

LStore: Scalable Storage for AI and Research Computing

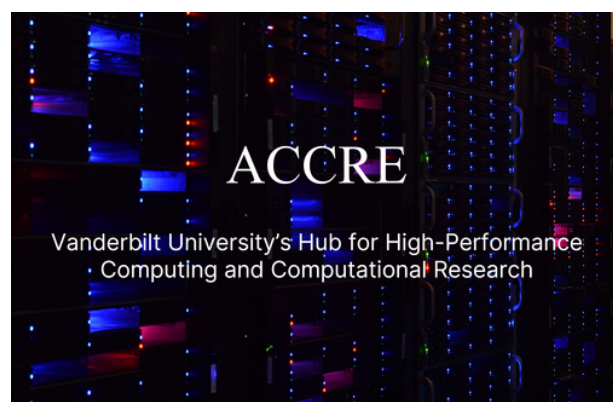
Multi-Petabyte Storage Designed for GPU-Accelerated Research Environments

Executive Summary

AI-driven research, machine learning, and large-scale simulation are increasing infrastructure demands across higher education. Research computing environments are expanding to support more GPU nodes, larger datasets, and more concurrent projects. [Institutions that previously operated at single-digit petabyte scale are now planning for multi-tens of petabytes, with consistent annual growth.](#)

While public cloud platforms offer convenience, long-term economics tell a different story. Storage that appears affordable when parked becomes significantly more expensive when accessed, analyzed, or moved. For AI research, where data is read repeatedly and at scale, recurring cloud costs can quickly exceed the cost of building on-prem infrastructure.

[ACCRES \(Advanced Computing Center for Research and Education\)](#), the premier resource for the high-performance computing research needs throughout Vanderbilt University, modernized its research data environment to address these pressures. By implementing a multi-tier architecture that includes large-scale LStore (Logistical Storage) capacity, ACCRES reduced storage costs, increased flexibility, and aligned infrastructure more closely with research workloads. This deployment demonstrates how a tiered LStore architecture can support large-scale university research.



Building on architectural principles demonstrated in environments such as ACCRES, SourceCode now offers LStore, a customizable on-premises architecture designed specifically for research institutions.

LStore enables universities to **optimize blended cost per terabyte, avoid proprietary appliance lock-in, scale AI and GPU workloads predictably, and align storage tiers with workload I/O characteristics** rather than pre-defined appliance configurations.

Rather than continuing to rent infrastructure in perpetuity, institutions can invest in scalable systems that they own, control, and expand incrementally.

Challenge

AI and machine learning workloads are now a primary driver of research infrastructure growth. GPU clusters generate large volumes of training data, checkpoints, and intermediate results. These workloads introduce sustained high-throughput I/O, metadata pressure, and multi-tenant concurrency that traditional capacity tiers were not designed to support.

Cloud storage can absorb short-term demand, but long-term economics become difficult to justify. Frequent read operations, cross-region transfers, and data rehydration charges compound over time, particularly in iterative AI training cycles. At multi-petabyte scale, annual operating expense can exceed the cost of building on-premises infrastructure, without contributing to institutional asset ownership.

At the same time, many universities operate fragmented storage environments. Performance tiers, capacity tiers, and archive systems are often managed separately. Moving data between them introduces operational friction, slows research workflows, and increases administrative overhead.

Legacy appliance platforms add another constraint. Expansion typically requires large capital increments, proprietary hardware, and recurring software licensing commitments. Scaling becomes a political and budgetary exercise rather than a technical decision.

The result is a common institutional tension: data growth is predictable and continuous, but storage architecture and funding models are not.

Solution

LSTORE (LOGISTICAL STORAGE) BY SOURCECODE

In environments such as Vanderbilt's ACCRE deployment, LStore has been used as a capacity-optimized storage tier for large research datasets.

SourceCode now offers a similar LStore architecture, a hardware-agnostic, scale-out storage architecture designed for AI-driven research computing environments.

Rather than deploying fixed-function storage appliances, LStore is built on commodity storage servers and structured as a tiered, workload-aligned system. It supports high-throughput GPU pipelines, large parallel file workloads, and cost-optimized bulk capacity within a unified design framework.

CLOUD

\$10mi

20PB

0% of funds applied to hardware

Raw capacity cost – \$0/TB

Usable cost – \$280/TB



LSTORE

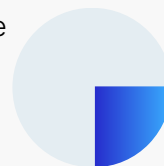
\$1.4mi

20PB

71% of funds applied to hardware

Raw capacity cost – \$48/TB

Usable cost – \$63/TB



ARCHITECTURE MODEL

LStore aligns storage tiers to workload characteristics:

Tier 0 – Server-Local NVMe

GPU- or CPU-attached NVMe for ultra-low-latency training and checkpointing.

Tier 1 – High-Performance Shared Storage

Commodity scale-out storage nodes designed to sustain parallel file system throughput, high metadata concurrency, and multi-project access patterns.

Tier 2 – Capacity-Optimized LStore Core

Multi-petabyte, cost-efficient disk-based storage for large files, research datasets, and sustained throughput workloads.

Tier 3 – Object / Archive (Optional)

S3-compatible storage for long-term retention, collaboration, and publishing workflows.

LStore deployments also integrate a global namespace and data orchestration layer, such as Hammerspace, to unify performance, capacity, and archive tiers under a single data environment. This enables transparent data movement across GPU-local, shared, and capacity tiers without application-level reconfiguration.

This tiered structure enables data placement based on I/O profile rather than historical purchasing decisions. Data can move between tiers based on performance requirements, cost sensitivity, and retention policy.

ECONOMIC MODEL

LStore is designed for incremental, predictable expansion. Capacity can be added in smaller petabyte-scale increments without appliance replacement cycles or forced refresh events. The model prioritizes hardware investment over recurring licensing overhead and avoids proprietary lock-in.

In prior LStore deployments, including large research environments, usable storage costs have been approximately \$63 per TB, with incremental 1PB expansion costs under \$60,000 depending on drive configuration and density. At multi-petabyte scale, this model has demonstrated the potential to reduce blended cost compared to sustained cloud operating expense for active research datasets.

DESIGN PRINCIPLES

LStore is built around:

- Commodity server infrastructure
- Horizontal scalability
- Metadata partitioning and high concurrency support
- Encryption at rest
- Optional global namespace integration for tier orchestration

The architecture is designed to support sustained multi-petabyte growth without requiring forklift upgrades or disruptive data migrations.

SourceCode engineers work with customers to co-design each deployment based on GPU density, expected dataset growth, concurrency profile, backup policy, and budget

Results

In the ACCRE environment at Vanderbilt University, the LStore architecture enabled the institution to expand multi-petabyte storage capacity while controlling cost growth associated with traditional storage platforms.

In environments where LStore architectures have been deployed, such as the ACCRE research environment, measurable improvements in cost structure and scalability have been observed. By placing large research datasets on commodity-based, capacity-optimized tiers and reserving high-performance media for latency-sensitive workloads, institutions can significantly lower blended cost per terabyte. In comparable environments, cost reductions of 40 to 50 percent have been achieved while sustaining required throughput for AI and simulation workloads.

GPU nodes retain direct access to NVMe for active training and checkpointing, while shared tiers support high-concurrency parallel workloads. This reduces checkpoint latency and minimizes pipeline stall time during iterative model training.

Incremental petabyte-scale expansion removes the need for large capital refresh cycles. Capacity can be added in smaller units, aligning infrastructure growth with research demand and funding cycles. By avoiding proprietary storage appliances and relying on commodity hardware, institutions maintain control over procurement, lifecycle management, and long-term scaling strategy.

Conclusion

For institutions planning sustained AI growth, storage architecture directly impacts throughput, cost predictability, and long-term operational flexibility. LStore provides a scalable, hardware-agnostic alternative. By aligning storage tiers to workload behavior and enabling incremental expansion on commodity infrastructure, universities can reduce long-term cost while maintaining performance for GPU and parallel workloads.

Evaluate your current storage model against an Lstore architecture. [Talk to a SourceCode storage expert.](#)