

NeuralMesh™ Accelerates Quant Research and Backtesting

NeuralMesh™ delivers the low latency and extreme concurrency quantitative trading teams need to generate more trading strategies, backtest them faster, and act before markets move

Challenges

- Slow iteration limits strategies
- Latency drains trading profit
- Idle GPUs waste capital daily
- Competitors can access the same data
- Legacy storage stalls backtests
- Tiering forces data tradeoffs
- Tick archives strain budgets

Solution

NeuralMesh™ is a fully distributed, shared-nothing infrastructure where every server handles data, metadata, and I/O directly, with no dedicated metadata nodes and no fixed data paths.

NeuralMesh continuously rebalances data and metadata across the cluster, holding predictable, ultra-low latency as more quants, models, and nodes come online, even through node failures.

NeuralMesh Axon™ runs directly on the NVMe already inside your GPU servers, and native multi-protocol access moves data from ingestion to compute with no copies.

Benefits

- Run more simulations at once
- Reproduce backtests point in time
- Move before the market does
- Reclaim idle GPU capacity
- Scale without new hardware
- Ultimately: **make more money**

Turn research velocity into profit for quantitative trading

When **70% of global hedge funds now deploy machine learning models** and all firms have access to the same tick, time-series, and alternative data, quantitative trading alpha comes from research velocity, not information. And when storage is slow, iteration cycles shrink and teams test fewer strategies than better-resourced rivals.

Execution latency matters, but the durable edge against other trading firms is the time spent testing strategies. With NeuralMesh, research cycles run in minutes and hours, not hours and days. This allows firms to move from research to production seamlessly whether workflows run on-premises, cloud, or hybrid. NeuralMesh removes the data replication overhead of legacy NAS and holds ultra-low latency through peak market activity, extracting more from the infrastructure you already run so research keeps compounding.

NeuralMesh **Augmented Memory Grid™** drives concurrency across hundreds of signal, risk, and backtesting models at once, so expensive GPUs and CPUs stay saturated. Quants pull data on demand with no manual tuning and no wait on IT, so research teams move at their own pace.

*“Our existing Network Attached Storage was too slow to effectively cover analytics workloads. **WEKA proved to be 10× faster** and allowed for seamless backup to our object store.”*

— WEKA Capital Market Customer