

NeuralMesh™ Powers Sovereign AI Data Storage at National Scale

NeuralMesh™ ends the tradeoff between sovereignty and performance, feeding GPUs at AI scale across on-premises, classified, and hybrid environments while data stays within national borders.

Challenges

- Hyperscaler storage routes data through non-sovereign control planes
- Multi-tenant performance degradation
- Overprovisioning inflates cost and power
- Weak checkpointing puts weeks of training at risk of hardware failure
- Legacy NAS starves GPU clusters

Solution

NeuralMesh™ runs as a distributed, shared-nothing architecture with no centralized controllers. Every node serves data, metadata, and I/O in parallel across classified on-premises, cloud, and hybrid environments. **NeuralMesh Axon™** pools NVMe already inside GPU servers into a single storage cluster, activating existing capacity instead of purchasing a separate tier. **Augmented Memory Grid™** extends the key-value cache beyond GPU memory, so tokens aren't reloaded from storage.

Benefits

- Federate data across classified and authorized cloud networks
- Encrypt data end-to-end with zero operator access
- Meet data sovereignty mandates
- Drive GPU utilization past 90%
- Host data ingestion, training, and inference in a single platform

Put an end to the sovereign AI storage vs. performance tradeoff

National governments, government-backed cloud operators, and NVIDIA Cloud Partners are racing to build homegrown AI capability without surrendering control of their data. That control has to hold across classified networks, under end-to-end encryption, and under direct operational oversight, while still hitting GPU-scale performance. Most AI infrastructure asks programs to choose one or the other.

NeuralMesh answers with one platform. Everything is encrypted, access-controlled, and deployable on-premises, in classified networks, and across hybrid and multicloud environments. Neuralmesh supports multiple namespaces in a single physical cluster, spanning ingestion through inference, unifying AI programs and multi-tenant clouds under a single storage infrastructure.

"NeuralMesh enabled us to deploy composable storage leveraging an operator-driven architecture that delivers secure multi-tenancy across our AI cloud infrastructure."

— Philip Lin, CEO, YTL AI Cloud

