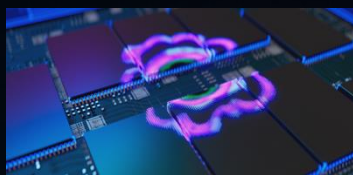


MEXT Performance Benchmarking

How MEXT AI-Powered Predictive Memory Works



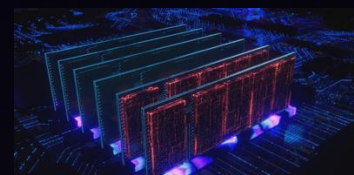
MEXT continually offloads cold memory pages from DRAM to a 20x lower-cost tier, Flash



MEXT's AI engine tracks related pages and predicts which pages in Flash are likely to get requested by the application



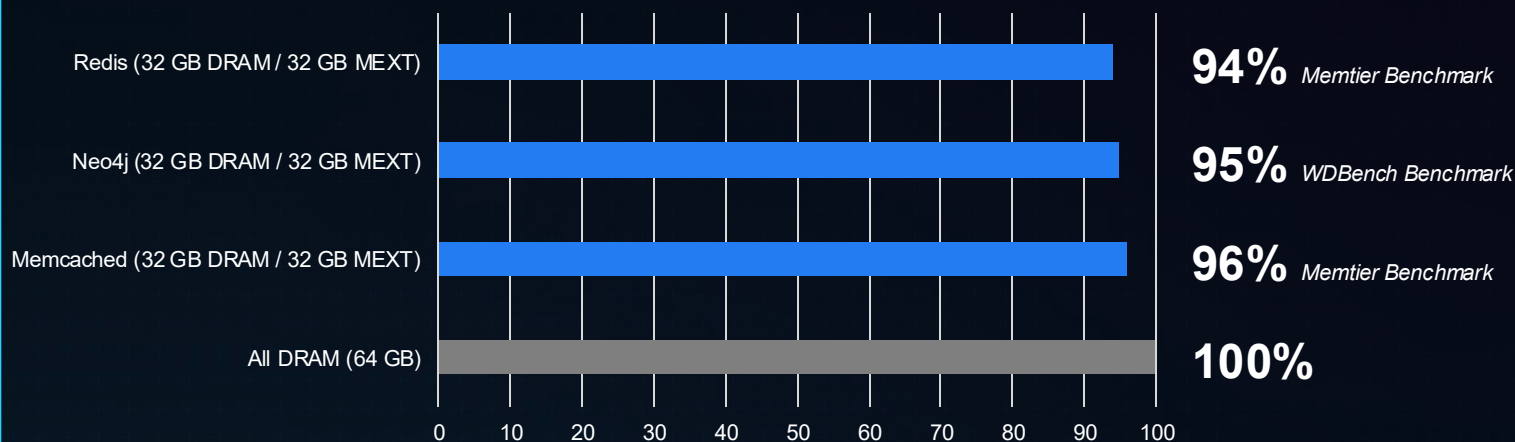
MEXT transparently pushes those pages back into DRAM—even before they are requested by the application



From the application's point of view, all relevant pages can be found in DRAM—maintaining performance within a smaller DRAM footprint

MEXT Performance vs. All DRAM System

With MEXT, customers can achieve comparable application performance using 50% of the DRAM of their original configuration. For instance, when comparing 64 GB DRAM system to a 32 GB DRAM + 32 GB Flash-with-MEXT system, performance for Redis hits 94% of the all-DRAM configuration, Neo4j hits 95%, and Memcached hits 96%.



Performance / \$ with MEXT

If we look at the amount of performance achieved per dollar spent, we see that leveraging MEXT yields significant efficiencies.

For instance, if we take the Performance / \$ of an all-DRAM system to be a reference level of X, MEXT enables a Performance / \$ of 1.7X for Redis, 1.7X for Neo4j, and 1.7X for Memcached.

