NY-P2, S6NX, etc.
112 Gbps	0.0015 ~ 0.002	MEGTRON8 (N), EM892K, TU943SN, IT998G, etc.	NY-P4N, S9N, etc.
112 Gbps	0.0015 ~ 0.002	MEGTRON8 (N), EM892K, TU943SN, IT998G, etc.	NY-P4N, S9N, etc.
224 Gbps	0.0001 ~ 0.0015	MEGTRON8 (U), EM892K2, IT998GSE, etc.	NY-P4, S9N2, etc.
224 Gbps	0.00005 ~ 0.0001	MEGTRON9 (U/Q), EM896K2/K3, IT999GSE2/3, etc.	NY-P5N2/Q, S10GN2/Q, etc.

2.2.2.6 High-Speed Raw Cable

Higher-speed raw cable requires smaller dimensions and better SDD21 loss. Designs therefore often select insulation materials with lower dielectric constant (Dk) and dissipation factor (Df), or reduce loss by enhancing coupling between conductors.

Solid insulation cannot achieve a Dk below 2.0 due to material limits; conventional PE/PP/FEP materials have Dk values of 2.03-2.3. To reduce loss, high-speed raw cable insulation has evolved from single-layer insulation to double-layer insulation, foamed insulation, or dual-core co-extrusion structures.

Typical SI indicators are also becoming more stringent: for example, center impedance matched to the connector with tolerance tightened from $\pm 5 \Omega$ to $\pm 2 \Omega$; intra-pair skew reduced from 10 ps/m to 2 ps/m; and smooth SDD21 curves with bandwidth rising to 40 GHz, 67 GHz, even 110 GHz without spikes.

The table below gives the SI performance of typical high-speed copper cables without ground wires; double-ground or single-ground constructions can also be selected as needed.

Table 2.3 High-speed raw cable application information

AWG	Conductor OD	Insulation structure	Ground wire	Tape dimensions (mm)	Impedance (Ω)	112G SDD21 @26.56GHz	112G SDD21 @53.12GHz	224G SDD21 @84GHz	448G SDD21 @110GHz
32	0.203 SC	FEP dual-core co-extrusion	None	1.55×0.89	92	7.9-8.4 dB	11.5-12.0 dB	16.0 dB	19.5 dB
30	0.285 SC	FEP+FEP	None	1.70×1.20	92	6.1 dB	8.8 dB	12.3 dB	14.5 dB
28	0.35 SC	Foam FEP+PE	None	1.99×1.18	92	5.4-5.6 dB	8.3 dB	10.2 dB	11.9 dB
26	0.394 SC	Foam FEP+PE	None	2.10×1.14	92	4.7-5.1 dB	8.0 dB	9.7 dB	11.3 dB
25	0.455 SC	Foam FEP+PE	None	2.55×1.32	92	4.1-4.4 dB	7.3-7.7 dB	8.4 dB	9.8 dB

2.2.3 Outlook for Next-Generation Interconnect Technology Paths

With the explosion of AI large models and ten-thousand-GPU clusters, GW-scale AIDC interconnect faces rigid requirements for high bandwidth, low latency, and low cost, pushing per-lane rates toward 448 Gbps. Interconnect technology is currently making parallel breakthroughs along the dual paths of copper and optical.

Copper interconnect is undergoing system-level innovation: from coding modulation (PAM4/6/8) to architecture (evolving toward various orthogonal forms) to product structure upgrades. The core focus, CPC (co-packaged copper), uses an integrated "chip-connector-copper cable" design that, compared with traditional solutions, offers lower link loss, lower power, lower system cost, and both high density and maintainability. Meanwhile, the evolution of PCB and raw cable — with the introduction of new materials such as high-frequency laminates and foamed insulation — helps copper interconnect continually push the boundaries of interconnect distance and tackle high-speed interconnect challenges at 110 GHz and above.

Optical interconnect is evolving from pluggable modules toward NPO/CPO. CPO technology co-packages the optical engine with the electrical chip, so electrical signals travel only a few millimeters, greatly improving bandwidth density and reducing bit error rate. It is key to breaking through panel I/O density limits, but the industry ecosystem is still in an exploratory phase and the overall supply chain is incomplete, requiring upstream and downstream partners to jointly mature product commercialization.

In the future, "copper-optical synergy" will be the long-term landscape: copper interconnect holds short-distance in-rack scenarios while optical interconnect handles long-distance cross-rack communication — complementary strengths that together build a high-efficiency AIDC foundation.

2.3 AI Network and Switching System

2.3.1 Scaling

The AI Scale-Out network is evolving toward ever-larger cluster architectures, driven mainly by two forces.

1. Technical demand: advanced large models such as GPT, Claude, and DeepSeek have proven that a model's intelligence is directly and positively correlated with its parameter scale, training data volume, and total compute. To pursue stronger reasoning, generalization, and multimodal capabilities, hyperscale clusters of more GPUs must be built, placing extremely high demands on the Scale-Out horizontal scalability of the underlying network.

2. Economic efficiency: in fierce global AI competition, shortening model training cycles has become a core competitive advantage. Larger GPU clusters significantly accelerate training tasks, and the resulting efficiency gains directly reduce the operating cost per unit of compute, achieving dual optimization of efficiency and cost.

2.3.1.1 Rapidly Evolving Switch Chip Bandwidth Drives Generational Leaps in AIDC Networks

The coordinated development of AIDC switch chips and networking architectures is the cornerstone supporting the evolution of ten-thousand-GPU clusters. Its evolution path clearly reflects the shift from "using architectural complexity to compensate for insufficient bandwidth" to "using chip performance to drive architectural simplification."

Early architecture-compensation stage: limited by per-chip switching bandwidth, the industry had to adopt complex architectures such as chassis-box networking and three-tier box-box networking to meet early ten-thousand-GPU cluster networking needs.

Current chip-driven simplification stage: as switch chip performance improves, per-chip bandwidth has moved from 12.8T into the 51.2T and even 102.4T era. This makes simple, flat two-tier box-box networking the mainstream, significantly reducing management complexity and cost.

The core driver of this shift is the continuous process and architecture breakthroughs of chip suppliers represented by Broadcom. By adopting advanced process nodes such as 5 nm and 3 nm and successively releasing chip families from 51.2T to 102.4T, they provide higher integration and bandwidth performance, laying a solid hardware foundation for building simple, efficient AIDC networks.



Figure 2.25 Broadcom AIDC chip roadmap

To match the rigid demand for switch chip capacity to evolve from 51.2T to 102.4T and beyond, per-lane interface rates and signal modulation technologies are undergoing a critical generational transition.

Port rate and SerDes technology trends: 400G ports still dominate the industry today, with 800G about to explode. Constrained by chip process and ecosystem maturity, 112G SerDes-based technology platforms will persist in China for some time. As 224G SerDes gradually lands, NPO (near-packaged optics) modules are expected to deploy at scale starting in 2027, and CPO (co-packaged optics) switches are expected to reach large-scale commercial use after 2028.

Co-evolution of SerDes and modulation: per-lane SerDes rates are rapidly migrating from 56 Gbps to 112 Gbps, and the Tomahawk 5 chip fully supports 112 Gbps. The next-generation Tomahawk 6 will further introduce 224 Gbps SerDes. In modulation, PAM4 has become the mainstream standard for high-speed interconnect, with twice the transmission efficiency of traditional NRZ. In future 102.4T and above chips, PAM4 will continue to combine with 112G/224G SerDes to build a high-bandwidth, high-energy-efficiency physical layer foundation.

2.3.1.2 Elastic AIDC Network Architecture Supporting Hyperscale Expansion

The Scale-Out network of an AI compute cluster aims to achieve efficient, lossless interconnection of GPU servers. Its networking architecture must scale elastically with cluster size and falls mainly into the following modes.

1. Single-switch direct connection: for small-scale inference. Scenario: small inference or development/testing scenarios with tens of GPUs. Solution: connect all GPUs directly to a single high-performance switch, forming the simplest non-blocking network. For example, H3C's S9827-128DH series switches based on the Tomahawk 5 chip support up to 128 400G ports per device, meeting hundred-GPU cluster needs.

Value: easily supports hundred-GPU clusters with quick deployment and simple O&M. 2. Two-tier CLOS architecture: supporting medium and large training clusters. Scenario: medium and large training clusters of hundreds to ten thousand GPUs — the current mainstream solution. Solution: the classic Spine-Leaf two-tier flat architecture. By selecting Spine/Leaf switches of different bandwidths (12.8T to 102.4T) and linearly increasing device counts, flexible, linear scaling from 512 to 16K GPUs is achieved while maintaining constant low latency and high bandwidth.

Value: a three-hop non-blocking network with low latency; simpler design, deployment, and O&M; and easy fault localization.

Table 2.4 Two-tier CLOS Scale-Out network architecture scale

Spine/Leaf configuration	Quantity	400G NIC access scale
Spine: 12.8T, 32×400G	Spine: 16 units	512
Leaf: 12.8T, 32×400G	Leaf: 32 units	
Spine: 25.6T, 64×400G	Spine: 32 units	2K
Leaf: 25.6T, 64×400G	Leaf: 64 units	
Spine: 51.2T, 128×400G	Spine: 64 units	8K
Leaf: 51.2T, 128×400G	Leaf: 128 units	
Spine: 102.4T, 128×800G	Spine: 64 units	16K
Leaf: 102.4T, 128×800G, access split 1-to-2	Leaf: 128 units	

Note: in the 102.4T solution, Leaf switches use 800G ports with breakout cables (e.g., 1-to-2) to connect 400G GPU NICs, doubling port density.

3. Three-tier CLOS architecture: building hundred-thousand-GPU hyperscale clusters. Scenario: single hyperscale training clusters of tens of thousands to a hundred thousand GPUs. Solution: add a Core layer above the two-tier architecture, forming a Core-Spine-Leaf three-tier architecture. The architecture scales horizontally based on POD (performance-optimized module) units, each containing a two-tier CLOS inside. By stacking multiple PODs interconnected via the Core layer, near-unlimited linear scalability is achieved — the infrastructure for training next-generation trillion-parameter models.

Value: an irreplaceable hyperscale expansion solution that two-tier CLOS cannot replace. However, it also has issues such as more hops and added latency, complex design/deployment/O&M, and high cost.

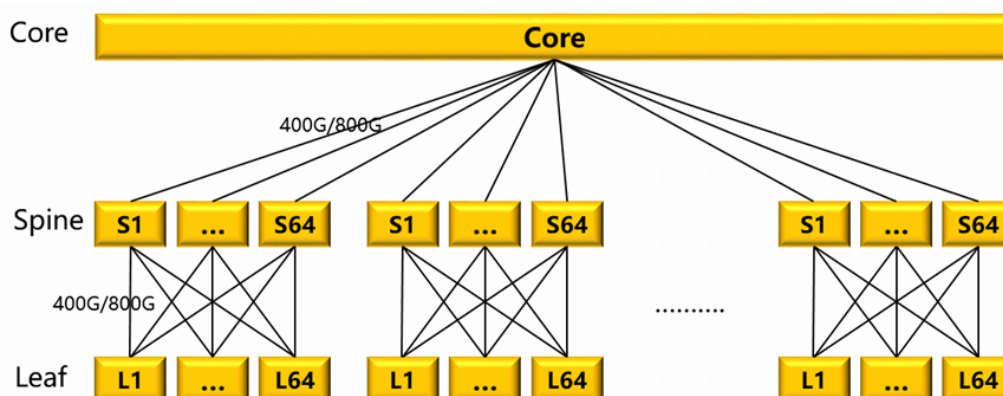


Figure 2.26 Three-tier CLOS Scale-Out network architecture

4. Multi-plane networking: a heterogeneous expansion solution for high-performance NICs. Scenario: paired with high-performance SmartNICs supporting port aggregation and splitting (e.g., NVIDIA CX8 800G, currently supporting 2×400G), using 2 planes — two independent two-tier CLOS networks — to provide tens-of-thousands-GPU super clusters. Solution: connect a server's multiple NIC ports to multiple independent two-tier Spine-Leaf network planes. Each plane provides the server with a complete network

path, achieving cross-plane bandwidth aggregation and fault isolation. Based on the Tomahawk 6 chip, a single plane can provide 32K 400G interfaces.

Value: using simple, low-cost two-tier CLOS, it is still possible to build tens-of-thousands-GPU super clusters.

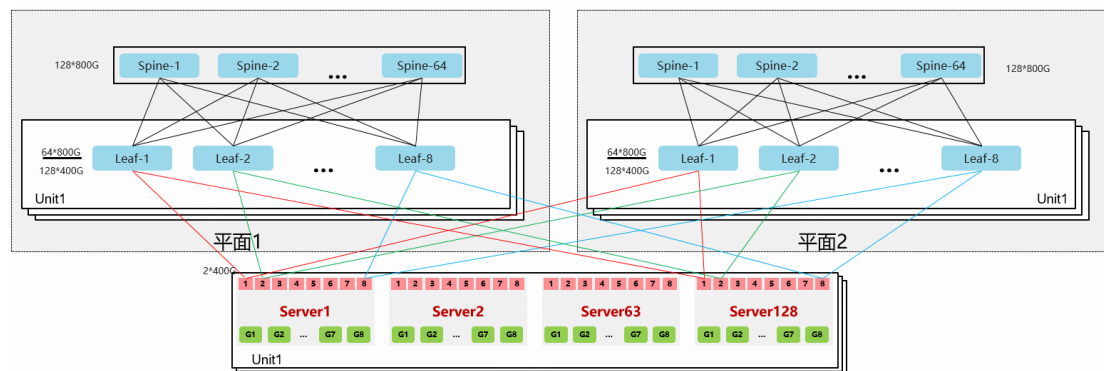


Figure 2.27 Dual-plane Scale-Out network architecture

2.3.2 Communication Efficiency

AIDC network traffic is dominated by a few fixed collective communication operations such as All-Reduce and All-to-All. The traffic has typical characteristics of low entropy (few flows with uniform features), burstiness, and long flows, which easily trigger instantaneous congestion in the switching network.

Because GPU training follows a strict synchronization paradigm, all nodes must complete collective communication before entering the next round of computation, making the "bucket effect" likely. End-to-end latency jitter caused by congestion therefore directly slows the overall training task, becoming the core performance bottleneck of AIDC networks.

To overcome this, the industry focuses on improving network communication efficiency and has proposed two technical paths — end-network coordinated scheduling and pure network-side scheduling — aiming for finer load balancing and congestion control.

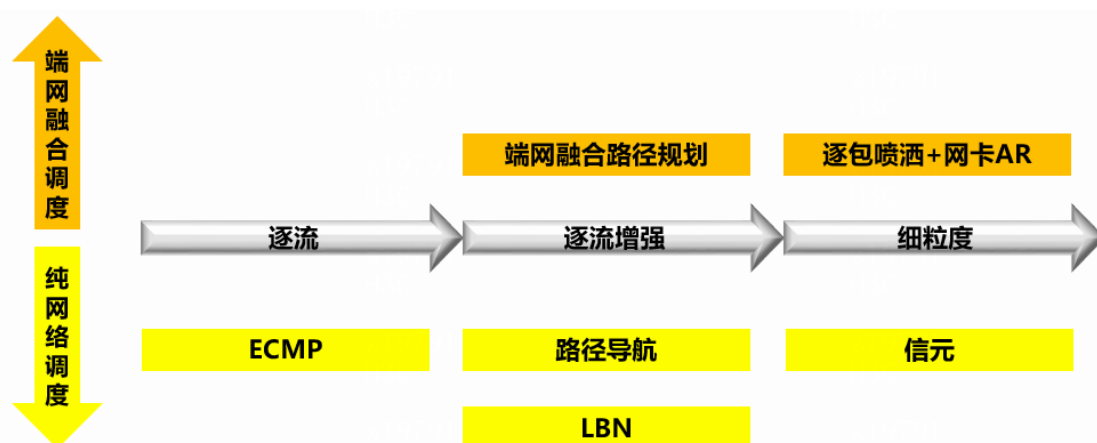


Figure 2.28 Scale-Out network load balancing technology system

2.3.2.1 Pure Network-Scheduled Load Balancing Technologies

Refers to technical solutions in which network devices (mainly switches) autonomously implement traffic optimization scheduling, without any cooperation or awareness from servers or NICs.

Table 2.5 Pure network-scheduled load balancing technologies

Item	Description
ECMP	Switch N-tuple hash algorithm; currently the most widely used method in the industry. Due to the low-entropy characteristic, load balancing may be suboptimal; it can be tuned via multiple QPs, hash algorithm adjustment, and other means.
Path navigation	The controller collects switch information, discovers congestion points and contributing flows, and replans to move congested flows to lightly loaded paths. It only suits scenarios with fixed communication relationships and performs poorly in unpredictable-traffic scenarios such as MoE.
LBN	Ingress-port-based hashing ensures uniform Leaf uplink traffic. In large-scale networks where multiple Leafs share the same Spine for downlink, congestion probability is relatively high; suits small networks whose end sides do not support multiple QPs.
Cell switching	Cell switching splits traffic into cells at the Leaf, sprays them evenly across all Spines, and reassembles them, achieving load balancing. Naturally heterogeneous and decoupled, it requires no NIC support and adapts to any traffic.

2.3.2.2 End-Network Converged Scheduling Load Balancing Technologies

A technical architecture that achieves global traffic optimization through information exchange and joint decision-making between the network side (controllers, switches) and the compute side (SmartNICs, collective communication libraries).

Table 2.6 End-network converged scheduling load balancing technologies

Item	Description
End-network converged path planning	The controller obtains GPU communication relationships via the CCL and pre-plans non-overlapping network paths to avoid

	congestion. It suits training scenarios with fixed communication patterns and performs poorly in unpredictable-traffic scenarios such as MoE.
Per-packet spraying + NIC AR (Adaptive Routing)	Packet-level uniform spraying paired with the receiving NIC's AR out-of-order reassembly capability delivers superior load balancing. It currently relies mainly on NVIDIA's closed ecosystem and is expected to proliferate as domestic and Broadcom NICs support AR.

2.3.3 Reliability and Maintenance

2.3.3.1 Switch Node-Level Reliability and Maintenance

AIDC networks typically consist of tens to hundreds of box-architecture switches. Their hardware design must ensure that key components (such as power supplies and fans) support M+N redundancy, with fault monitoring and O&M capabilities for core components such as CPUs and memory.

2.3.3.2 Network Architecture-Level Reliability and Maintenance

Network-internal reliability: relies primarily on ECMP for fast self-healing. Taking the H3C S9827-128DH as an example, a 64-path ECMP group can form between Spine and Leaf; when a device or link fails, traffic automatically switches to other paths, achieving lossless self-healing. Combined with data-plane fast convergence, failover packet loss can be controlled at the millisecond level, ensuring business continuity.

Network-edge reliability: multi-plane redundant architecture improves access reliability. Server NICs are dual-homed to Leafs in two planes; normal traffic is forwarded in isolation across the dual planes. When an access link or device fails, end-side CX8 & NCCL escape mechanisms provide fault redundancy, avoiding single points of failure at the network edge.

2.3.3.3 Link-Level Reliability and Maintenance

AIDC networks rely on tens of thousands of optical modules for interconnection, and their transmission stability directly affects the continuity of AI training and inference tasks. To ensure reliability, the industry adopts two levels of assurance: pre-training and during-training.

Pre-training PRBS stress testing: after networking is complete and before services go live, PRBS (pseudo-random binary sequence) patterns are widely used to test all optical modules at scale. This screens out substandard optical modules in advance, avoiding service interruptions caused by hardware quality problems at the source.

During-training optical link monitoring: once services are running, optical modules undergo probabilistic performance degradation over time and with environmental factors, which can cause port failures or link flaps. The industry therefore widely monitors key optical module indicators in real time — signal-to-noise ratio (SNR), forward error correction (FEC) counts, symbol error rate (SER), and bit error rate (BER). By analyzing degradation trends in these indicators, potentially faulty modules can be flagged and replaced

early, enabling predictive maintenance that safeguards transmission quality and significantly reduces the risk of AI task interruption.

Through this combined strategy of "strict screening before go-live" and "continuous monitoring during operation," AIDC networks can build a highly reliable optical interconnect layer, providing a stable foundation for large-scale AI services.

2.3.3.4 Service-Level Reliability and Maintenance

AIDC networks impose higher requirements for congestion monitoring, which traditional second-level port statistics on switches cannot meet. Limitations of traditional monitoring: switches can detect network congestion through statistics such as ECN marking and PFC message rates, indicating performance bottlenecks. But this monitoring is coarse-grained: it can neither capture instantaneous micro-burst traffic at the millisecond to hundred-millisecond level nor locate the specific service flows causing congestion, making sudden congestion hard to observe and display effectively.

Micro-burst monitoring capability: to solve this, AIDC switches must collect port buffer statistics at millisecond granularity usage and forwarding rates. This allows the system to accurately capture and report instantaneous traffic peaks when micro-bursts occur, preventing them from being "smoothed out" by second-level averages, thereby enabling effective burst monitoring and problem localization.

Flow-level monitoring and tuning: to further locate the root cause of congestion, flow-level monitoring must be enabled. By building flow tables for service flows and collecting per-flow forwarding latency, the "congestion-contributing flows" with significantly increased latency can be precisely identified. Combined with the network controller, these flows can be dynamically redirected to lightly loaded paths, achieving real-time, state-based load balancing tuning and optimizing network performance.

In summary, by combining millisecond-level port monitoring with flow-level latency analysis, AIDC networks can achieve closed-loop management from "detecting congestion" to "locating the root cause" to "dynamic tuning," effectively addressing the challenges of high-burst, low-latency AI traffic.

2.3.4 Future Outlook

2.3.4.1 Network Management and Control Software Supporting AIDC Services

AIDC network construction and O&M is a systems engineering effort spanning multiple service layers, requiring deep integration and coordination between the network and switches, NICs, GPUs, and other devices. Centered on automation, systematization, and intelligence, network management and control software provides three key capabilities.

End-network coordination for ultra-fast rollout: through global visualization and automated deployment, achieve unified management of devices, GPUs, and NICs, with one-click policy delivery and health stress testing, greatly improving large-scale network delivery efficiency.

Global visibility and network tuning: for AIDC traffic characteristics such as low entropy, burstiness, and elephant flows, achieve service-level congestion awareness and dynamic tuning for dynamic load balancing and global path optimization, improving network throughput and training efficiency.

Intelligent O&M safeguarding training: build a full-cycle assurance system covering pre-training inspection, during-training real-time monitoring (paths, congestion, hashing), and post-training fault demarcation and traceback, achieving fast sensing, localization, and optimization.

This forms a closed loop of "fast delivery - dynamic tuning - intelligent O&M," providing stable, efficient, and visible network support for AIDC services.

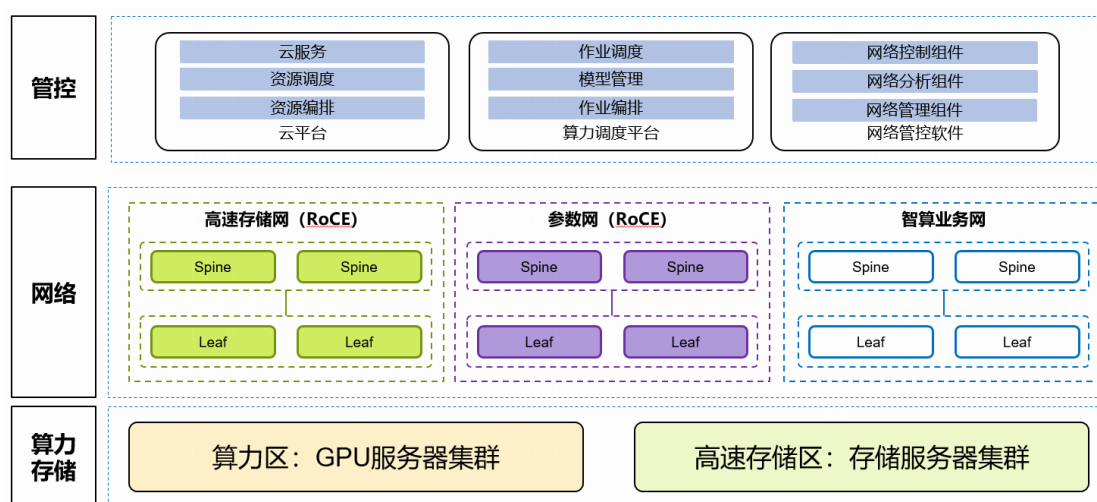


Figure 2.29 AIDC network management and control software architecture

2.3.4.2 A Healthy Ecosystem of Openness and Decoupling

AIDC networks are mainly composed of switches, NICs, and GPUs. Taking a typical GPU-to-GPU RDMA Write operation as an example, the DMA operation process is as follows.

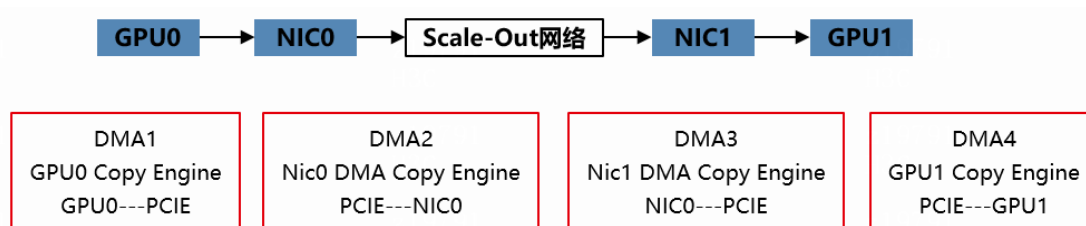


Figure 2.30 Participants in the RDMA process from a hardware perspective

AIDC network construction often faces closed-ecosystem and integration challenges. The current market mainly consists of different vendors supplying switches, NICs, GPUs, and other independent components, requiring system integration to build complete solutions. However, some large vendors, leveraging full-stack product capabilities, tend to promote closed "all-in-one" solutions, objectively limiting customers' opportunities to choose the industry's best technology and product combinations, and bringing problems such as single-source supply chains and uncompetitive costs.

To address this challenge, DDC (Diverse Dynamic Connectivity) technology emerged. Its core value lies in decoupling the network from upper-layer AI frameworks and heterogeneous compute. By fully integrating functions such as cell spraying, packet reassembly, and reordering into the network side, DDC lets customers obtain stable, high-performance transmission regardless of which brand or model of GPU and NIC they use — even if those lack advanced out-of-order handling — truly achieving "heterogeneous compatibility with worry-free performance." This technology not only breaks ecosystem lock-in but also provides key technical support for customers to build open, efficient, and independently controllable AIDC networks.

2.4 AI Storage Technology and Systems

2.4.1 Overview

For AI storage infrastructure construction in large equipment halls, the capability boundaries of storage nodes — which carry training hot data, inference caches, persistence and high-speed reads of vector index files, checkpoints, parallel file services, and metadata services — are jointly determined by the following factors: first, the scale of PCIe lanes and memory channels provided by the CPU and its NUMA topology characteristics; second, front-end NVMe media density and local flash pool capability; and third, the access method of 400G and emerging 800G-class RDMA networks.

Based on current deployment practices in large equipment halls, high-density all-flash nodes have become an important representative form of high-performance AI storage servers. Such nodes typically provide high-bandwidth data services for the hot data tier, with the design goal of systematically matching CPUs, SSDs, memory, and high-speed RDMA NICs.

2.4.2 AI Storage Technology

2.4.2.1 High-Speed Inference Storage Architecture

The high-speed inference storage architecture can be summarized as four tiers: GPU VRAM, CPU memory, local SSD, and network SSD. GPU VRAM carries the hottest KV cache and current request context; CPU memory serves as the active cache and swap-in/swap-out buffer; local NVMe SSDs provide low-latency capacity expansion; and network SSDs or remote flash pools handle cross-node sharing, elastic expansion, and warm/cold data.

The significance of this architecture is placing inference context in tiers by access hotness, so the GPU does not have to wait frequently for data due to long context, multi-turn dialogue, or agentic AI tasks.



Figure 2.31 Four-tier high-speed inference storage architecture schematic

The NVIDIA Inference Context Memory Storage Platform (CMX), released by NVIDIA at CES in early 2026, is a dedicated "inference context memory storage platform" for long-context LLM, multi-turn dialogue, and agentic AI inference. Its core is hardware acceleration via the BlueField-4 DPU, inserting a Pod-shared context tier — a "tier 3.5" (see table below) — between GPU HBM and traditional storage, specifically solving the problems of KV cache VRAM overflow, high latency, and high power consumption.

Table 2.7 NVIDIA divides the entire AI memory/storage hierarchy into G1-G4, with CMX inserted in the middle:

Tier	Location	Medium	Latency	Purpose
1	GPU local	HBM3e	~100 ns	Model weights, hottest KV
2	CPU near-memory	Vera CPU + LPDDR5X	~1 μ s	Overflow KV, medium hotness
3.5 (CMX, tier 3.5)	Pod shared	BlueField-4 + Flash	~5 μ s	Long-context KV, multi-turn sharing, agent memory
4	Cluster / cloud	Traditional NAS/S3	~ms	Cold data, persistence

2.4.3 AI Storage Nodes

2.4.3.1 CPU

In storage server architecture, the core value of x86 CPUs lies mainly in a mature I/O ecosystem, stable enterprise-class RAS features, and long-term adaptability to mainstream storage software stacks. For AI storage nodes, CPU selection must consider: how a single server supports direct attachment of large numbers of NVMe SSDs; and how to attach one or more high-speed RDMA NICs.

Platform planning starts with the number of PCIe lanes. Enterprise NVMe SSDs generally connect via PCIe x4 interfaces; 400G/800G-class high-speed NICs usually connect via PCIe x16 slots; in addition, lanes are consumed by boot drives, BMCs, onboard controllers, and other devices. In terms of PCIe 5.0 lanes directly available from the CPU (native CPU direct-attached lanes, not the final usable count on the

motherboard): Intel Xeon 6 (Granite Rapids, 6th Gen Xeon) offers up to about 96 lanes per socket; AMD EPYC 9005 (Turin) offers up to 128 lanes per socket, and 160 in dual-socket configurations. Note that dual-socket systems do not simply double the count — some lanes between the two CPUs are used for UPI/Infinity Fabric interconnect, and actual usable lanes are further consumed by boot drives, BMCs, and so on. Different platforms therefore differ significantly in single-socket direct-attach capability, dual-socket expansion approaches, and whether an external PCIe Switch is needed for expansion.

When a node deploys multiple high-speed RDMA ports, the affinity among CPU, NICs, and SSDs is especially critical. NICs should be in the same NUMA domain as the NVMe pool carrying the main I/O, reducing cross-socket data movement. With dual-socket architectures, fully evaluate the impact of cross-socket paths on bandwidth efficiency and tail-latency stability.

2.4.3.2 Memory

Two scenarios significantly increase memory capacity configuration. First, when the KV cache uses memory caching, its capacity is approximately linearly related to token count, model layer count, KV head count, per-head dimension, data type, and concurrent session count. In long-context or high-concurrency inference scenarios, cache demand easily exceeds per-GPU VRAM limits. The industry therefore typically adopts a four-tier offload strategy of GPU VRAM → CPU memory → local SSD → network SSD. CPU memory not only serves as the intermediate buffer layer holding swapped-out KV blocks but also handles prefetching and asynchronous swap-in/swap-out scheduling. For capacity planning, memory is generally configured at 3-6 times GPU HBM; for example, with 4 H100s (320 GB HBM total), 1-2 TB of DRAM is recommended, dedicated to inactive KV blocks swapped out of GPU VRAM.

Second, the metadata service layer of parallel file systems requires large amounts of memory. Memory capacity is closely related to the metadata hot set composed of inodes, dentries, stripe mappings, and client lock tables. Memory usage of client lock tables typically grows approximately linearly with concurrent client count; under high lock contention, frequent lock callbacks and lock upgrades can push growth beyond linear, so extra capacity should be reserved during planning.

2.4.3.3 NVMe SSD

At the local persistence media layer, NVMe SSDs are the core component of AI storage servers. The mainstream configuration direction is enterprise-class NVMe SSDs, commonly in U.2/U.3 and EDSFF form factors such as E1.S and E3.S. Scenarios such as AI training, inference, checkpoints, vector indexes, and high-concurrency metadata access have clear requirements for low latency, high-concurrency random read/write, and stable sequential throughput. SSD selection should therefore comprehensively evaluate steady-state performance, write consistency, power-loss protection, thermal management, and long-term serviceability.

By interface generation, PCIe 5.0 x4 has become the mainstream foundation of high-performance enterprise NVMe SSDs, with clear improvements over PCIe 4.0 in per-drive sequential read/write, random

IOPS, and queue concurrency, better handling loads such as training hot data, checkpoint writes, vector index loading, and inference cache expansion. For AI storage nodes, Gen5 SSDs mean not only faster individual drives but also higher node throughput with fewer drives, reducing structural pressure on drive bays, power, PCIe switches, and network egress.

PCIe 6.0 raises the link rate to 64 GT/s and introduces PAM4/FLIT encoding with stronger error correction, providing a higher host-side bandwidth ceiling for next-generation Gen6 NVMe SSDs. As Gen6 SSDs, PCIe 6.0 switches, CXL-related memory expansion, and 800G/1.6T networks gradually mature, AI storage nodes can deliver higher aggregate bandwidth with fewer PCIe lanes, fewer drive bays, or fewer access-tier hops — directly valuable for long-context inference, KV cache tiering, real-time feature reads, and multi-tenant vector retrieval.

Longer term, the PCIe 7.0 specification is evolving toward 128 GT/s link rates, and future higher-generation SSDs will push the local flash tier closer to memory-tier access capability. For AI storage, the value of advanced SSD technology is not simply chasing peak benchmark scores, but shortening the time GPUs wait for data, improving recovery speed after cache hits, reducing the expansion cost of the hot data tier, and providing a more reliable media foundation for storage-disaggregated GPU nodes (few or no local drives, storage pooled at rack level), network context caches, and rack-level flash pools.

On media routes, the applicable positions of TLC and QLC should be distinguished by workload type. TLC better suits the training hot data tier, checkpoint tier, metadata tier, and mixed read/write tiers — its steady-state write capability and endurance better fit long-term heavy-load scenarios. QLC better suits read-intensive, capacity-first AI scenarios, such as the inference cache expansion tier, vector database primary data tier, RAG raw index tier, and read-only dataset tier.

Error diagnosis and early failure prediction for NVMe SSDs can be implemented through standardized management capabilities for device monitoring. The NVMe-MI specification requires unified interfaces for NVMe device discovery, monitoring, configuration, and update across in-band and out-of-band paths, facilitating predictive maintenance of large-scale nodes. Meanwhile, FDP (Flexible Data Placement) technology reduces garbage collection and write amplification by letting the host give SSDs finer-grained write placement hints, further improving media lifespan, performance consistency, and space efficiency.

In form factor and power, new interfaces such as E1.S and E3.S are replacing traditional drive-bay descriptions. SNIA's public documentation on EDSFF emphasizes better adaptability in flash density, power delivery, thermal design, and serviceability compared with traditional form factors, helping accommodate more high-performance flash in a single node. For per-drive power, the focus is mainly on lower write amplification, more rational airflow paths, and more stable thermal control curves.

Table 2.8 Applicable positions of TLC and QLC in AI storage

Media type	Better-suited workloads	Main advantages	Key risks to control
TLC	Training hot data, checkpoints, metadata tier, mixed read/write cache tier	Stronger steady-state writes, higher endurance, easier tail-latency control	Higher cost; prioritize for tiers with the highest performance-consistency requirements

QLC	Inference corpus, vector database main body, read-intensive datasets, capacity expansion tier	Higher capacity efficiency, lower unit cost	Sustained rewrites and high-frequency random writes expose steady-state performance and lifespan constraints sooner
-----	---	---	---

2.4.3.4 RDMA NICs

In AI storage servers, high-speed RDMA NICs are core components as important as CPUs and NVMe SSDs. As the aggregation capability of local NVMe pools keeps improving, 400G/800G-class RDMA network interfaces have become the baseline configuration direction for high-performance AI storage nodes. New-generation high-speed NICs place higher demands on PCIe 5.0 or PCIe 6.0 motherboards and CPU platforms, so network planning must be designed as one with the host platform.

GPU Direct Storage (GDS) is a storage I/O acceleration technology from NVIDIA. Its principle is to establish a direct DMA data path between GPU VRAM and storage devices, bypassing the bounce buffer in CPU memory, thereby reducing CPU load and alleviating system bandwidth bottlenecks. In local storage scenarios, GDS directly opens the DMA path between GPU VRAM and NVMe SSDs; combined with RDMA NICs, AI storage nodes further shorten the GPU data path and reduce CPU involvement at the storage node.

Before introducing EBOF, JBOF, and storage disaggregation architectures, the boundaries of these terms must be clarified. JBOF (Just a Bunch of Flash) usually refers to an external flash enclosure composed of multiple NVMe SSDs, providing shared block access via NVMe-oF. EBOF (Ethernet Bunch of Flash) can be understood as the JBOF form carried over Ethernet/RoCE. Storage disaggregation designs GPU compute nodes without local persistent data drives, with the compute side accessing remote flash pools or parallel file systems over high-speed networks. All three point toward resource pooling that "decouples compute from flash." EBOF/JBOF/diskless, parallel file systems, and distributed storage all rely heavily on high-speed RDMA networks to guarantee sustained bandwidth and low latency for multi-client concurrent access. In inference systems, both remote network storage tiering of the KV cache and context sharing and migration depend on high-speed network support.

2.4.3.5 DPU NICs

The DPU-based high-speed inference storage architecture is exemplified by NVIDIA's Inference Context Memory Storage Platform. Driven by the NVIDIA BlueField-4 DPU and aimed at long-context, multi-turn, and agentic AI inference, the platform targets expanding GPU context capacity at rack level and even across node clusters, and supports high-bandwidth KV cache sharing. The network foundation uses NVIDIA Spectrum-X Ethernet, providing RDMA-based high-speed KV cache access; at the software layer, it coordinates with NVIDIA Dynamo to implement KV cache seamless flow among HBM, CPU memory, and off-chip storage via the NIXL asynchronous transfer library. The emergence of this platform shows the

DPU's role in AI storage evolving further from traditional protocol offload toward context storage, cache routing, and efficient data sharing layers.

2.4.3.6 Storage Drive Expansion

Storage drive expansion can be further understood beyond "enclosure expansion" as "flash resource pool expansion." JBOF places large numbers of NVMe SSDs in an independent flash enclosure, exposing shared access via NVMe-oF; EBOF emphasizes that the enclosure connects to the compute network mainly over Ethernet and RoCE. The core value of both is decoupling flash resources from compute nodes and managing them as an independent resource pool, providing composable flash access services to upper-layer nodes via NVMe-oF over RoCE or NVMe-oF over TCP. This architecture enables finer-grained fault-domain isolation during expansion, maintenance, firmware upgrades, and media replacement, and better suits managing flash as resource pools in large equipment halls.

The diskless architecture advances this pooling idea further: GPU compute nodes no longer have local drives for business data persistence; instead, models, datasets, vector indexes, checkpoints, and context caches are placed in remote flash pools, parallel file systems, or converged object/file storage. Some implementations still keep a small amount of local NVMe for system boot, logs, or temporary caches, but the core data path relies on high-speed network access to remote storage. This approach especially suits AI clusters that need elastic expansion, rapid replacement of stateless compute nodes, unified flash capacity management, and fault-domain partitioning by role (training, inference, retrieval). Combined, JBOF/EBOF and diskless architectures can form a multi-tier AI storage system from local cache to remote flash pools.

Chapter 3 Network and High-Speed Interconnect Layer

This report places the physical layer and connectors of high-speed interconnect (media, termination, cables, PCBs) in Chapter 2 (IT Facility Layer), and places interconnect protocols and network paradigms (the protocols, topologies, and software stacks of Scale-Up/Out/Across) in this chapter. Section 2.3 focuses on the hardware forms and deployment of switching systems, while Section 3.2 in this chapter focuses on Scale-Out protocols and tuning — the two complement each other.

The interconnect system of an AIDC cluster can be divided into three levels: Scale-Up vertical scaling, Scale-Out horizontal scaling, and Scale-Across long-distance scaling. The Scale-Up architecture focuses on resource integration within a single logical compute node, achieving unified addressing among multiple GPUs and breaking through the bandwidth and latency bottlenecks of traditional buses. Scale-Out relies on lossless networks for inter-node communication, emphasizing cluster-scale expansion. The Scale-Across architecture further extends Scale-Out clusters distributed in different regions, forming a cross-regional compute pool, commonly used in hyperscale AIDCs of over ten thousand GPUs. The three interconnect networks form a complementary hybrid compute architecture, greatly improving node cluster compute density and parallel computing efficiency.

3.1 Scale-Up Technology Analysis

3.1.1 Research Background and Significance of Scale-Up

Current large models impose stringent requirements on the computing performance, storage capacity, and data interaction capability of the underlying compute infrastructure. Compute demand is not merely stacking more chips, but a rigid requirement for low latency, strong consistency, and high compute density. The scalability of GPU computing power and storage space directly determines the efficiency of model training, inference, and multi-node collaborative computing. Traditional interconnect architectures generally suffer from prominent bandwidth bottlenecks, high communication latency, and limited scalability, making them ill-suited to the high-frequency interaction of massive data in model training and inference scenarios. Against this industry backdrop, Scale-Up vertical scaling has become a key technical path for improving compute power.

In the early development of AIDC cluster infrastructure, vertical scaling architectures generally used the PCIe bus for interconnection. With continuous iteration of compute and interconnect technologies, traditional architectures have gradually evolved toward high-performance supernode architectures. The mainstream Scale-Up interconnect protocols and technology systems in current supernode scenarios mainly include NVLink, OISA, UALink, SUE, ESUN, CLINK, and Co-Fabric.

3.1.2 Inter-Card High-Speed Interconnect Protocols and Technical Standards

3.1.2.1 NVLink Protocol

NVLink is a high-speed interconnect technology proposed by NVIDIA in 2014. Unlike traditional architectures where GPUs must forward data through the CPU, NVLink uses a GPU direct-connection design; paired with the dedicated NVSwitch chip, multiple GPUs can exchange data directly, greatly improving communication efficiency. NVLink supports memory coherence and unified memory addressing, consolidating the local VRAM of multiple GPUs into one large unified memory pool. The technology also integrates hardware-accelerated protocols that optimize collective communication operations such as All-Reduce, reducing communication latency in multi-GPU coordination and improving data interaction efficiency and stability.

As the benchmark Scale-Up solution, NVLink has undergone multiple iterations with continuously leaping performance. From the first-generation NVLink 1.0 at 160 GB/s bidirectional to the current commercial NVLink 5.0 at 1.8 TB/s bidirectional — more than a tenfold improvement — the next-generation NVLink 6.0 is expected to reach commercial availability in 2026, further raising bidirectional bandwidth to 3.6 TB/s. At the system deployment level, the NVIDIA Blackwell Ultra NVL72 rack achieves all-to-all topology among 72 GPUs, with total in-rack communication bandwidth of 130 TB/s.

Although NVLink excels in performance, it has obvious shortcomings. The core problems are a highly closed ecosystem, strong vendor lock-in, and poor compatibility. Although the NVLink Fusion open interconnect solution was introduced in 2025, vendors face strict licensing thresholds and it has not yet been deployed at scale. Once adopted, users are deeply locked into NVIDIA's hardware system. Second, costs are high: NVLink requires dedicated NVSwitch chips, proprietary backplanes, and power and cooling design, imposing strict requirements on overall rack architecture, with deployment and O&M costs higher than other general-purpose solutions. These limitations create high implementation and adaptation thresholds in construction.

3.1.2.2 Omni-directional Intelligent Sensing Express Architecture (OISA)

China Mobile proposed the OISA (Omni-directional Intelligent Sensing Express Architecture) technology system for AIDC Scale-Up scenarios in 2024 and released the updated OISA 2.0 protocol in 2025. The technology supports efficient interconnection of 1024 GPUs within a supernode, with inter-card bandwidth breaking through the TB/s level and latency reduced to hundreds of nanoseconds, providing strong support for data-intensive AI applications such as large-model training and inference.

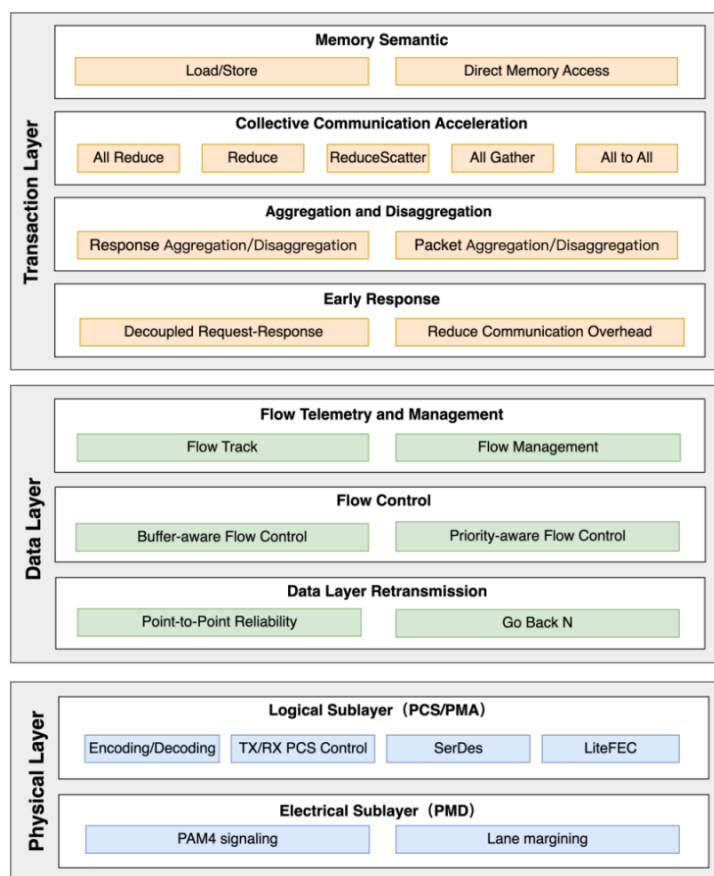


Figure 3.1 OISA protocol architecture

As shown in Figure 3.1, the OISA protocol stack consists of a transaction layer (TL), a data layer (DL), and a physical layer (PL). The transaction layer, at the top of the stack, interacts with the GPU's on-chip network, receiving transaction-related signals from the GPU, encapsulating them into transaction layer packets (TLPs), and forwarding them to the data layer. The data layer performs flow control and intelligent traffic awareness to ensure transmission stability. The physical layer uses an Ethernet-based implementation, leveraging mature Ethernet physical-layer technologies (such as SerDes) to support high-speed transmission between GPU devices.

The OISA 2.0 protocol has several features: first, native memory semantics, allowing AI chips to access data across nodes without barriers; second, innovative TLP packet reconstruction technology that aggregates small compute transactions to improve bandwidth utilization; third, intelligent awareness, monitoring the interconnect status of all GPUs and switch chips in real time to provide data support for dynamic routing, congestion control, and intelligent O&M; fourth, hardware acceleration of collective communication, offloading complex operations such as collective communication in AI training to the switch chip through dedicated hardware engines, bringing significant performance improvements to scenarios such as large-model training.

By building a unified, open technology platform, OISA provides a new domestic solution for supernode inter-card interconnect and is a representative Scale-Up technology in China. It will effectively stimulate

industrial chain innovation, drive technical upgrades and iterative innovation of related products, and build a new AIDC industry ecosystem of positive interaction and continuous evolution.

3.1.2.3 SUE Technology

To achieve efficient transmission between XPU in large-scale scenarios, Broadcom released the SUE (Scale-up Ethernet) protocol in April 2025; Figure 3.2 shows the SUE protocol stack architecture. In addition to the standard mode, SUE provides a streamlined mode called SUE Lite, which simplifies the packet format and header length, keeping only the headers needed for data forwarding to reduce hardware overhead and transmission latency.

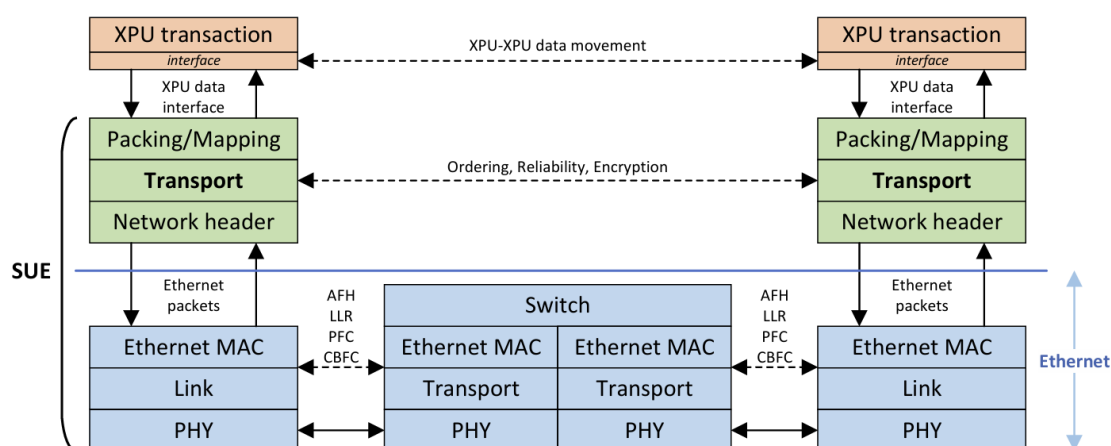


Figure 3.2 SUE protocol architecture

In transmission reliability design, SUE supports credit-based flow control (CBFC), priority-based flow control (PFC), and link-layer retransmission (LLR), building an end-to-end lossless transmission system. Working together, the three significantly reduce latency jitter and packet loss, meeting the stringent transmission stability requirements of large-model training and inference scenarios.

In June 2025, the Open Compute Project (OCP) announced the SUE-T (Scale-up Ethernet Transport) working group to improve the Scale-up Ethernet interconnect protocol system. SUE-T will inherit and continue advancing the transport-layer core work from parts of the original SUE effort, focusing on end-side transport-layer protocols and achieving decoupling of the transport and switching layers.

At the hardware ecosystem level, SUE adapts to Broadcom's Tomahawk Ultra and Tomahawk 6 series switch chips, leveraging their forwarding performance to provide ample bandwidth and low latency for large-scale Scale-Up clusters. The deep coordination of SUE with switch chips makes Ethernet-based Scale-Up architecture the core support for AI supernodes moving from technical solution to commercial deployment.

3.1.2.4 UALink Protocol

In October 2024, to break the monopoly of proprietary protocols and enable efficient collaboration between accelerators and switch chips from different vendors, nine major companies — AMD, AWS,

Astera Labs, Cisco, Google, HP, Intel, Meta, and Microsoft — announced the UALink (Ultra Accelerator Link) Consortium. The UALink protocol, the consortium's core technology, officially released version 1.0 in April 2025 and updated to version 2.0 in April 2026, and is open to consortium members. The protocol architecture is shown in Figure 3.3: from top to bottom, the protocol layer, transaction layer, data link layer, and physical layer.

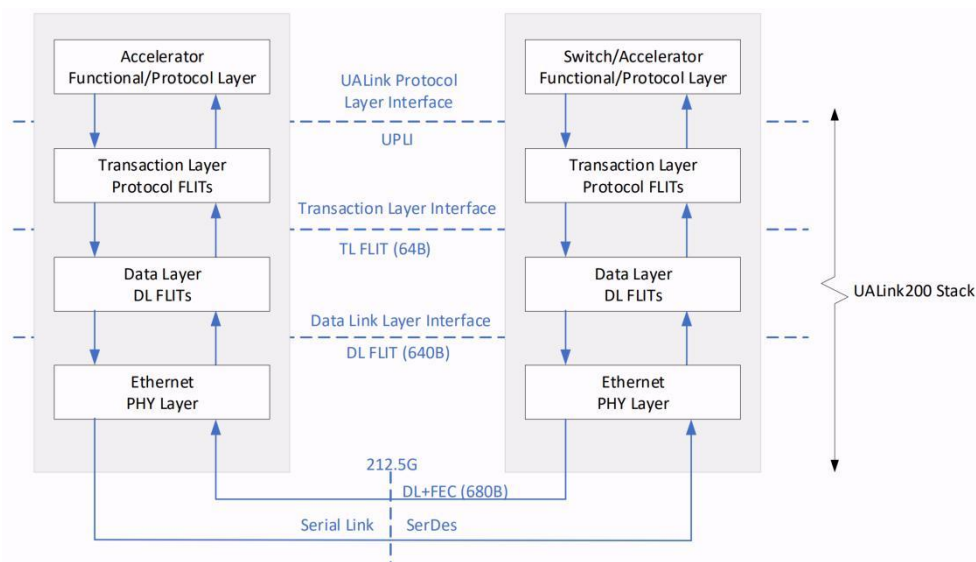


Figure 3.3 UALink protocol architecture

The UALink protocol layer, also called UPLI (UALink Protocol Level Interface), is the interaction interface between the protocol and the accelerator. Its transaction layer splits protocol-layer instructions into standard transaction-layer transfer units, TL Flits (64 bytes), and passes them to the data link layer. The data link layer adds a CRC and frame header to each Flit, encapsulating it into a data-link-layer transfer unit, the DL Flit (640 bytes), and passes it to the physical layer; link retransmission and credit flow control mechanisms also ensure transmission reliability. Its physical layer is based on IEEE 802.3dj with some modifications, allowing the protocol to directly reuse mature Ethernet SerDes, cable, and other supply chain resources, reducing hardware mass-production and deployment costs.

UALink integrates the UALinkSec security mechanism, supporting end-to-end data encryption and trusted execution environment adaptation, enabling tenant traffic isolation and meeting the security compliance requirements of cloud-native multi-tenant scenarios. However, UALink does not support coherence protocols; data consistency must be designed and implemented by vendors themselves, for example through software or other means. In industry adoption, the first products supporting the UALink specification are expected to reach commercial availability in 2026-2027, becoming a core support for opening up AI compute infrastructure.

3.1.2.5 Other Inter-Card Interconnect Protocols

PCIe technology, with its mature supply chain, controllable cost, and strong ecosystem compatibility, became the quickly adopted technical route in the early development of AIDC infrastructure. Its typical

architecture uses full-mesh topology within a single node, with inter-node multi-machine communication via PCIe switches, addressing the core needs of trillion-parameter large models for large VRAM space and low-latency communication. However, limited by its protocol stack architecture and transmission characteristics, nodes cannot achieve native point-to-point direct connection, affecting bandwidth utilization.

IEIT deeply modified the PCIe protocol to develop the Co-Fabric interconnect protocol, building 32-GPU, 64-GPU, 128-GPU, and larger Scale-Up AIDC systems. Based on general bus protocols, Co-Fabric adds an agent layer and an adapter layer, breaking the extension limits of native bus protocols and supporting fully switched non-blocking interconnection of over a hundred GPUs. Through local virtual devices, it transparently maps remote GPU VRAM, building a globally unified address space. Through efficient data forwarding within the switching domain, communication latency is compressed to the hundred-nanosecond level. Taking a 64-GPU supernode as an example, the token generation speed when running the DeepSeek R1 model reaches 8.9 ms, and per-GPU throughput reaches 631.4 tokens/s when running DeepSeek V3.

In September 2025, led by ODCC with the China Academy of Information and Communications Technology and Tencent taking the lead, the ETH-X interconnect protocol technical report was released. At the transaction layer, the specification defines the implementation of PAXI (Peer-to-Peer AXI), extending AXI bus transactions to remote nodes over Ethernet; the data link layer defines the frame structure and retransmission mechanism; and the physical layer is compatible with the PMD and PCS definitions of standard Ethernet IEEE 802.3. ETH-X aims to provide Scale-Up interconnect capability based on the Ethernet ecosystem.

At the OCP Summit in October 2025, 12 tech giants including AMD, Arista, and Broadcom jointly launched the ESUN (Ethernet for Scale-Up Networking) consortium, focusing on standardizing lower network layers in accelerators and switches. Coordinating with SUE-T, which focuses on the transport layer, they jointly promote the scaled application of open Ethernet in supernodes, breaking the monopoly of proprietary protocols.

In November 2025, under the guidance of China's Ministry of Industry and Information Technology, the China Electronics Standardization Institute launched the Compute Link (CLink) joint initiative. Through standardization-led computing technology innovation, it jointly promotes the continuous iteration of the compute interconnect bus protocol standards family. CLink covers core content such as Scale-Up overall architecture design, communication semantics, and flow control, providing unified technical specifications for large-scale AIDC clusters and improving scaled deployment efficiency.

In addition, there are the Scale-Up open ecosystem ALS (ALink System) initiated by Alibaba Cloud; ByteDance's Ethernet-based EthLink solution supporting Load/Store semantics and RDMA semantics; Huawei's published interconnect protocol for supernodes — the UB (Unified Bus) 2.0 technical specification, enabling peer-to-peer direct communication among heterogeneous processing units such as CPUs, GPUs, NPUs, DPUs, and switches; the high-throughput Ethernet ETH+ initiated by Alibaba Cloud

and the Institute of Computing Technology, whose protocol standards support open decoupling; Tanwei Xinlian's self-developed ACCLink, with related products expected to complete technical verification by the end of 2026; ZTE's self-developed OLink technology, compatible with RDMA standards; and Hygon's fully open interconnect bus protocol (HSL), which reduces coordination latency and adaptation costs between GPUs and CPUs.

3.1.3 Future Development Trends and Application Potential

Scale-Up technology will continue to optimize along three directions: the physical layer, the protocol stack, and boundary extension. As the underlying core support, the 448 Gbps single-lane physical-layer technology based on the OIF CEI-448G framework completed its standards framework release in November 2025. This technology doubles link bandwidth without changing port density, greatly improving per-rack compute density, but it also imposes extremely high signal-integrity requirements on PCB materials and connectors, requiring low-loss materials and optoelectronic fusion solutions for stable deployment. Carrying the physical layer's evolution needs, optical interconnect will become the core solution for breaking through the limits of electrical interconnect; technologies such as NPO and CPO will gradually reach commercial deployment — see Section 2.2 for details.

In-network computing in the protocol stack has achieved initial deployment in leading vendors' large-model training clusters, implementing standardized hardware offload of collective communication operators and reducing end-side resource consumption. However, it also faces problems of limited operator programmability, poor cross-vendor compatibility, and rising chip design complexity and power consumption, and coordination mismatch between flow control and compute scheduling may trigger new congestion risks. The technology will keep iterating with ecosystem co-building.

Advancing in parallel is the extension of architectural boundaries: Scale-Up will expand from a single rack to a two-tier network across racks, breaking through per-rack compute ceilings and simplifying the networking complexity of large-scale clusters. However, increased hops bring higher latency, and issues such as multi-path load balancing, adaptive routing, packet reordering, broadcast storm risks from enlarged broadcast domains, and significantly expanded single-point-of-failure impact scope require balancing cluster scale against interconnect performance. Overall, Scale-Up technology will keep evolving toward higher bandwidth, lower latency, and better adaptation.

3.2 Scale-Out Network Technology

3.2.1 Scale-Out Network Requirements and Goals

3.2.1.1 Scale-Out Requirements

An AIDC's Scale-Out network must simultaneously support cross-node communication for large-scale distributed training and inference tasks. Core requirements include ultra-high bandwidth, extremely low latency, and lossless transmission. Training tasks mainly generate Scale-Out traffic from gradient All-Reduce in data parallelism and cross-node transfers in pipeline parallelism (tightly coupled communication such as tensor parallelism is mainly carried by the Scale-Up domain); inference tasks, with the rise of PD disaggregation, MoE expert parallelism (All-to-All), and cross-node KV cache sharing, impose increasingly strong bandwidth and latency requirements on Scale-Out.

1. High bandwidth requirement. When network bandwidth is insufficient, waiting time caused by data-parallel communication can account for 55-85% of model training iteration time, becoming the main performance bottleneck. To complete training synchronization within milliseconds or microseconds, each compute node needs extremely high network egress bandwidth — for example, 400 Gbps, 800 Gbps, or even 1.6 Tbps.
2. Extremely low latency requirement. Collective communication often requires all worker nodes to wait for the slowest communication to complete before entering the next iteration, so network tail latency directly determines the progress of the entire training job. In addition, inference PD disaggregation and KV cache transfer are highly sensitive to end-to-end latency and jitter. The Scale-Out network must therefore provide microsecond-level end-to-end latency with minimal jitter.
3. Lossless transmission requirement. Packet loss triggers excessive retransmission, causing network performance to degrade sharply. The Scale-Out network therefore implements lossless transmission of workload traffic based on the high-performance RDMA communication protocol.

3.2.1.2 Goals of the Scale-Out Network

The goal of the Scale-Out network is to maximize the overall communication performance of distributed training and minimize XPU (GPU, TPU, LPU, NPU, etc.) idle waiting time caused by the network. Its design targets include 1) supporting bandwidth up to 100 GB/s, and 2) achieving end-to-end round-trip latency below 10 microseconds over distances of several hundred meters.

3.2.2 Network Operating System SONiC

SONiC (Software for Open Networking in the Cloud) is an open-source, Linux-based network operating system initiated and open-sourced by Microsoft in 2015, now a sub-project in the OCP networking domain. SONiC aims to provide highly scalable network operating solutions for data centers through

software-hardware decoupling. It supports multiple ASIC chips and CPU architectures, breaking the hardware-binding limitations of traditional operating systems. SONiC has strong ecosystem support spanning chip vendors such as Broadcom and Mellanox and equipment vendors such as H3C, Dell, and Arista, and has been deployed in production by hundreds of large organizations including Baidu, Tencent, and LinkedIn.

SONiC's core architecture is built on the Switch Abstraction Interface (SAI). Using Docker container technology, network function modules (such as BGP) are packaged as independent processes, achieving complete decoupling among software, hardware, and protocol modules. This allows operators to share the same software stack across hardware from multiple vendors and quickly add new components or third-party software.

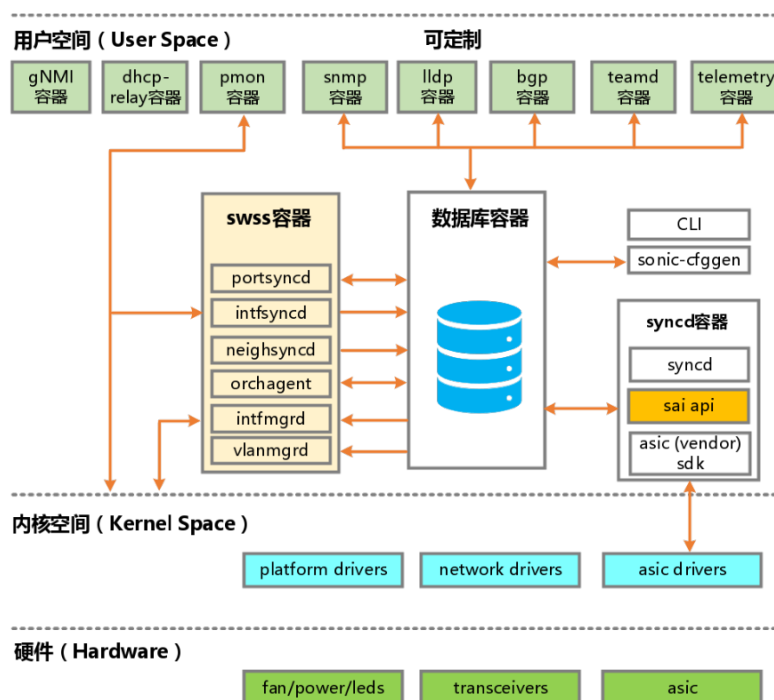


Figure 3.4 SONiC system architecture

In Scale-Out networks, SONiC's standardized interfaces and containerized architecture fit the needs of building large-scale AIDC clusters: it handles high-speed interconnection and dynamic configuration across multiple nodes, supports high-throughput scenarios such as RDMA, and enables linear expansion of hardware nodes without changing software. As the commercial ecosystem matures, SONiC is expected to become the mainstream network operating system for AIDCs, with a position similar to Linux in the server domain.

3.2.3 Scale-Out Protocols and Software Stack

3.2.3.1 Communication Protocols

AIDC Scale-Out networks mainly use communication protocols such as InfiniBand, RoCEv2 (RDMA over Converged Ethernet v2), and Ultra Ethernet (UE), along with a variety of application protocols.

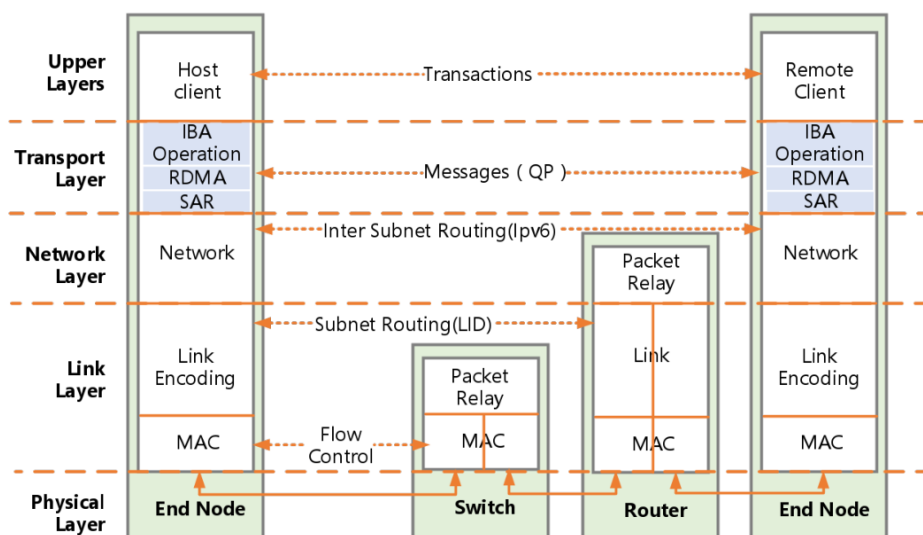


Figure 3.5 InfiniBand protocol stack

1. **InfiniBand RDMA.** InfiniBand RDMA is a network interconnect protocol designed for high-performance computing and AI clusters. It allows one computer's NIC to directly read from and write to another computer's memory, bypassing the operating system kernel and CPU protocol processing, thereby reducing latency. The InfiniBand link layer uses credit-based flow control to achieve lossless transmission and guarantees packet ordering within a single virtual lane path. RDMA operations are controlled via the InfiniBand Verbs interface and executed by RDMA NIC hardware, providing low-latency, high-bandwidth data transfer. In addition, most key data transfer operations in the protocol stack are offloaded to hardware for acceleration. Although facing competition from Ethernet solutions such as RoCEv2 and UEC, InfiniBand RDMA still holds an advantage in AI scenarios with extreme performance requirements.

2. **RoCEv2.** In AIDC Scale-Out networks, RoCEv2 is currently the main Ethernet standard. RoCEv2 is an implementation of RDMA that runs over Ethernet and supports IP routing across Layer 3; the transport layer uses UDP, and the application layer supports collective communication libraries and RDMA APIs. RoCEv2's greatest advantage over RoCEv1 is routability across subnets, enabling broader interconnect expansion in Scale-Out environments.

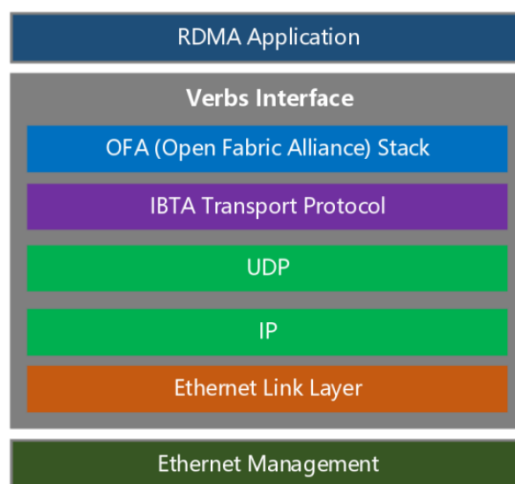


Figure 3.6 RoCEv2 protocol stack

RoCEv2 is based on the connectionless UDP transport protocol and lacks TCP-like reliability and congestion control, so it usually needs to be combined with protocols such as PFC, ECN, and DCQCN to ensure reliable communication and high-performance transmission in large-scale XPU clusters. RoCEv2 fits data centers' existing IP/Ethernet protocol stacks, avoiding the extra cost and complexity of deploying a dedicated InfiniBand network.

3. UE. Ultra Ethernet (UE) defines a set of Ethernet-based standards and specifications to meet the high-bandwidth, low-latency communication needs of HPC and AI applications. In particular, UE is the Scale-Out network recommended by OCP Open Systems for AI.

The UE software layer defines APIs and specifications, with its core supporting the libfabric (v2.0) interface.

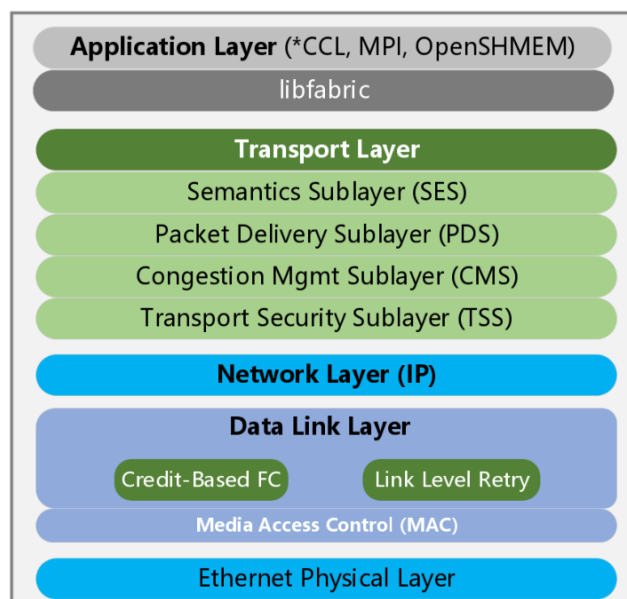


Figure 3.7 UE protocol stack

As shown above, the transport layer is UE's innovation core. It supports RDMA and aims at hyperscale expansion to millions of endpoints, consisting of the semantic sublayer, packet delivery sublayer, congestion management sublayer, and others.

- The semantic sublayer (SES) maps libfabric APIs to network operations (such as RDMA read/write), supporting deferred send and message aggregation protocols to optimize offload efficiency for AI collective communication libraries;
- The packet delivery sublayer (PDS) handles reliability and ordering, implementing connection management through ephemeral packet delivery contexts (PDCs) and supporting zero-RTT context establishment — the first arriving packet can establish the reliability context;
- The congestion management sublayer (CMS) provides end-to-end congestion control, supporting sender-side and receiver-side control. A key CMS feature is packet spraying: by varying the entropy value (EV), packets of a single flow are distributed across all available paths, eliminating network hot spots.

The UE network layer uses standard IPv4 and IPv6 protocols to ensure compatibility with existing data center routing architectures and management tools. Packet Trimming (PT) is a key optional network-layer feature: when a switch buffer becomes congested, PT trims the packet payload and forwards only the header — an extremely fast explicit loss signal that lets the receiver immediately sense the loss and request retransmission, greatly reducing tail latency.

The UE link layer protocol remains compatible with IEEE 802.3 while introducing optional performance enhancements, including Link Layer Retry (LLR) and Credit-Based Flow Control (CBFC). The Link Layer Discovery Protocol (LLDP) is used to negotiate with communication peers to enable LLR and CBFC. A comparison of InfiniBand, RoCEv2, and UE is shown in the table below.

Table 3.1 Comparison of InfiniBand, RoCEv2, and UE

	InfiniBand	RoCEv2	UE
Current stage	Mature / high-end market monopoly	Mature / large-scale deployment	Introduction / standards development
Core advantages	Extreme performance, full-stack optimization	Openness, cost, ecosystem, fast iteration	Open architecture, future-oriented design
Core disadvantages	Closed, expensive, slow iteration	O&M complexity from ecosystem openness	Not yet deployed at scale; uncertainty remains
Future development	Market may gradually be taken by RoCE	De facto market mainstream for a long time	Important future investment and technology bet

communication functions. In addition, the RDMA API can be used to directly operate the RDMA network to customize libraries, storage, and network optimization functions.

3.2.3.2 Network Congestion Control

1. Adaptive routing. Adaptive routing is an end-to-end fine-grained load balancing protocol implemented through close cooperation between switches and NICs, aiming to solve the elephant-flow collisions and network congestion common in AI training. Scale-Out adaptive routing supports per-packet multi-path transmission, forwarding packets to available ports simultaneously based on egress queue availability, achieving packet spraying and receiver-side reordering, reducing link congestion and greatly improving effective network bandwidth utilization.

2. Closed-loop control. For many-to-one (Incast) congestion caused by multiple senders simultaneously sending data to a single receiver, the Scale-Out network cooperates with congestion control mechanisms such as PFC, ECN, DCQCN, and packet trimming, controlling the data injection rate through feedback signals to achieve closed-loop congestion control in Incast scenarios.

3. Network telemetry. AIDC Scale-Out routing and control protocols rely on measurement of network state. Switches detect congestion hot spots in real time through telemetry and feed the information back to NICs. NICs execute congestion control with microsecond-level latency, finely regulating the traffic sources causing congestion to avoid network congestion. Real-time telemetry also ensures performance isolation in multi-tenant environments.

3.2.4 Performance Evaluation and Tuning

3.2.4.1 Scale-Out Network Performance Metrics

Scale-Out network performance metrics include: 1) bandwidth and latency in RDMA bisection scenarios; 2) large-scale All-Reduce and All-to-All bandwidth; 3) effective bandwidth utilization; 4) scaling efficiency of distributed training jobs — the improvement ratio in job completion time when the number of XPU cards expands from N to 2N; 5) network resilience — network performance when some links are down; and 6) performance before and after multi-tenant network virtualization.

3.2.4.2 Performance Tuning Technologies

Scale-Out network performance tuning technologies include traffic conditioning, topology-aware routing, end-to-end coordination, and RDMA fault recovery. Traffic tuning dynamically adjusts ECN thresholds, PFC thresholds, and so on for different traffic patterns; topology-aware routing uses routing mechanisms based on job communication matrices and topology information (such as multi-path traffic engineering) to improve load balancing protocols; end-to-end coordination performs full-stack joint tuning across software communication libraries, NIC drivers, and switch OSes; and RDMA fault recovery protocols provide fault detection and recovery capabilities, reducing the impact of failures on compute tasks through backup multi-paths.

3.2.5 Industry Frontier Cases and Trends

3.2.5.1 AIDC Scale-Out Network Practices

1. InfiniBand-based Scale-Out networks. NVIDIA DGX SuperPOD uses InfiniBand to build large-scale distributed AI clusters and implements Scale-Out networking through modular design, adopting adaptive routing, multi-path traffic sharing, and hardware-based network self-healing and error detection to guarantee high-performance RDMA communication.
2. RoCEv2-based Scale-Out networks. Azure and Meta use RoCEv2 + customized Ethernet to achieve large-scale AI Scale-Out networks with performance close to InfiniBand, but more open and lower in cost.
3. UEC-based Scale-Out networks. Broadcom's Scale-Out network solution is based on UEC, using enhanced RDMA, multi-path, and programmable congestion control to upgrade Ethernet into a high-performance interconnect system for million-node AI clusters.
4. Heterogeneous interconnect and customized solutions. Google's OCS Scale-Out network essentially replaces the traditional multi-tier electrical switching network with "reconfigurable optical paths + SDN control," achieving low-latency, high-bandwidth, and low-cost interconnection for AI clusters. OCS is purely physical-layer switching and performs no protocol parsing or electrical signal processing on packets, so OCS is transparent to encapsulation protocols such as RoCE.
5. VNET Scale-Out case. VNET's AIDC adopts a two-tier non-blocking Spine-Leaf architecture to build a Scale-Out network with a 1:1 oversubscription-free ratio. The entire network is based on RoCEv2 lossless Ethernet technology. Thousand-GPU commercial clusters have been deployed in national computing hubs and bases such as Ulanqab, Zhangjiakou, the Yangtze River Delta, and the Guangdong-Hong Kong-Macao Greater Bay Area, with smooth expansion to ten-thousand-GPU scale; end-to-end communication latency is stably controlled within 5 μ s, and peak network utilization exceeds 90%.
6. Purple Mountain Laboratories Scale-Out case. Purple Mountain Laboratories built an ultra-wide lossless AIDC network based on self-developed 51.2T Ethernet white-box switches running the UniNOS network device operating system, using a two-tier Spine-Leaf architecture to build an end-to-end 400G Scale-Out ultra-

wide RoCEv2 network, deploying PFC and ECN for lossless transmission, and adopting adaptive routing with per-packet load sharing, achieving link bandwidth utilization of no less than 95%.

3.2.5.2 Future Trend Outlook

InfiniBand and RoCEv2 remain the typical high-performance Scale-Out network protocols today, while protocols such as UE and MRC (Multipath Reliable Connection) are iterating rapidly and may become the mainstream in the future. RDMA-related protocols deployed on OCS are also an important evolution direction for AIDC Scale-Out networks. In addition, Scale-Up and Scale-Out protocols show a trend of coordinated evolution — for example, ESUN and UEC are collaborating to align open standards.

3.3 Scale-Across Network Technology

3.3.1 AIDC Interconnection Requirements

3.3.1.1 Compute Expansion

Scale-Across networks extend AI fabrics beyond a single AIDC, with connection distances spanning tens of kilometers. Through Scale-Across networks, operators can interconnect multiple AI data centers, forming hyperscale AI infrastructure across multiple locations.

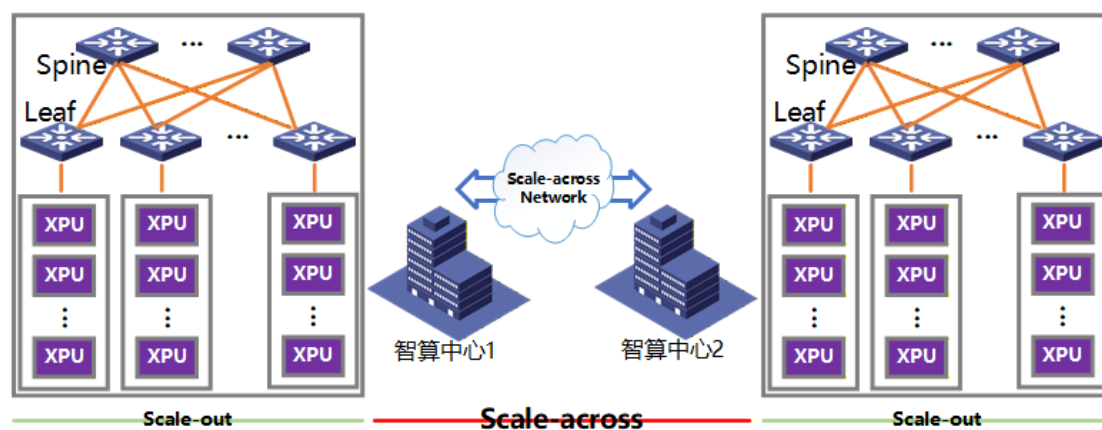


Figure 3.8 Scale-Across network schematic

Scale-Across compute expansion connects the compute of multiple data centers distributed across different geographic locations through Scale-Across networks, making them logically appear as a unified, schedulable hyperscale computing fabric. The demand mainly stems from the physical and engineering limits a single data center faces when training trillion-parameter-class large models. In addition, as AI business shifts toward inference, compute needs to be distributed in metropolitan areas closer to users; Scale-Across technology can organize inference nodes across multiple metropolitan areas into a unified resource pool, enabling distance-prioritized scheduling and dynamic load allocation.

3.3.1.2 Physical Constraints

Beyond compute expansion demand, the need to connect multiple data centers via Scale-Across networks also comes from the physical constraints of a single data center, including power, cooling, and footprint limits.

1. Power supply constraints. Per-rack power is evolving from ~150 kW today toward 200-500 kW (and even higher) (see Section 1.3.2.1 for details). A single AIDC can hardly support the enormous power and infrastructure consumption required by millions of XPU.
2. Cooling challenges. Traditional air cooling can no longer meet the heat dissipation needs of ultra-high-power racks. Liquid cooling systems are required, placing higher demands on single equipment halls in terms of physical space and engineering implementation.
3. Footprint constraints. The physical space of a single building or campus is limited, making it difficult to keep expanding to accommodate million-XPU clusters.

3.3.2 Lossless RDMA

3.3.2.1 Distance-Adaptive Load Balancing

Compared with single-data-center networks, the core challenge of Scale-Across networks is the significant difference in link distances, causing obvious latency differences among RDMA paths. Meanwhile, load easily concentrates on some links, and traditional static routing can hardly make full use of network resources. Efficient dynamic load balancing mechanisms are therefore needed to improve performance and stability.

The Distance-Adaptive Load Balancing (DALB) protocol significantly improves the performance and reliability of hyperscale Scale-Across networks through dynamic path weight computation and adaptive traffic allocation based on distance and load. The DALB protocol adaptively allocates traffic according to path distance and current network load, and by dynamically monitoring network state and periodically adjusting weights, hot links are balanced and multi-path resources fully utilized, enhancing the stability of InfiniBand RDMA, RoCEv2, and UE RDMA protocols in Scale-Across networks. When a link becomes abnormal, the DALB protocol can automatically lower its weight or switch traffic to other paths, ensuring communication continuity and overall network reliability.

3.3.2.2 End-to-End Congestion Control

Scale-Across networks involve communication across racks, Pods, or data centers, with complex topologies and significant differences in link latency and bandwidth. End-to-end congestion control (E2E CC) protocols monitor network congestion between communicating parties and adaptively adjust the sending rate to avoid link overload. These protocols typically rely on Explicit Congestion Notification (ECN):

when a network device detects excessive queue length or potential congestion, it marks packets with ECN, and the sender accordingly reduces its sending rate, preventing packet loss and latency accumulation.

In large-scale, heterogeneous, multi-path network environments, end-to-end congestion control can be combined with multi-path traffic engineering (MPTE), spreading traffic across multiple links and dynamically adjusting traffic allocation ratios according to each path's congestion state, so that less-congested paths carry more traffic. By combining end-to-end congestion control and ECN notification with multi-path traffic scheduling, Scale-Across networks can sense and respond to congestion in advance, achieving lossless RDMA transmission.

3.3.3 Wide-Area Deterministic Interconnection

3.3.3.1 Precise Latency Management

Latency management in wide-area Scale-Across networks is achieved through mechanisms such as network resource reservation, multi-path traffic allocation, load balancing weight optimization, fast congestion notification, and elastic bandwidth allocation, relying on close cooperation between compute and network scheduling. The network scheduler reserves bandwidth resources for each communication path in advance, and combines multi-path traffic engineering (MPTE) with backup path strategies for fine-grained traffic control. When a link is about to become congested, network nodes can identify the affected link and notify the sending node, which dynamically adjusts its sending rate accordingly, suppressing latency jitter while maintaining high throughput and end-to-end lossless performance. In addition, elastic bandwidth and multi-tenant scheduling policies allow network resources to grow or shrink dynamically with compute task demand.

3.3.3.2 Network State Awareness

In Scale-Across networks, network state awareness (such as in-band telemetry) is the core performance assurance mechanism for hyperscale distributed computing systems. The protocol incorporates the physical distance of Scale-Across links and the resulting propagation delay into the flow control model, allowing each link's data injection rate to be optimized according to its propagation characteristics, thereby minimizing tail latency. It can also automatically identify network topology and intelligently distribute traffic among multiple feasible paths. In-band telemetry provides sub-microsecond real-time monitoring, feeding congestion signals directly back to senders so NICs can quickly complete rate adjustments, and providing end-to-end health monitoring and troubleshooting capabilities.

3.3.4 Industry Frontier Solutions and Trends

3.3.4.1 Long-Distance Cross-AIDC AI Training Solutions

As shown in the figure below, the industry deploys XPU in multiple equipment halls and interconnects them via DCI equipment, forming a Scale-Across network for cross-DC collaborative training.

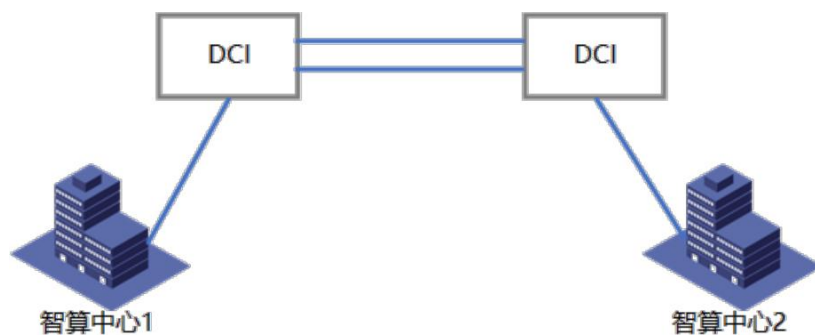


Figure 3.9 Long-distance AIDC interconnection schematic

Long-distance cross-DC collaborative training has been deployed in North America; limited by technology, current distances are all within 100 km. The challenges and solutions for 100 km long-distance cross-hall training scenarios are shown in the table below.

Table 3.2 Challenges and solutions for long-distance cross-hall training

100 km scenario technical challenges	Description	Solutions
Long-distance congestion notification efficiency	For congestion on devices in the near-end AIDC, the total path of ECN marking - CNP congestion notification - sender rate reduction takes 1 ms	Near-end DCI equipment supports FastCNP, replying to CNPs by proxy and achieving sender rate reduction at the μ s level
Cache backlog at the long-distance congestion proxy	For congestion on devices in the far-end AIDC, CNP congestion notification to sender rate reduction takes 1 ms in total, during which the far end triggers PFC and the far-end DCI must wait 1 ms before releasing buffers; a single port needs 50 MB of buffer and overflow packet loss is likely	DCI equipment uses models with large HBM buffers, such as Broadcom's StrataDNX (Dune) series chips or router NP chips

3.3.4.2 Scale-Across Network Practices

1. NVIDIA Scale-Across case. NVIDIA introduced Spectrum-XGS Ethernet technology, optimized specifically for Scale-Across scenarios, enabling geographically distributed facilities to operate like a single "giga-scale" AI factory. In NCCL All-Reduce collective communication tests between data centers 10 km apart, Spectrum-XGS performed 1.9 times better than traditional Ethernet.

NVIDIA's Scale-Across solution supports distance-adaptive congestion control protocols, introducing algorithms that sense inter-site distance and dynamically optimize flow control to minimize the tail latency of long-distance transmission. The topology-aware load balancing in the solution automatically adjusts traffic distribution across Scale-Across fabrics, eliminating network hot spots and maximizing

throughput. Through control logic based on global telemetry, the solution ensures deterministic latency and jitter quality of service.

2. Broadcom Scale-Across case. In the Scale-Across network domain, Broadcom proposed a solution based on open Ethernet standards, aiming to integrate dispersed AI clusters into a hyperscale compute pool supporting millions of XPU nodes through metro- or regional-level long-distance interconnection.

Broadcom's solution combines switch-side dropped-packet congestion notification (DCN), priority flow control (PFC), and explicit congestion notification (ECN), together with end-to-end traffic visibility, to address the oversubscription challenge caused by Scale-Across link bandwidth typically being smaller than intra-cluster bandwidth. The solution uses deep-buffer technology at the Scale-Across level to absorb latency and micro-burst congestion from long-distance transmission, using GB-level buffering to ensure RDMA traffic maintains lossless performance over long distances.

3. VNET Scale-Across case.

VNET builds AIDC clusters in cities such as Beijing and Suzhou with a three-tier CLOS-like network architecture — a Hub-Spoke-Edge topology. Hubs are core hub AIDCs deployed in near-suburban areas with abundant power and accessible green electricity, carrying ten-thousand-GPU training and inference clusters. Spokes are medium-sized radial AIDCs serving as relay hubs for metro compute, evenly deployed in the core locations of each urban administrative district. Edges are edge AIDCs reaching down into urban core business districts and industrial parks, deployed close to users and data sources, providing sub-millisecond low-latency edge intelligence services. At the bottom layer, an all-optical switching network reconstructs the metro compute bus, achieving one-hop direct all-optical interconnection across all Hub-Spoke-Edge nodes, with optical-layer end-to-end transmission latency stably controlled within 0.5 ms, building a non-blocking AI fabric within the city.

VNET proposed the Universal Cross Connection (UCC) ultra-interconnect network protocol. Based on standard Ethernet and RDMA technologies, it deeply optimizes and lightweights network and transport protocols, and has a built-in endogenous network security mechanism based on a public/private key system, achieving trusted, manageable, and controllable metro connections and ensuring high-speed, secure cross-node compute scheduling.

4. Purple Mountain Laboratories Scale-Across case. Purple Mountain Laboratories built a multi-compute-center cross-wide-area interconnected RDMA network based on the CENI network. RDMA gateway devices are deployed at compute center egresses; the gateways use Purple Mountain Laboratories' self-developed 12.8T programmable white-box switches running the UniNOS network device operating system, interconnected via 400G network ports. The FastCNP fast congestion control technology is deployed on the gateways, achieving fast congestion control feedback and solving the end-side RDMA NIC flow control problem caused by wide-area long-distance transmission latency. Elephant-flow load balancing is also deployed on the gateways, implementing lossless splitting of RDMA service traffic based on the five-tuple plus QPID, with dynamic orchestration through SRv6 multi-path policies to achieve precise service flow segmentation and optimal path matching, improving wide-area link utilization. The CENI network provides

400G interconnection with deterministic network transmission. The final result: over 100 km long-distance transmission, network throughput stays above 90%.

3.3.4.3 Future Trend Outlook

The Scale-Across network for AIDCs is currently in a stage of vigorous development. It is not only a logical extension of Scale-Out, but also a key path for breaking through the physical limits of single data centers and enabling the continuous evolution of trillion-parameter large models and multi-agent systems.

Driven by organizations such as the Ultra Ethernet Consortium (UEC), Scale-Across network protocols will more likely be built on Ethernet standards. Meanwhile, an open ecosystem is also an important direction for the evolution of Scale-Across protocols.

Chapter 4 Unified Operations Management

Limited by resources and time at the current stage, this version of the report mainly covers supernode-level management (rack level, node level) based on BMC/Redfish in the operations management layer, and does not yet fully elaborate data-center-level topics such as campus-level DCIM, power and environment monitoring, and cluster-level intelligent autonomous operations (AIOps); these will be supplemented in future versions. We are well aware of the importance of this layer to GW-scale AIDCs and sincerely welcome experts and vendors in data center operations, DCIM, power and environment, and intelligent O&M to join the working group and build and improve this chapter with us.

4.1 Current Status and Analysis of AIDC Operations

Current AI supernode per-rack maximum power density is about 150 kW and evolving toward 200-500 kW. Liquid cooling, RoCE lossless networks, and large-scale server clusters have become standard. Services are dominated by large-model training, inference, and compute leasing, with requirements for device online rate, rapid fault handling, and remote control far exceeding traditional IDCs. AIDC operations are in a critical transition period from traditional IDC hosting to high-density, high-reliability, intelligent O&M.

1) O&M model: gradually shifting from manual on-site inspection to a combination of remote centralized O&M + automated O&M + AI predictive O&M, with on-site staff only performing hardware replacement and emergency repairs.

2) Core demands: compute clusters must not be interrupted; a single server/GPU failure requires minute-level localization, isolation, and recovery; massive numbers of servers (ten-thousand/hundred-thousand scale) need unified management. 3) O&M layering: the industry generally divides into in-band O&M (operating systems, services, cluster software) and out-of-band O&M (BMC) (hardware fundamentals, power-on, BIOS, health status) — the two complement each other, and BMC has become the O&M foundation of AIDCs.

4.2 Full-Rack Supernode Management Architecture

Supernode system management software features clear layered control and integrated integration. Its technical architecture covers, from top to bottom, four levels: rack-level control, node-level management, network management, and centralized operations control.

1) Rack-level management extends the control scope to the rack-level microenvironment, achieving power management, distributed CDU cooling regulation, and unified monitoring of environmental sensors such as temperature and humidity. 2) Node-level out-of-band management relies on the baseboard management controller and standard protocols such as IPMI and Redfish to achieve deep monitoring and

control of server hardware status. 3) The network manager faces the high-performance interconnect network, providing automatic topology discovery, health inspection, and visual O&M capabilities to ensure stable and efficient data center network operation. 4) Finally, the centralized operations console aggregates all-domain management capabilities, providing operations staff with a global view and supporting unified cross-domain policy delivery and automated O&M. Supernode system management software uses unified modeling and telemetry mechanisms, building a unified hardware model based on protocol standards such as Redfish, abstracting all components in the supernode — CPUs, accelerator cards, memory, NICs, and PCIe/CXL devices — into standardized data models, and collecting fine-grained telemetry data based on these models.

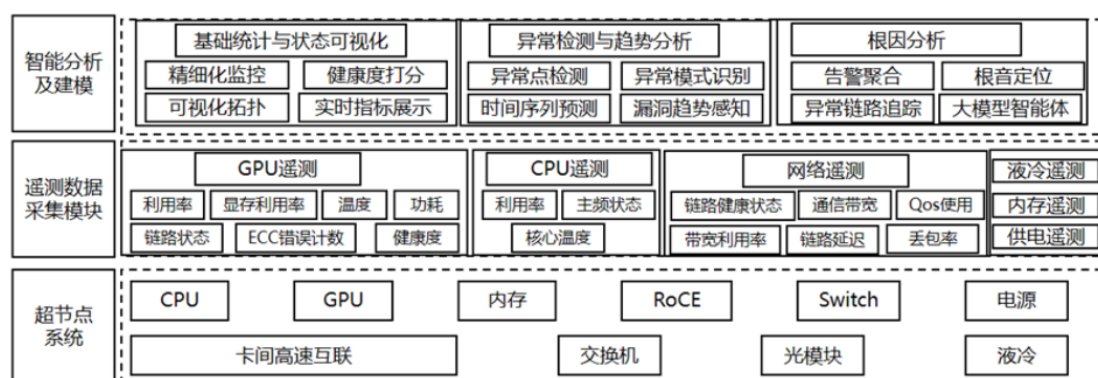


Figure 4.1 Full-rack supernode management functional framework

4.3 Supernode Rack-Level Management

The full rack is managed through the PMC (Power Management Controller) module for power shelf management, PSU (Power Supply Unit) management, and CDU (Coolant Distribution Unit) system management.

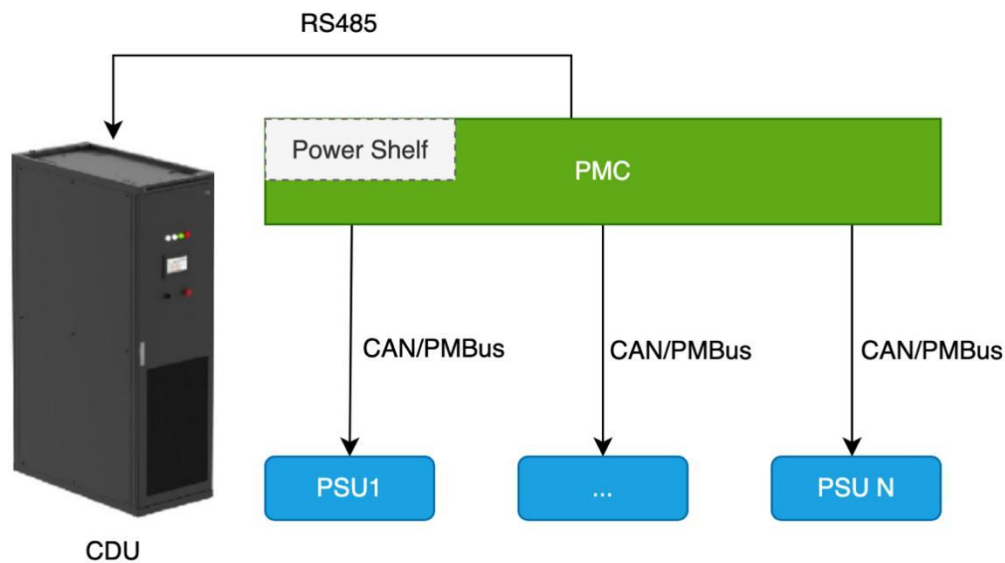


Figure 4.2 Rack-level management block diagram

4.3.1 Power System Management

The power shelf must support sensor monitoring of core voltage, current, input power, output power, ambient temperature, and shelf output busbar clip temperature, as well as monitoring of the total number of PSUs in the power shelf, with power shelf health status logging and alarm functions.

For each PSU in the power shelf, PSU asset query must be supported, including firmware, manufacturer, serial number, FRU, and more. PSU health monitoring must support voltage, current, temperature, and power monitoring with logging and alarms, support PSU black-box log download, and support service-transparent PSU firmware upgrades.

4.3.2 Liquid Cooling System Management

Leak detection sources for full-rack supernodes must cover the full-rack manifold inlet and outlet quick disconnects, flanges, and node core components (CPU/GPU, etc.), supporting leak monitoring and logging/alarm functions.

For scenarios where a distributed CDU independently supplies liquid to the full rack, CDU asset retrieval must be supported, including vendor, product name, serial number, firmware version, liquid level, ambient temperature, liquid temperature, fans, and power management, plus fault monitoring and logging/alarms for core components such as leaks, temperature, fans, and power — with logs retained through power loss. CDU management must support remote and local control of CDU operating parameters and remote CDU firmware upgrades, with the external interface protocol supporting Modbus RTU/TCP.

4.3.3 Rack Management Protocols

The rack PMC (Power Management Controller) connects to the upper-layer operations management platform via the Redfish, IPMI, and SMASH CLP protocols to monitor rack health status and attributes.

4.4 Supernode Node-Level Management

Node-level BMCs achieve observability and controllability of supernode infrastructure by uniformly orchestrating compute trays, switch trays, and power and cooling subsystems.

Node-level management mainly includes: core component status monitoring and lifecycle management, PCIe switch management and topology configuration, and power and cooling control.



Figure 4.3 Node-level management block diagram

4.4.1 Core Component Management

1) Baseboard management. The BMC is responsible for unified monitoring and management of baseboard basic resources. The system continuously collects board-level telemetry data through interfaces such as PMBus, I2C, GPIO, and ADC, and implements rapid detection and reporting of abnormal states through event-driven mechanisms. Mainly includes:

a) Voltage, current, VR, and power monitoring; b) power-on sequence detection and alarms; c) FRU management with support for modifying major fields; d) firmware lifecycle management supporting firmware upgrades and version management for BMC, BIOS, CPLD, PSU, GPU FW, and more — BMC upgrades have no impact on running services.

e) BIOS configuration management, supporting out-of-band BMC modification of BIOS Setup options. 2) CPU management. CPU management mainly covers CPU lifecycle management, operating status monitoring, and fault detection, including:

a) CPU temperature and power monitoring; b) CPU asset information management; c) CPU presence, fault detection, and alarms. 3) GPU management. GPUs are the core compute resources in AIDC supernodes; node-level BMCs need GPU identification, status monitoring, and fault management capabilities.

GPU management capabilities include:

a) GPU asset management; b) GPU temperature and power monitoring; c) GPU presence detection, utilization and VRAM utilization monitoring, and monitoring of recoverable errors, unrecoverable errors, and ECC mode.

4) SmartNIC management. As AI data centers evolve toward intelligent networking, SmartNICs have become key components in supernodes. The BMC can achieve unified DPU management via MCTP, NC-SI, PCIe VDM, and other methods.

The node-level management system needs to support:

a) SmartNIC asset management; b) SmartNIC temperature and power monitoring; c) SmartNIC presence detection, port link status, negotiated rate, and other information monitoring. 5) Memory management. AI training scenarios place extremely high demands on memory stability; the BMC must continuously monitor:

a) Memory temperature monitoring; b) memory asset information monitoring, including vendor, serial number, capacity, bit width, maximum frequency, current frequency, and more;

c) memory presence detection and recoverable/unrecoverable error monitoring. 6) PCIe switch management. In AI supernode scenarios, PCIe fabric stability directly affects GPU-to-GPU communication performance and overall computing efficiency. The BMC must continuously monitor PCIe link health, with real-time detection and alarming for events such as link downgrades, link flapping, error recovery, and unexpected interrupts, ensuring high availability and stable operation of the PCIe fabric.

PCIe switch management mainly includes:

a) link width/speed status monitoring; b) port traffic detection; c) AER error collection; d) switch firmware management. In addition, the system should support dynamic PCIe topology reconstruction and state synchronization. In scenarios such as GPU reset, PCIe recovery, link anomalies, or topology changes, it can automatically rediscover and update topology information, ensuring the BMC-side topology view stays consistent with actual hardware state.

PCIe switch topology management mainly covers the following device types:

a) PCIe Switch; b) GPU Endpoint; c) RNIC Endpoint; d) PHY Retimer; e) NVMe.



Figure 4.4 PCIe SW topology management

4.4.2 Power Management

The node BMC can perform power control on the host server, including power on, forced power off, soft power off, reset restart, and forced restart. For SmartNIC configurations, SmartNIC power on/off control must be supported, as described below.

- 1) Power on: equivalent to a short press of the power button while the server is off; system components power on in sequence per the designed timing.
- 2) Soft power off: equivalent to a short press of the power button while on. This requires correct OS configuration — setting the short-press power button action to direct shutdown — for this function to take effect. Some OSes pop up a confirmation prompt after receiving the soft power-off signal, and shutdown only proceeds after user confirmation in the OS.
- 3) Forced power off: equivalent to a long press (over 5 s) of the power button to force the server off, with CPUs/GPUs powered down synchronously.
- 4) Reset restart: the system restarts immediately, simulating a short press of the power reset button. 5) Forced restart: power off then power on — equivalent to a long press of the power button to shut down, then after waiting 10 s, a short press of the power button to power on.
- 6) SmartNIC power on/off: SmartNICs carry bare-metal and virtual machine services. When the BMC detects a SmartNIC is present, normal power operations require that the SmartNIC not be powered down.

When power cycling the SmartNIC is needed, additional commands must be provided to power the SmartNIC on/off, ensuring new SmartNIC features take effect.

4.4.3 Thermal Management

For node component cooling — apart from liquid-cooled parts — components relying on fan cooling have their fan speeds adjusted by the BMC to maintain normal operating temperatures. Two speed control modes, manual and automatic, must be supported:

- a) Manual fan control: set fan control mode to manual, allowing custom fan speed adjustment;
- b) automatic fan control: set fan control mode to automatic for intelligent automatic speed regulation. For thermal anomaly handling, the following scenarios must be covered:
 - a) BMC hang — fan speed must be raised from the abnormal state to a safe speed;
 - b) during BMC restart or IPMI service abnormalities, the safe speed must be maintained until the BMC finishes starting and the thermal module works normally;
 - c) during firmware updates, thermal risk must be avoided (if the thermal module cannot regulate speed normally during the upgrade, fans must be raised to the safe speed starting from the file upload);
 - d) thermal module abnormality — maintain the safe speed until normal speed regulation resumes. If core components (CPU, GPU, SmartNIC, etc.) have 3 or more consecutive abnormal temperature readings (zero value or unreadable), abnormal speed regulation must be triggered immediately, adjusting to the safe speed.

4.4.4 Log Diagnostics

The BMC must have one-click log collection, automatically completing multi-source log collection, packaging, and upload or export, facilitating rapid problem localization and O&M analysis.

Log collection content should include:

- a) System Event Log (SEL):

records hardware sensor fault events for the baseboard, CPU, GPU, memory, and more. b) Audit/operation logs:

records non-query operations performed on the device. c) HOST logs:

records BIOS and OS boot and runtime SOL logs. d) BMC logs:

records BMC runtime logs, including dmesg, trace logs, internal module exception logs, and more. e) Comprehensive device diagnostic data:

records system component asset information, firmware version numbers, historical sensor data, BIOS option data, thermal speed regulation logs, CPU key status register information, the last screen before a crash, and more.

4.4.5 Alarm Settings

The BMC must support alarm log configuration, including syslog alarm settings, SNMP settings, and email alarm settings, with the following specific requirements:

1) SNMP settings. Forward alarm logs externally via SNMP Trap; must support configuration parameters such as host identifier, Trap version, alarm level, and target device IP.

SNMP Get/Set is used to read device logs; must support SNMP V1/V2C/V3 configuration, community names, and more. 2) Syslog alarm settings. Forward alarm logs externally via syslog remote forwarding; must support configuration of target device IP, alarm level, transport protocol, and more.

3) Email alarm settings. Forward alarm logs externally via email; must support SMTP server address, port, alarm level, and more.

4.4.6 Remote Access

The BMC provides remote access and operation of servers, including graphical access, command-line access, and remote media, without needing a local display.

1) Graphical access. Remotely view and operate the host screen via HTML5 KVM and VNC; VNC supports password changes.

2) Command-line access. Redirect the server's serial port to the SOL console for input/output operations.

3) Virtual media. Support remote mounting of ISO image files, CD/DVD, HDD, and folder-mounted images for remote media use on the host server. Virtual media is off by default; the virtual USB NIC device must automatically turn off after entering the OS; virtual media can be opened via IPMI/Redfish commands or the web.

4.4.7 Users and Security

The BMC is an important management entry point for AI infrastructure. The node-level management system needs a complete identity authentication, permission control, network access protection, and certificate security system to ensure the security and reliability of the supernode platform.

1) User and permission management. The node BMC supports multi-level user management and role-based access control (RBAC), assigning differentiated management permissions to different user roles, effectively reducing misoperation risk and achieving secure isolation and fine-grained access control of system resources.

The node-level management system supports a typical three-level permission model, with permissions divided as follows:

Function module	Administrator	Operator	Regular user
User configuration	✓	×	×

Configure own account	✓	✓	✓
General configuration	✓	✓	×
Power control	✓	✓	×
Remote media	✓	✓	×
Remote KVM	✓	✓	×
Security configuration	✓	×	×
Debug diagnostics	✓	×	×
Query functions	✓	✓	✓

2) Password security policy. To improve system security, the BMC supports password complexity policy configuration, with complex password requirements enabled by default. The password policy supports: minimum password length limits, password complexity verification, password validity period management, historical password reuse detection, and login failure lockout mechanisms. 3) System firewall. The node BMC supports a built-in system firewall to limit illegal access and network attack risks. Administrators can restrict access to management interfaces from specific sources according to the data center network plan, improving management network security.

The system supports the following types of firewall rules: IP address firewall rules, supporting access control based on IPv4/IPv6 blacklists/whitelists; MAC address firewall rules, supporting access control based on MAC address blacklists/whitelists; and port firewall rules, supporting access restrictions on management service ports. 4) SSL certificate management. To ensure communication security of management interfaces, the node-level BMC supports SSL/TLS certificate management for secure communications such as HTTPS.

The system supports generating CSR (Certificate Signing Request) files through the web interface, and supports certificate import and update.

This effectively protects management data transmission security, preventing sensitive information leakage and man-in-the-middle attacks.

4.4.8 BMC Network Management

The node BMC supports flexible network configuration and management capabilities, meeting out-of-band management needs in different data center network deployment scenarios. It supports IPv4 and IPv6 dual-stack network protocols and dynamic switching between static IP and DHCP modes.

To improve the reliability and availability of the node-level management network, the BMC network service has a self-recovery mechanism. Through this mechanism, the node-level BMC can effectively improve the stability and continuous operation of the management network, reducing the impact of network anomalies on remote O&M and cluster management. The network recovery schemes are as follows:

- 1) When the system detects abnormal loss of an IPv4 or IPv6 DHCP address, the network service automatically and periodically sends DHCP/DHCPv6 requests to reacquire a valid network address, ensuring the out-of-band management network recovers promptly.
- 2) When an IPv4/IPv6 network service process exits abnormally, crashes, or is lost, the system can automatically perform service restart and network recovery operations.

4.5 Development Trend Outlook

In the future, the BMC will gradually evolve from a traditional out-of-band management module into an infrastructure control node integrating standardized interfaces, hardware observability, automated O&M execution, secure and trusted control, and scenario coordination. Driven by new compute scenarios such as AIDCs, liquid cooling, and GPU clusters, the BMC's capability boundaries will keep expanding, making it a key component of server, rack, and even cluster-level O&M systems — the BMC is moving from an auxiliary management component to a core control node of intelligent compute infrastructure.

The main development trends can be seen in five directions: Standardization — from IPMI to Redfish-led, unified interfaces across all machine types, reducing the difficulty of managing multiple devices. Intelligence — from "monitoring" to "diagnosis + decision + coordination": the BMC no longer just "collects data" but gradually gains hardware autonomous O&M capabilities.

Security — from management entry point to security boundary: BMC security requirements will grow ever higher; the future BMC is not just an "O&M entry" but an important part of the server hardware chain of trust.

Scenario orientation — evolving deeply toward AIDC/liquid cooling/GPU clusters: the BMC's responsibilities will expand significantly, from a "single-machine view" to a coordinated "node + rack + cluster" view.

Platformization — from single-point management to large-scale orchestration: the future value of the BMC lies not in logging into an individual machine's page, but in whether it can be used in batch, automated, platformized ways.

Conclusion

The GW-scale Open AIDC Framework Technical Report is a systematic response to the rapidly evolving needs of compute infrastructure in the AI era. Following the main themes of openness, efficiency, greenness, and intelligence, it organizes the construction of GW-scale AIDCs into four interdependent layers — the infrastructure layer, the IT facility layer, the network and high-speed interconnect layer, and the system software and resource management layer — emphasizing holistic co-design of energy, computing, storage, networking, cooling, and operations from the planning stage.

One core judgment runs through this report: the competitiveness of a GW-scale AIDC increasingly depends on system-level reliability, maintainability, energy efficiency, and sustainability, rather than single-point peak performance. To address the exponential growth of compute demand, this technical report distills several common evolution directions:

- a) campus-level, phaseable, replicable open reference architectures to reduce cross-entity collaboration costs;
- b) high-power-density building and structural design for racks above 140 kW and campus total power at the gigawatt scale;
- c) green power supply and backup systems represented by high-voltage direct current (HVDC, ± 400 V/800 VDC) and "source-grid-load-storage" coordination;
- d) liquid cooling systems represented by cold plate and immersion, and their coupling with campus heat recovery and renewable energy;
- e) multi-level, open and decoupled high-speed interconnect and network systems across Scale-Up/Scale-Out/Scale-Across;
- f) Redfish as the unified interface for supernode-oriented intelligent autonomy and unified operations management.

These directions embody the intrinsic requirements of system-level co-design and reflect China's unique strengths in new energy supply, high-voltage direct current, manufacturing supply chains, and large-scale engineering organization. The working group hopes this technical report will build industry consensus, inviting partners across all links — component design, chips, complete systems, civil engineering, power supply and distribution, cooling, networking, storage, and operations — to participate jointly, continuously improving along the path of "requirements - architecture - interfaces - specifications," and, on the basis of aligning with OCP's global methodology, promoting systemic solutions that can be understood and adopted internationally, achieving the co-construction of Chinese practice and the global open system.

Glossary

This glossary collects the key terms and abbreviations used in this technical report for the reader's unified reference.

AIDC (AI Data Center): a new-generation data center form characterized by system-level co-design and oriented toward large-scale AI training and inference workloads.

BBU/BESS: rack-level Battery Backup Unit and campus-level Battery Energy Storage System.

BMC (Baseboard Management Controller): performs out-of-band monitoring and control of server hardware; the O&M foundation of AIDCs.

CDU (Coolant Distribution Unit): provides controlled flow, pressure, and temperature for the liquid cooling secondary side.

DPU/SmartNIC: data processing unit / smart network interface card, handling network, storage, and security offload.

GW-scale AIDC: a hyperscale AIDC whose campus total power reaches the gigawatt (10^9 W) level, requiring systems engineering coordination at the campus or even regional scale.

HBD (High Bandwidth Domain): a dedicated network providing high-bandwidth, low-latency communication between accelerators.

HVDC: high-voltage direct current supply (e.g., ± 400 VDC, 800 VDC), used to reduce transmission current and line losses and improve system efficiency.

NVLink/UALink: high-speed inter-card interconnect protocols for accelerators; UALink is an open-standard GPU-to-GPU interconnect specification.

OAM/UBB: Open Accelerator Module and Universal Base Board — core specifications of the OAI system.

OCP (Open Compute Project): an open-source hardware and system design collaboration community for data centers.

OCTC: the Open Compute Technology Committee, affiliated with the China Electronics Standardization Association (CESA).

OISA/SUE: inter-card interconnect technology paths such as the Omni-directional Intelligent Sensing Express Architecture (OISA) and Scale-up Ethernet (SUE).

ORv3 (Open Rack v3): the OCP 21-inch open rack specification for high-density, liquid-cooled, and high-power systems.

Pod: in parallel computing, a group of compute nodes / resource domains scaled vertically (Scale-Up) to act as a single logical unit.

PUE/WUE: Power Usage Effectiveness and Water Usage Effectiveness — core energy and water efficiency metrics for data centers.

Redfish: an open RESTful hardware management standard for out-of-band management, telemetry, and automation.

RoCEv2: RDMA over Converged Ethernet v2, a protocol carrying RDMA over Ethernet.

Scale-Up: vertically scaling the number of accelerators within a single Pod via high-bandwidth interconnect (e.g., NVLink, UALink, OISA, SUE).

Scale-Out: horizontal scaling via high-performance networks, interconnecting multiple Pods/racks into larger clusters (e.g., RoCEv2, UEC).

Scale-Across: compute expansion for cross-AIDC, long-distance interconnection, requiring distance-adaptive load balancing and wide-area deterministic interconnection.

SONiC: an open-source network operating system (Software for Open Networking in the Cloud).

SST (Solid State Transformer): converts medium-voltage AC directly into controlled DC, shortening the power supply chain.

TCS/FWS: Technology Cooling System (secondary side) and Facility Water System (primary side).

UEC (Ultra Ethernet): the Ethernet system optimized for AI/HPC, proposed by the Ultra Ethernet Consortium.

Supernode: a full-rack or cross-rack system in which multiple accelerators are tightly coupled via a high-bandwidth domain and presented externally as a single logical compute unit.

East Data, West Computing: the national computing hub node layout project. Leveraging the endowment differences between the east (dense demand) and the west (rich resources, especially green power and land), it guides the east's non-real-time, latency-insensitive compute demand to the west in an orderly way, optimizing the cross-regional allocation of national compute resources.

Computing-power and electricity coordination: mechanisms and engineering practices for cross-domain coordinated dispatch of compute loads and power systems. Treating compute centers as adjustable loads, it matches their spatiotemporal distribution with power (especially new energy) supply and participates in grid regulation, thereby improving compute center energy efficiency and grid stability.

Cold plate/immersion: the two mainstream liquid cooling methods — cold plates conduct heat through direct chip contact; immersion dissipates heat by submerging IT equipment in dielectric fluid.

References and Further Reading

This technical report draws on the methodology and format of the OCP Open Systems for AI strategic initiative and connects with the open computing work of the OCP China Community and OCTC. Readers are encouraged to consult the following sources:

- a) Open Compute Project, "Open Systems for AI: Blueprint for Scalable Infrastructure," v1.0.0, 2025.
- b) Open Compute Project projects and specifications (Server/DC-MHS, OAI, Open Rack v3, Cooling Environments, Hardware Management/Redfish Profiles, Networking, etc.), <https://www.opencompute.org>.
- c) Open computing standards and technical reports of the Open Compute Technology Committee (OCTC) / China Electronics Standardization Association.
- d) The GW-scale Open AIDC Working Group's public work list and supporting technical documents (published together with the official release).

About OCTC and the OCP China Community

The Open Compute Technology Committee (OCTC) is affiliated with the China Electronics Standardization Association (CESA).

It is dedicated to promoting the research, formulation, and industrial implementation of open computing technical standards, fostering open innovation in data centers from cloud to edge.

The OCP China Community is the open collaboration platform of the Open Compute Project in China, bringing together cloud computing vendors, hyperscale data center operators, telecom and colocation service providers, enterprise users, and product and solution ecosystem partners. Guided by the four tenets of impact, efficiency, scale, and sustainability, it drives major technology evolutions represented by AI/ML, optical interconnect, advanced cooling, composable memory, and silicon, extending the benefits of hyperscale-led innovation to the broader industry.

The GW-scale Open AIDC Working Group was established under the joint guidance of the above two organizations. Tailored to China's national conditions and engineering practice, it focuses on co-building the open reference architecture and specifications for GW-scale AIDCs, and welcomes all industry sectors to join and co-build.

Learn more: <https://www.octc.net>, <https://www.opencompute.org>