

A RoCE-based network framework for science workloads in HEPS data center

Shan Zeng¹, Tao Cui¹, Fazhi Qi¹, and Yanming Wang¹

¹Institute of High Energy Physics,CAS, Computing Center, 100049 Beijing, China

Abstract. According to the estimated data rates, it is predicted that 800 TB raw experimental data will be produced per day from 14 beamlines at the first stage of the High-Energy Photon Source (HEPS) in China, and the data volume will be even greater with the completion of over 90 beamlines at the second stage in the future. Therefore, designing a high-performance, scalable network architecture plays a crucial role in the efficient output of scientific tasks. We designed a RoCE-based network framework for science workloads in HEPS data center, which provides high-performance network connectivity both between HPES Data Acquisition system(DAQ) and HEPS data center, and between the compute and storage system within the HEPS data center.

1 Introduction

1.1 Introduction to HEPS and its data

HEPS is currently under construction in Huairou, Beijing, which represents China’s first fourth-generation synchrotron radiation facility and one of the most advanced light sources globally. HEPS is designed to serve diverse fields, including: structural biology, quantum materials, energy storage, environmental science, ultrafast phenomena and etc[1]. The official commissioning date is currently scheduled for 2025.

Table 1. Estimated data volumes for first-batch beamlines at HEPS.

Beamlines	Burst data volume per day	Average data volume per day
B1 Engineering materials beamline	600 TB	200 TB
B2 Hard X-ray multi-analytical nanoprobe (HXMAN) beamline	500 TB	200 TB
B3/B4/B5/B6/B8/B9/BA/BC/BD/BE	1-80 TB	0.2-11.2 TB
B7 Hard X-ray imaging beamline	1000 TB	250 TB

1 Corresponding author: zengshan@ihep.ac.cn

BB Pink small-angle X-ray scattering	400 TB	50 TB
BF Test beamline	1000 TB	60 TB

HEPS produces highly heterogeneous datasets characterized by extreme data volumes, with daily raw data volume projected to range from several megabytes to multiple terabytes per beamline, depending on experimental parameters and detector configurations[2]. Notably, burst data generation rates may transiently reach up to several terabits per second (Tb/s) during high-intensity synchrotron radiation experiments. Tables 1 and 2 provide detailed estimates of data volume and data throughput rates for the first operational beamlines at HEPS, stratified by instrument type and application domain.

Table 2. Estimated data generation rates for first-batch beamlines at HEPS.

Beamlines	Burst data generation rate	Average data generation rate
B1/B6/B7/B9	160-320 Gb/s	47.5 Mb/s - 34 Gb/s
B2/B3/B8/BA/BB/BC/BE	8-80 Gb/s	285 Mb/s-19 Gb/s
B4 Hard X-ray coherent scattering beamline	3.2 Tb/s	4.75 Gb/s
B5/BD	47.5-800 Mb/s	0.025-95 Mb/s
BF Test beamline	3.2 Tb/s	34 Gb/s

1.2 Network challenges

The network underpins scientific computing and storage. The massive data volumes generated at HEPS, paired with ultra-high burst data rates (ranging from Mb/s to Tb/s per beamline), necessitate a network infrastructure that sustains high bandwidth consistently while minimizing latency. Real-time data transfer from beamline storage to central repositories requires low-latency, high-throughput connectivity to prevent bottlenecks during burst operation. Additionally, the diverse computing modes (online, offline, and AI-driven) introduce heterogeneous traffic patterns, such as bursty short-reads for AI training, large sequential writes for offline analysis, and interactive low-latency queries for online processing. So the key challenges for network system designing include:

- **Bandwidth Scalability:** Supporting concurrent data ingestion from 14+ beamlines with varying rates.
- **Latency Sensitivity:** Ensuring sub-millisecond response times for real-time feedback loops in AI and online experiments.
- **QoS Differentiation:** Prioritizing latency-critical flows (e.g., live imaging) over bulk transfers.
- **Burst Handling:** Mitigating transient traffic surges without degrading performance.

This requires a hybrid network architecture integrating fiber-optic backbones and RDMA(Remote Direct Memory Access)-enabled fabrics, complemented by intelligent traffic engineering and machine learning-driven predictive analytics to dynamically allocate resources.

2 RDMA technology

RDMA is a network technology that enables direct memory access between two nodes in a computer network without involving the CPU of either node. This eliminates the need for data to be copied from system memory to a CPU buffer and then transmitted over the network, drastically reducing latency, overhead, and power consumption[3]. Two dominant RDMA protocols are InfiniBand (**IB**) and RDMA over Converged Ethernet (**RoCE**).

2.1 IB

IB is a high-speed networking architecture designed for data-intensive applications requiring ultra-low latency, high bandwidth, and efficient CPU offloading. Developed by the InfiniBand Trade Association (IBTA), It is an isolated infrastructure which enables RDMA between nodes using dedicated hardware (switches, NICs) and a proprietary protocol stack(see in Fig.1), minimizing CPU involvement and achieving microsecond-level latency[4]. IB has become a cornerstone in supercomputing, AI training, and large-scale data centers to ensure low interference and scalability for massive parallel workloads.

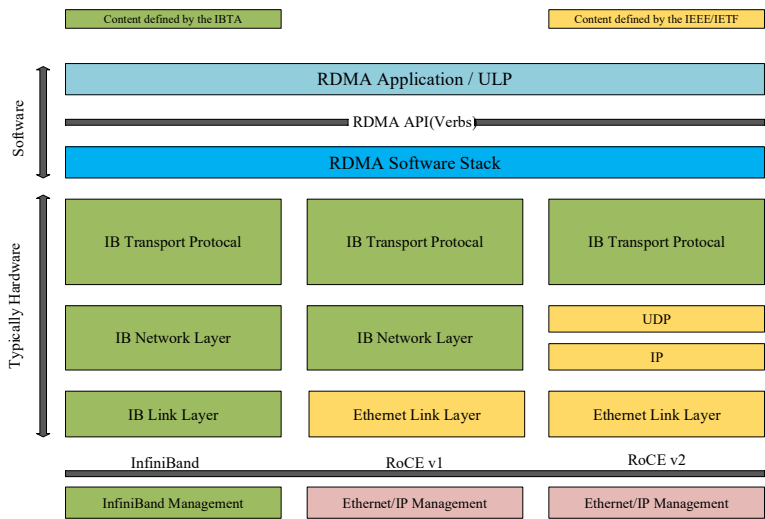


Fig.1. Underlying ISO Stacks of the Flavors of RDMA

2.2 RoCE

RoCE is a protocol that delivers RDMA capabilities over standard Ethernet networks, combining data transfer efficiency with the cost and flexibility of IP-based infrastructure. It eliminates the need for specialized hardware like InfiniBand, allowing seamless integration into existing Ethernet environments. Nowadays, RoCE is widely adopted in cloud data center, storage networks (NVMe over Fabrics), and edge computing, enabling efficient inter-node communication for Kubernetes clusters, distributed storage systems, and real-time analytics while balancing performance and cost-effectiveness[5].

2.3 Key differences between IB and RoCE

Table 3. shows the main differences between IB and RoCE through several dimensions, including infrastructure, protocol stack, cost, latency and scalability.

Table 3. Key differences between IB and RoCE.

Feature	IB	RoCE
Infrastructure	Dedicated (fiber, proprietary switches)	Shared (Ethernet, IP routers)
Protocol Stack	Proprietary (Verbs API)	Standard Ethernet/RoCE protocols
Cost	High	Low
Latency	Ultra-low (~1 μ s)	Low but depends on network conditions
Scalability	High (massive clusters)	Moderate (limited by Ethernet topology)

3 Applications inside HEPS DC

HEPS data center supports a diverse suite of applications to enable efficient data management and computational workflows, which including:

- **Beamline Storage Filesystem:** Provides on-site data storage for experimental data generated by beamlines.
- **Central Storage Filesystem:** Serves as a centralized repository for long-term data archiving and cross-facility access.
- **Tape Library Storage:** Offers cost-effective, high-capacity backup for irreplaceable datasets.
- **MDW (Mamba Data Worker):** A dedicated service responsible for writing raw experimental data from beamlines to the beamline storage filesystem[6].
- **Data Transfer System:** Orchestrates the migration of data from the beamline storage filesystem to the central storage filesystem, followed by backup to the tape library system.
- **Kubernetes (K8s) Cluster:** Hosts containerized applications for both online data processing (e.g., real-time analysis) and offline batch jobs (e.g., simulation workflows).
- **HPC Cluster:** Delivers high-performance computing resources for computationally intensive tasks such as molecular dynamics simulations and machine learning training.

4 Data center network architecture design

DC network architecture ensures seamless integration between data acquisition, storage, and computation, enabling researchers to efficiently manage large-scale datasets and execute complex scientific workflows within the HEPS ecosystem.

4.1 RoCE-based converged Ethernet network

We designed a Spine-leaf architecture network, which enables seamless coexistence of standard Ethernet traffic (e.g., TCP/IP-based services) and RoCE traffic (RDMA-over-Ethernet for low-latency, high-throughput computing) by leveraging advanced QoS mechanisms, priority-based flow control (PFC), and congestion management protocols. As illustrated in Fig.2, network architecture employs a two-tier spine-leaf topology consisting of two spine switches and eight leaf switches, provisioning 25G and 100G Ethernet connectivity to accommodate varying bandwidth demands. Gateway functions for the access layer are collocated at the leaf nodes, minimizing latency and optimizing traffic routing between compute/storage resources and the core network. Between spine and leaf layers, dynamic routing protocols establish adaptive paths, while Equal-Cost Multipath (ECMP) distributes traffic across redundant links, ensuring load balancing, maximizing throughput, and enhancing fault tolerance. This design achieves scalable, low-latency communication with efficient bandwidth utilization, making it suitable for data-intensive applications such as AI training, HPC, and real-time analytics.

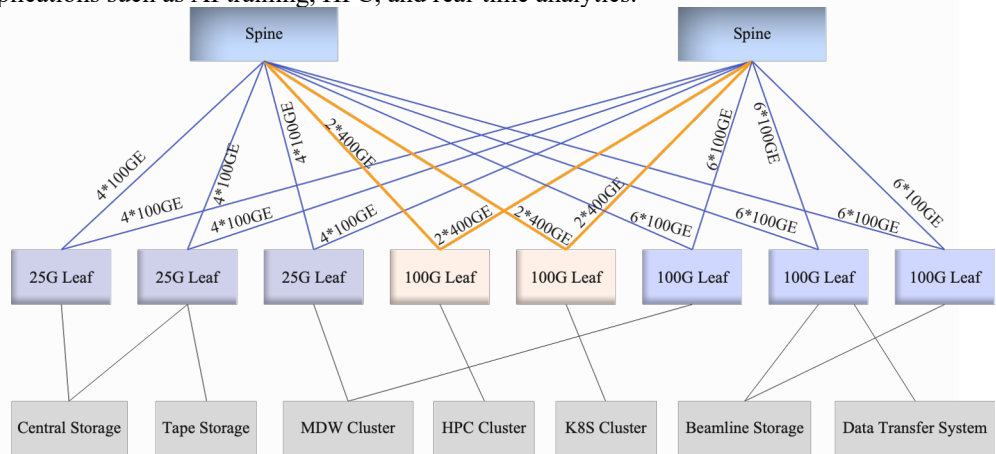


Fig.2. DC Network Architecture

4.2 DCQCN Implemented for efficient congestion control

Data Center Quantized Congestion Notification(DCQCN) is a critical congestion control framework, standardized under IEEE 802.1Qau. Unlike traditional TCP-based approaches, DCQCN is specifically optimized for lossless, low-latency RDMA traffic, addressing the challenges of buffer overflow and incast congestion in modern data centers. A key innovation of DCQCN lies in its compatibility with Priority-based Flow Control (PFC), which prevents buffer starvation across multiple traffic classes[7]. By integrating with PFC, DCQCN balances congestion control with traffic prioritization, ensuring latency-sensitive RDMA flows coexist efficiently with best-effort TCP traffic. Additionally, DCQCN’s adaptive rate adjustment mechanism supports bursty workloads typical in AI training, HPC, and cloud-scale analytics, where instantaneous bandwidth demands can fluctuate drastically. So we implemented DCQCN in HEPS DC to reduce latency and optimize bandwidth utilization.

5 Operation monitoring

We developed a comprehensive network performance evaluation framework that integrates multiple components to ensure thorough monitoring and optimization of the HEPS

data center network, as illustrated in Fig. 3. The core concept of this framework is to capture integrated networking operation data from syslog, network telemetry[8-9], and xFlow[10]. By analyzing this big data, the framework functions as an intelligent system, delivering a robust network monitoring service. This service includes network health assessments, failure reports, optimization suggestions, and more. The overall system operation monitoring dashboard is shown in Fig. 4.

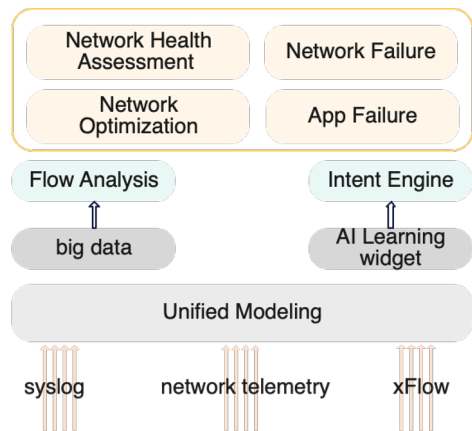


Fig.3. Network Performance Evaluation Framework

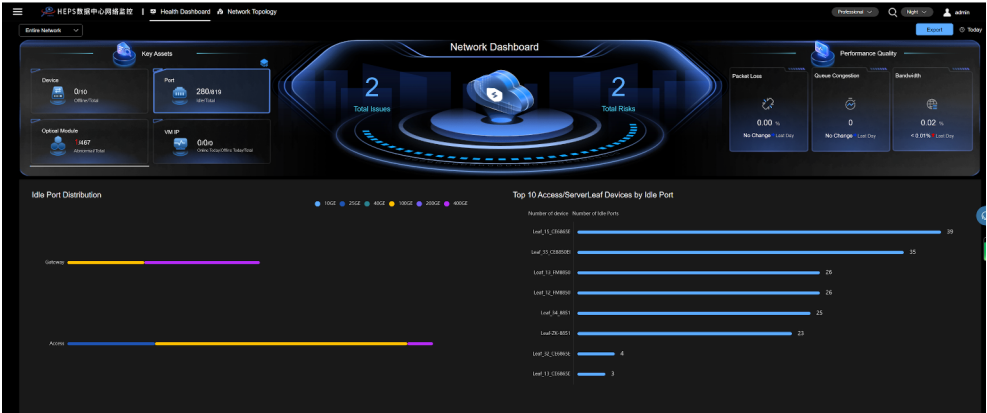


Fig.4. Network Health Overview Dashboard

6 Conclusion and future work

HEPS will officially start data acquisition in 2025. Based on preliminary test results, the RoCE-based high-performance data center network we designed can meet the data analysis and storage requirements of HEPS Phase I. Going forward, as data acquisition progresses, we will also optimize network performance based on the analysis of business network traffic models, thereby providing more efficient and flexible data center network services.

7 Acknowledgements

This paper is funded by the National Natural Science Foundation of China (No. 12175258)

References

1. Z.Zhao. *Fourth-Generation Synchrotron Light Sources: Principles and Applications*. *Nature Reviews Physics*, **3**(8), 555–569(2021).
<https://doi.org/10.1038/s42254-021-00335-0>
2. X. Li et al., “A high-throughput big-data orchestration and processing system for the High Energy Photon Source,” *J Synchrotron Rad*, vol. 30, no. 6, Art. no. 6, Nov. 2023, doi: 10.1107/S1600577523006951
3. A.Kalia, M.Kaminsky, & D.G.Andersen (2016). Design guidelines for high performance RDMA systems. *USENIX Annual Technical Conference (USENIX ATC)*, 437–450. <https://www.usenix.org/conference/atc16/technical-sessions/presentation/kalia>
4. S.Zeng, F.Qi, L.Han. Research and Evaluation of RoCE in IHEP Data Center[J].*The European Physical Journal Conferences*, 2021.DOI:10.1051/epjconf/202125102018.
5. Li, Y., Miao, R., Liu, H., Zhu, Y., Yu, M., & Zhang, C. (2020). *Understanding RDMA over commodity Ethernet at scale*. *USENIX Symposium on Networked Systems Design and Implementation (NSDI)*, 647-665. <https://www.usenix.org/conference/nsdi20/presentation/li-yuliang>
6. Liu Y , Geng Y D , Bi X X ,et al.Mamba: a systematic software solution for beamline experiments at HEPs[J]. 2022.DOI:10.48550/arXiv.2203.17236.
7. *Microsoft Azure Networking Team*. (2016). DCQCN: Congestion control for RDMA in cloud datacenters. *Microsoft Research Whitepaper*. <https://www.microsoft.com/en-us/research/publication/dcqcn-congestion-control-for-large-scale-rdma-deployments/>
8. Handigol, N., Heller, B., & Jeyakumar, V. (2023). *P4TELEMETRY: A universal framework for data plane telemetry*. *ACM SIGCOMM Conference*, 45-58. <https://doi.org/10.1145/3570361.3570367>
9. Benson, T., Akella, A., & Maltz, D. A. (2010). Network traffic characteristics of data centers in the wild. *ACM Internet Measurement Conference (IMC)*, 267-280. <https://doi.org/10.1145/1879141.1879175>
10. [1] Li B , Springer J , Bebis G ,et al.A survey of network flow applications[J].*Journal of Network & Computer Applications*, 2013, 36(2):567-581.DOI:10.1016/j.jnca.2012.12.020.