

Accelerating AI Adoption at Scale:

Neysa Networks Achieves Speed, Safety, and Efficiency with WEKA

India's fastest-growing AI Cloud provider eliminated storage bottlenecks and achieved enterprise-grade reliability by using NeuralMesh™ by WEKA.

Neysa Networks is on a mission to accelerate AI adoption by delivering AI-native infrastructure built for speed, control, and scale. Through Velocis, their AI Acceleration Cloud, Neysa helps enterprises and startups move from idea to scaled deployment, faster, safer, and more cost-efficiently.

Operating thousands of GPUs including NVIDIA H100s and H200s, Neysa supports customers processing over 100 million tokens daily. But as Neysa scaled to serve customers ranging from innovative startups to major enterprises, their infrastructure needed to deliver enterprise-grade reliability without compromise.



The Challenge: Growing Demands Outpaced Open-Source Capabilities

As AI workloads scaled, Neysa observed that different workload categories required different storage characteristics. Ceph continued to serve object storage and broader platform needs effectively. However, high-performance GPU clusters demanded tighter latency consistency, stronger isolation, and predictable throughput under bursty traffic patterns.

To meet these requirements, Neysa evolved to a dual-layer approach, retaining open-source storage where appropriate, while integrating a specialized high-performance storage layer optimized for distributed AI workloads.

“Storage is not one-size-fits-all,” noted Anindya Das, Co-founder and CTO, Neysa Networks. “Our goal is to align storage architecture with workload behavior, so GPUs remain productive and customers experience consistent performance.”

The Solution: NeuralMesh by WEKA Delivers Performance, Flexibility, and Peace of Mind

Neysa evaluated multiple storage platforms against real AI workload characteristics—extreme parallelism, metadata-intensive pipelines, mixed IO profiles, and synchronized checkpoint bursts—rather than synthetic benchmarks or steady-state assumptions.

After detailed architecture reviews and production workload testing, Neysa selected NeuralMesh by WEKA as a complementary high-performance data layer purpose-built for GPU-intensive environments.

“NeuralMesh’s architecture is built specifically for GPU-driven, parallel AI workloads,” said Das. “It wasn’t about replacing Ceph. It was about adding a storage layer designed for deterministic behavior under GPU-driven pressure.”

Critical differentiators included:

- **Software-defined flexibility:** NeuralMesh could run on both HPE and Dell hardware, preventing vendor lock-in.
- **Enterprise-grade support:** WEKA provided the professional support structure Neysa needed to serve demanding customers with confidence.
- **Rapid deployment:** When Neysa faced tight customer delivery timelines, WEKA worked with the team to repurpose existing Dell hardware and implemented the solution in just two days.
- **Performance for AI workloads:** NeuralMesh delivered low and consistent latency under mixed read/write workloads, with near-linear scaling aligned with high-speed GPU networks—a unified file and object platform across the entire AI lifecycle.

The engineering engagement proved remarkably smooth, with Neysa operating **without a single support ticket for six months**.

The Impact: Enabling AI Innovation Without Limits

"NeuralMesh has been rock-solid, allowing our team to focus on building AI infrastructure rather than troubleshooting storage issues," said Das. "That reliability gives us the confidence to take on our most demanding customers."

The results? **Zero critical issues for six months, higher effective GPU utilization with shorter training cycles** across multiple hardware platforms, with **fewer stalled or failed jobs**.

Looking Forward: Pioneering AI Infrastructure Innovation

As Neysa continues its rapid expansion, the company is piloting WEKA’s Axon technology, positioning themselves as a lighthouse customer for cutting-edge storage innovation. WEKA provides the AI-ready storage foundation, while Neysa integrates compute, orchestration, observability, security, and cost control into a production-grade AI infrastructure system—enabling customers to move from idea to scaled deployment without compromise.