

NeuralMesh™ Powers Real-Time Sports and Prediction Markets

NeuralMesh™ delivers the ultra-low latency, high performance storage infrastructure to process data from live events, run prediction models faster, and scale user bases without downtime

Challenges

- Massive datasets combined with latency and performance requirements
- Slow checkpoints kill live platforms
- Storage fails under live event load
- Downtime loses revenue instantly
- Hard to quickly scale into new markets
- Data replication wastes capacity

Solution

NeuralMesh provides ultra low latency and high performance for real-time models to keep platforms up during peak loads.

NeuralMesh replaces legacy distributed storage with a high-performance, software-defined architecture built for hot, always-on workloads. NeuralMesh Augmented Memory Grid™ enables concurrency across hundreds of live models simultaneously without storage becoming the bottleneck.

Benefits

- Track odds, bets, and players live
- Ensure prediction markets reflect real-world outcomes to act before customers participate in the exchange
- Keep platforms live during peak events
- Scale into new sports, markets without rebuilding infrastructure from scratch
- Eliminate idle GPU and CPU waste

Storage built for always-on, high-stakes market platforms

Online prediction market platforms, from sportsbook operators like DraftKings, FanDuel, BetMGM, Caesars Sportsbook, and Kaizen Gaming, to more holistic prediction exchanges like Kalshi and Polymarket, run storage infrastructure stacks that require ultra low-latency and high levels of performance because every live event triggers a surge of reads, writes, and model updates that must be ingested and analyzed in milliseconds to ensure ever-changing odds are positioned for optimal profitability.

NeuralMesh eliminates the replication overhead of legacy systems, checkpointing live applications fast enough to sustain availability through peak loads. NeuralMesh's Augmented Memory Grid™ drives concurrency across live odds engines, fraud detection, and risk models simultaneously, without manual tuning. As platforms are stress tested and expand, the network grows stronger due to the distributed architecture. As platforms expand into new sports and geographies, NeuralMesh scales capacity and performance independently across on-premises and cloud deployments, with no architectural rework required.

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

— WEKA Capital Market Customer