



SAKURAONE: Open Ethernet-Based, AI-Ready HPC with 800GbE and SONiC

Fumikazu KONISHI¹, Yuuki Tsubouchi¹, Hirofumi Tsuruta¹, Takashi Inoue², Kiyohiro Kurosawa²

1) Research Center, SAKURA internet inc. 2) SAKURA internet inc., Japan

Background/Objectives

Japan's private sector requires production-grade AI computing on par with Big Tech, but with vendor-neutral infrastructure.

A key motivation is the use of open, disaggregated networking. This design improves vendor agility and supply chain resilience by decoupling NOS from the underlying ASICs, accelerates innovation through community-driven modular components, and enables lossless AI fabrics through RoCEv2 QoS (PFC, ECN).

Scalable overlays (EVPN/VXLAN) further support multi-tenant AI workflows, while open management frameworks provide automation and observability essential for transparent large-scale operations.

SAKURAONE, SAKURA internet's managed GPU supercomputer, was developed to:

- (i) demonstrate the scalability of an all-Ethernet fabric with SONiC for AI/HPC,
- (ii) provide reproducible AI software stacks for enterprises, and
- (iii) evaluate performance using community benchmarks.

At ISC~2025, SAKURAONE ranked 49th on the TOP500 (HPL) and, to our knowledge, is the only Top-100 system built entirely on an open 800GbE + SONiC networking stack.

System:

100 nodes; each: 2 Intel Xeon Platinum~8580+, 1.5TB DDR5, 8 times NVIDIA H100 SXM5 (80GB).

Network:

Rail-optimized leaf-spine; host links 400GbE; leaf-spine 800GbE; RoCEv2; SONiC on Broadcom Tomahawk5; NIC GPU affinity preserves NUMA locality for collectives.

Storage:

2PB all-flash Lustre on DDN ES400NVX2-NDR200; per-node dual 200GbE for storage.

Software:

Rocky Linux9, Slurm, CUDA~12.x, NCCL~2.2x, Apptainer/Singularity & Pyxis, curated conda ML stacks.

Benchmarks:

HPL, HPCG, HPL-MxP (HPL-AI class), IO500, MLPerf Training.

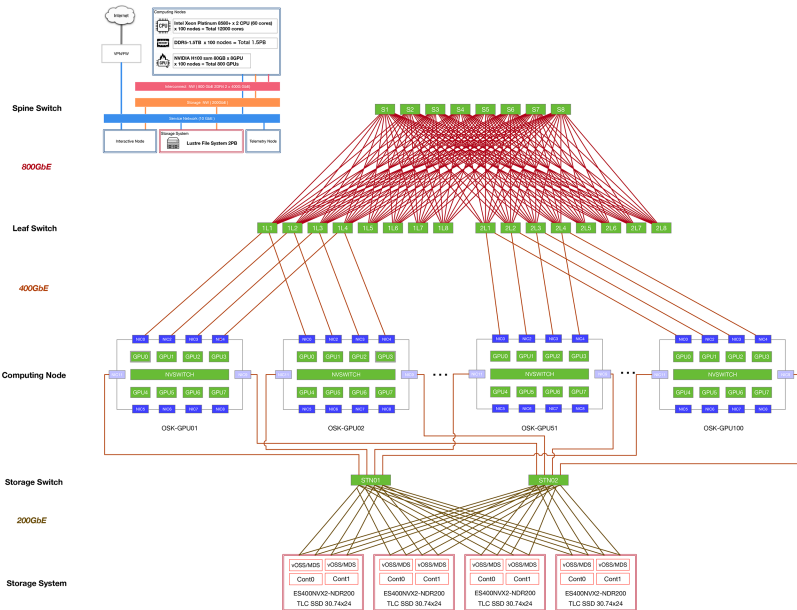


Table 1: SAKURAONE Computing Nodes

Name	Description
Chassis	Supermicro GPU SuperServer SYS-821GE-TNHR
CPU	Intel Xeon Platinum 8580+ x 2 CPUs (200W)
Core (per CPU)	32Cores
GPU	NVIDIA H100 SXM5 80GB x 8 GPUs
Memory (RAM)	DDR5-1.5TB (96GB x 16)
System storage (SAS)	372GB x 2
Data storage (NVMe)	7.68TB x 4
Local network interface	40GbE (800G x 2)
Interconnect network interface	NVIDIA Mellanox MCX75310AAS-NEAT ConnectX-7 400GbE/ND8 x 8
Storage network interface	NVIDIA Mellanox MCX75310AAS-NEAT ConnectX-7 400GbE/ND8 x 2

Table 3: SAKURAONE Interconnect Network

Name	Description
Network technology	Gigabit Ethernet (GbE)
Ethernet switch speed grade	800 GbE 2DR4 2 x 400G GbE
Protocol	RoCEv2 (RDMA over Converged Ethernet)
Network topology	Rail-Optimized Topology
Switch Chassis	Edge-core networks AIS800-640
Switch Capability	51.2Tbps full-duplex
Software Stack	SONiC (Software for Open Networking in the Cloud)
Switch Chip	Tomahawk 5 chipset powered by Broadcom
Leaf Switch	16 chassis
Spine Switch	8 chassis

Table 2: Classification of NIC Usage and GPU Connectivity Characteristics

NIC	Device Name	Primary Usage	GPU Connectivity Type
NIC0	mils_0	High-speed inter-node communication	NOE (via GPU PCIe domain)
NIC1	mils_1	High-speed inter-node communication	NOE (via GPU PCIe domain)
NIC2	mils_2	High-speed inter-node communication	NOE (via GPU PCIe domain)
NIC3	mils_3	High-speed inter-node communication	NOE (via GPU PCIe domain)
NIC4	mils_4	High-speed inter-node communication	NOE (via GPU PCIe domain)
NIC5	mils_5	High-speed inter-node communication	NOE (via GPU PCIe domain)
NIC6	mils_6	High-speed inter-node communication	NOE (via GPU PCIe domain)
NIC7	mils_7	High-speed inter-node communication	NOE (via GPU PCIe domain)
NIC8	mils_8	Storage network (dedicated IO path)	P2B
NIC10	mils_10	Storage network (dedicated for redundancy)	P2B (logical, multi-bridge path)
NIC9	mils_11	Management network (e.g., SSH)	SYS

This classification is based on mrd-ai-mlt. Usage: -mg-outputs, analyzing PCIe connectivity and NUMA affinity between GPUs and NICs. NIC0-NIC7 are directly assigned to GPUs by using NOE and PCIe paths for optimized inter-node communication. NIC8 and NIC10 are likely dedicated to storage traffic and are accessed via longer PCIe paths or logical bridging. NIC9 is connected to the system via a 10GbE level path that crosses NUMA domains.

Table 4: SAKURAONE Storage System

Name	Description	Count
Chassis	DDN ES400NVX2	4
Controller	Active Dual Controller	2
CPU	Ice Lake CPUs	2
NVMe	24 Drive (PCI Gen4)	24
Drive	TLC SSD 30.72TB	16
Interface	200 GbE	4

Item	Value
Matrix size (N)	2,706,432
Block size (NB)	1024
Process grid (P x Q)	16 x 49
Total processes	784
Total GPUs	784
HPL version	HPL~NVIDIA 25.4.0
Execution time (sec)	389.23
PFLOPS	33.95 PFLOPS
FLOPS per GPU	43.31 TFLOPS
Max GEMM performance (single GPU)	55.34 TFLOPS
GPU SM count	132
GPU peak clock frequency	1980 MHz

Table 8: HPL-MxP Benchmark Summary

Item	Value
Benchmark version	HPL-MxP-NVIDIA 25.4.0
Matrix size N	2,989,056
Block size NB	4096
Process grid (P x Q)	24 x 32
Total processes	768
Peak clock frequency	1980 MHz
GPU SM count	132
Observed Rmax	3.3986e+08 GFLOPS
Rmax per GPU	44,2520.81 GFLOPS
LU-only	5,3919e+08 GFLOPS
LU-only per GPU	702,074.99 GFLOPS
Precision mode	Sluggish FPS (sluggish-type = 1)
Validation result	PASSED (5.01e-05 < 1.6e+01)

Table 10: MLPerf Training (GPT-3) Benchmark Summary

Item	32 Nodes	96 Nodes
Total GPUs	256	768
Data Parallelism	4	6
Tensor Parallelism	4	8
Pipeline Parallelism	16	16
Context Parallelism	1	1
Sequence Parallelism	TRUE	TRUE
Global batch size	1024	2304
Micro batch size	2	6
Time-to-train (min)	105.31 (unverified)	41.86 (unverified)
Model FLOPS Utilization (%)	38.3	35.9
Training Throughput (tokens/s/GPU)	707.62	714.23
Computational Throughput (TFLOPS/GPU)	757.13	710.73

Table 7: HPCG Benchmark Summary

Item	Value
Benchmark version	HPCG 3.1
Total distributed processes	784
Threads per process	16
Global problem dimensions (nx x ny x nz)	4096 x 3584 x 3808
Number of nonzero terms	1.51 trillion
Total memory used (GB)	39,961.4
Memory for linear system and CG (GB)	35,169
Peak memory bandwidth (observed)	3.316 TB/s
Total GFLOPs (raw)	437,361
GFLOPs (with convergence overhead)	404,964
Final validated HPCG GFLOPs result	396,295

Table 9: Comparison of IO500 Results: 10 Nodes vs 96 Nodes

Benchmark	10 Nodes	96 Nodes
io-easy-write (GiB/s)	262.91 (340.96 s)	198.80 (355.68 s)
mdtest-easy-write (kiOPS)	204.44 (347.00 s)	256.64 (363.81 s)
io-hard-write (GiB/s)	15.84 (354.60 s)	24.61 (491.75 s)
mdtest-hard-write (kiOPS)	120.84 (340.05 s)	151.59 (332.84 s)
find (kiOPS)	1976.05 (56.39 s)	2637.17 (54.01 s)
io-easy-read (GiB/s)	365.71 (245.15 s)	305.86 (231.13 s)
mdtest-easy-star (kiOPS)	358.75 (197.40 s)	463.13 (200.14 s)
io-hard-read (GiB/s)	205.64 (311.23 s)	255.31 (48.82 s)
mdtest-hard-star (kiOPS)	262.43 (157.19 s)	408.13 (124.54 s)
mdtest-easy-delete (kiOPS)	168.19 (422.12 s)	198.91 (468.91 s)
mdtest-hard-read (kiOPS)	205.39 (200.53 s)	310.87 (162.97 s)
mdtest-hard-delete (kiOPS)	92.29 (445.91 s)	111.28 (453.80 s)
Bandwidth Score (GiB/s)	133.03	139.80
IOPS Score (kiOPS)	248.74	327.84
Total IO500 Score	181.91	214.09

Table 11: MLPerf Training (Llama2 70B LoRA) Benchmark Summary

Item	1 Node	8 Nodes	64 Nodes	96 Nodes
Total GPUs	8	64	512	768
Data Parallelism	2	8	64	96
Tensor Parallelism	4	4	4	4
Pipeline Parallelism	1	1	1	1
Context Parallelism	1	2	2	2
Sequence Parallelism	True	True	True	True
Global batch size	8	64	96	96
Micro batch size	1	1	1	1
Time-to-train (days)	28.44 (unverified)	4.79 (unverified)	1.94 (unverified)	1.26 (unverified)

Results

The system achieved a sustained HPL (FP64) performance of 33.95 PFLOP/s (Rmax) on 100 nodes equipped with 800 GPUs. For HPCG, it delivered 396.295 TFLOP/s.

On HPL-MxP (FP8), the system reached 339.86 PFLOP/s, corresponding to approximately 442.5 TFLOP/s per GPU, and 539.19 PFLOP/s in the LU-only phase.

In IO500 (production), the system scored 181.91 on 10 nodes, showing strong metadata scaling. Despite higher latency than InfiniBand/Slingshot, the rail-optimized 800 GbE + RoCEv2 + SONiC fabric sustained competitive AI/HPC throughput, confirming open-stack scalability at Top500 scale.

For MLPerf Training, the system achieved scores of 105.310 unverified (32 nodes) and 41.862 unverified (96 nodes) on GPT-3, while on Llama 2 70B-LoRA it achieved 28.444 unverified (1 node), 4.789 unverified (8 nodes), 1.941 unverified (64 nodes), and 1.256 unverified (96 nodes), respectively. Result not verified by MLCommons Association.

Conclusions

SAKURAONE shows that an open, SONiC-based 800 GbE fabric can deliver Top500-class FP64 and competitive AI performance at national scale while preserving vendor neutrality and cost efficiency.

Future work includes

- (i) NCCL/MPI tuning with rail-aware routing and NUMA pinning, and
- (ii) topology-aware placement with adaptive congestion control.

Reference

Fumikazu Konishi. Sakuraone: Empowering transparent and open ai platforms through private-sector hpc investment in japan, 2025. URL <https://arxiv.org/abs/2507.02124>.