



What Is Abaco?

The GPU and memory-centric rack architecture behind PNNL's new AI system, and Lqid's role in it

The Problem

Compute, storage, and networking all became poolable resources. Memory did not. It stayed welded to the CPU socket, sized at purchase, stranded there for the life of the server. So capacity sits idle in one server while the server beside it pages out to NVMe, and some workloads have nowhere good to run at all.

AI widened the gap. KV cache for long-context inference, vector and graph indexes, and the intermediate data in AI-assisted simulation all grow faster than HBM capacity or per-socket DIMM slots. Buying more GPUs to get HBM to cache data is the most expensive available answer to a capacity problem.

"For several decades, the HPC community has focused on distributed memory system architectures... it has led to 'orphaned' applications that require large directly addressable memory capacity."

JAMES A. ANG, PH.D., CHIEF SCIENTIST FOR COMPUTING, PNNL

Architecture At A Glance

Layer	What Is In It
Compute	Up to 16x AI compute servers (Supermicro Hyper X14/H14 or Dell R770 class), dual-socket, PCIe Gen5, 4 GPUs and one CXL host adapter card each, plus 1 fabric manager host.
Fabric	Liquid PCIe Gen6 and CXL 2.0 switching fabric and host bus adapters. Enable multiple CXL/PCIe x16 uplinks per chassis, expandable fabric with PCIe and CXL TOR switching.
Memory	4x memory chassis. Each is a Liquid EX-5410C 4U enclosure with 10x Primemas PMA CXL add-in cards at 4TB per card: 160 Micron RDIMMs, up to 40TB per chassis.
GPU	3x GPU chassis. Each is a Liquid EX-5410P 4U enclosure with 10x Nvidia H200 NVL GPUs: additional device types including NVMe, FPGA, and DPU can also be supported.
Software	RHEL 9.1. Liquid Matrix® as fabric manager, with Kubernetes and Slurm plug-ins. Micron famfs kernel for coherent shared access.

Each rack draws support up to 6,000W per chassis in a fully air-cooled system.

Target Workloads

KV cache for LLM inference (NVIDIA Dynamo), graph analytics (Pomerty), large databases (RocksDB), vector and graph RAG, in-memory analytical caches, and HPC codes such as computational chemistry and molecular biology.

What Liquid Contributes

Micron supplies the DRAM and integrates the system. Liquid supplies what makes a large pool of DIMMs behave like a resource instead of an inventory item.



Liquid EX-5410C Memory Expansion Chassis

A 4U enclosure with ten CXL add-in card slots and up to 40TB of DRAM. Five CDFP Gen5 x16 uplink ports reach as many as 16 hosts at 640GB/s aggregate bandwidth and roughly 500ns latency. In a Liquid rack-scale configuration, up to 200TB can go to a single host or be shared across as many as 16 nodes. The EX-5410C is the memory chassis in every Abaco configuration.



Liquid EX-5410P GPU Expansion Chassis

A 4U enclosure with ten PCIe add-in card slots for GPUs and other PCIe devices. Four CDFP Gen5 x16 uplink ports reach as many as 16 hosts at 512GB/s aggregate bandwidth and roughly 500ns latency. In a Liquid rack-scale configuration, up to 30x GPUs can go to a single host or be shared across as many as 16 nodes. The EX-5410P is the GPU chassis in every Abaco configuration.



Liquid Matrix® Software

The fabric manager, and the reason the pool is composable rather than statically partitioned. Matrix is the industry's only unified interface for real-time deployment, management, and orchestration of composable GPU, memory, and storage. Plug-ins connect it to Kubernetes and Slurm; it also works with Ansible, VMware, and Nutanix NKP. A scheduler can request tens of GPUs and terabytes of memory on-demand.



Liquid PCIe Gen5 and CXL 2.0 Fabric

A dedicated PCIe Gen5 and CXL 2.0 switch and host bus adapters carry low-latency, high-bandwidth traffic between hosts and the pool. Dual-fabric configurations pool GPUs and memory on the same platform, and either fabric runs on its own.

How The Pool Becomes Shareable: famfs

Expanding memory is one problem. Letting several hosts work on the same resident dataset is harder, and it is where Liquid and Micron built something that did not exist before: the industry's first architecture supporting true memory sharing. famfs, Micron's open-source fabric-attached memory file system, presents the Liquid pool to the host OS as a file system with software cache coherence across hosts.

Two consequences matter when sizing a system. Code that already reads files reaches pooled memory through the same interface, so no application rewrite is required. And hosts share one copy of a dataset in place instead of each pulling its own, which cuts data movement and removes the storage round trip.

Measured Results

Workload	Result	Detail
Graph Analytics	Up to 30x faster	Processing time on a graph database query set fell from roughly 8 hours to 15 minutes.
KV Cache LLM Inference	Up to 7x tokens/sec	Certain KV cache workloads, with NVIDIA Dynamo tiering to the Liquid pool.
In-memory Analytical Cache	3.9x throughput	Proven on single-server transaction throughput. Target for the shared configuration is 10x.
Agentic LLM Inference	5.5x tokens per dollar	Projected for future agentic inference workloads. Not yet measured.

Micron lab testing, representative workloads on production hardware and software. Configuration: Intel Xeon 6 host, 2 NVIDIA A100 per server, 2TB native DDR5-5600, 30TB Micron NVMe, and a Liquid chassis with 10 Primemas add-in cards delivering 40TB at 250GB/s read and 434ns latency. RHEL kernel 6.14.11 with weighted software interleaving and famfs; NVIDIA Dynamo tiering. Models: Llama 3.1 8B and GPT-OSS-20B, SLO 2.2s TTFT / 50ms TPOT.

Abaco at Pacific Northwest National Laboratory

PNNL's memory-centric program began with Crete, a 15TB active-memory testbed that came online in August 2025. Abaco scales that approach by more than an order of magnitude. Liquid provides the pool exposing hundreds of terabytes of coherent active memory to AI for Science workloads. The first user-facing application is AI-integrated computational chemistry; PNNL and Micron have already brought the exascale-ready ExaChem code up on CXL memory.

Learn More

What Is Abaco:

<https://www.liqid.com/products/abaco>

Liquid Fabric and Software:

<https://www.liqid.com/products/liqid-matrix-software>

Liquid HCL (Hardware Compatibility List):

<https://www.liqid.com/resources/library?tab=hcl-tab>

