

Hammerspace Tier 0

COMPETITIVE BRIEF

Introduction

Hammerspace Tier 0 turns GPU server-local NVMe into high-performance shared storage, managed and protected by Hammerspace. Tier 0 storage is up to 10x faster than networked storage - so you can reduce checkpointing time for AI training and HPC, and improve response times for inferencing and agentic AI. It can be activated in hours - on-prem or in the cloud - with no clients or agents to install, so you can put GPUs to work immediately. And it enables optimal use of assets you already own so you can reduce the need for external flash storage - to reduce costs, and gain back the power and rack space those systems would otherwise consume.

This Competitive Brief provides a high-level comparison between Hammerspace Tier 0 and select competitors - Vast Data, Weka, and DDN with Lustre - followed by a more detailed side-by-side comparison against Weka’s Converged Mode offering.

	Hammerspace	Vast Data	Weka	DDN with Lustre
Local NVMe Option	Yes, Tier 0	No Tier 0 equivalent	Yes, Weka Converged Mode	Yes - Persistent Client Cache
Client Connectivity	No client to install - just Linux.	N/A	Requires Weka client	Requires Lustre client
Operational Complexity	Deploy on standard Linux without modification.	N/A	Complex, kernel module dependent	Complex, kernel module dependent
Use of Compute Resources	Minimal, runs in kernel space	N/A	High overhead on compute nodes	Low overhead, but kernel coupled
GPU Utilization Efficiency	Maximized via local data access, up to 10x faster than external storage	Bottlenecked by shared storage and external network	High	Variable; limited by Lustre client
Data Orchestration	Global namespace, policy-driven tiering to disk or NVMe	Limited only to Vast Data Platform	Limited, flash-to-object tiering on NVMe only	Requires custom scripting
Hybrid / Multi-Cloud Ready	Yes; single namespace across sites / clouds. Data is orchestrated based on policy	Limited. Vast clusters in the cloud lack performance and scale.	Data must be tiered to object and re-hydrated	No global namespace, no ability to orchestrate data between sites/clouds



Hammerspace Tier 0 Compared to Weka Converged Mode

This table provides a more in-depth, side-by-side comparison between Hammerspace Tier 0 and Weka Converged Mode.

Capability	Hammerspace Tier 0	Weka Converged Mode
Architecture Model	Disaggregated metadata/data	Metadata + data services co-located on compute nodes
Client Connectivity	No client to install - just Linux.	Requires Weka Client
Compute Resource Consumption	Lightweight client; metadata services run elsewhere	CPU/memory overhead; runs in user space
GPU Utilization Efficiency	High	High
Data Orchestration	Global namespace, policy-based data movement	No built-in data movement or orchestration across nodes
Hybrid / Multi-Cloud Readiness	Cloud-native, supports hybrid with unified namespace	Must tier data into S3 and then rehydrate for use in the cloud.
Data Locality Optimization	Automatically places data near jobs (AI/ML aware)	Static; tied to filesystem distribution logic
Minimum Deployment	Any number of GPU servers	WEKA will require 128 GPU servers minimum or a deal size of \$300K TCV.

Summary

Hammerspace Tier 0 is the only option for local NVMe that is standards-based, does not require a proprietary client, and minimizes compute resource consumption in order to maximize GPU utilization and accelerate AI pipelines on-prem and in-cloud.

