



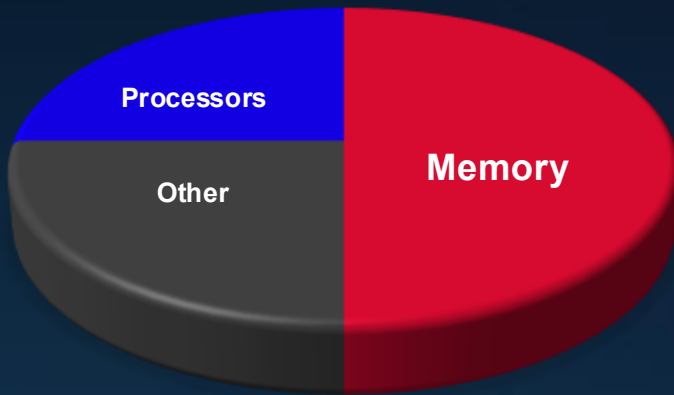
Smart Memory Tiering

Harnessing AI to radically transform the cost,
capacity, and availability of memory

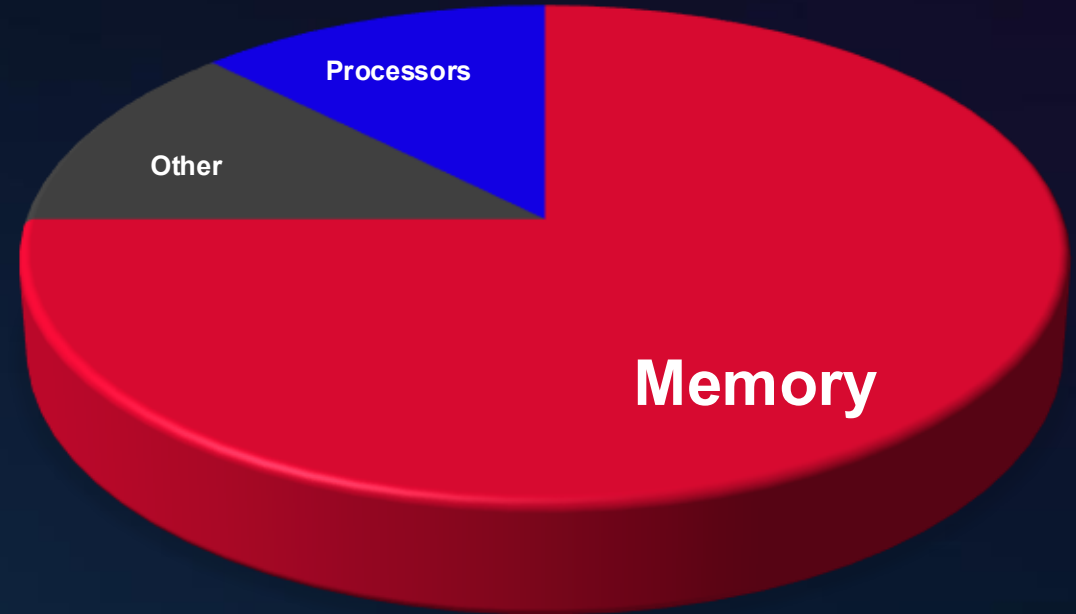
Gary Smerdon

Founder & CEO, MEXT
gary.smerdon@mext.ai

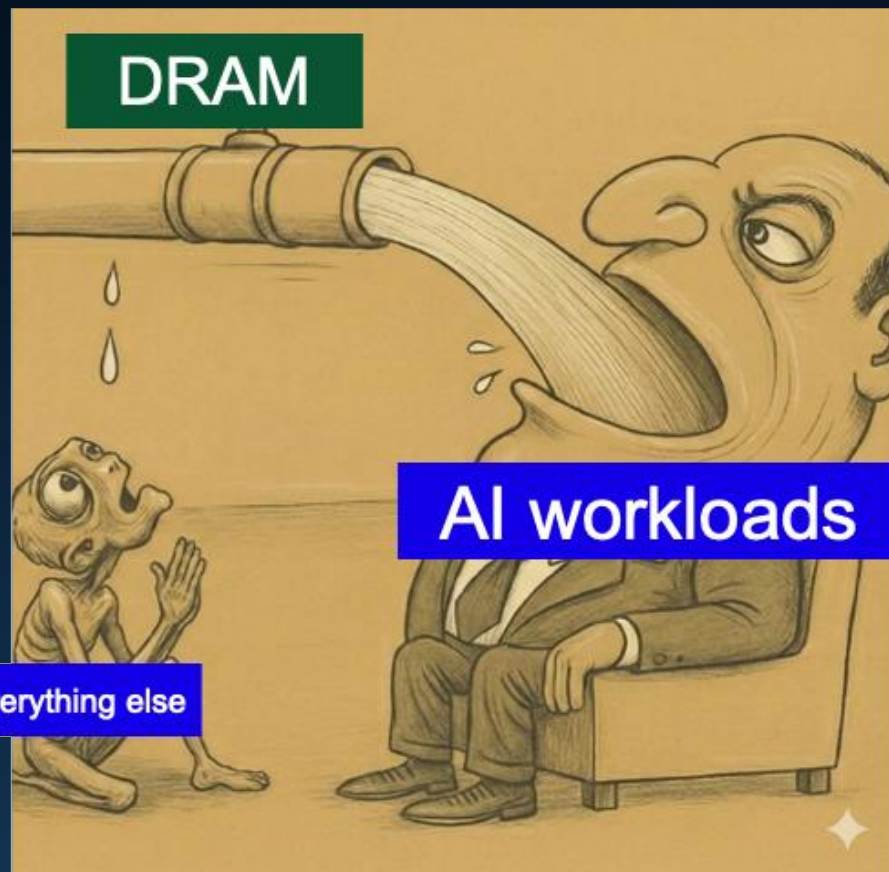
DRAM Has Become a Global Crisis



2023:
DRAM accounts for 50%
of cost of each server



2026:
DRAM accounts for 60-90%
of cost of each server



**RAM I WANT
TO BUY**



**RAM I CAN
BUY**



But Over Half of DRAM Is Wasted!

50%



Percentage of DRAM that
Half of VMs NEVER access

25%

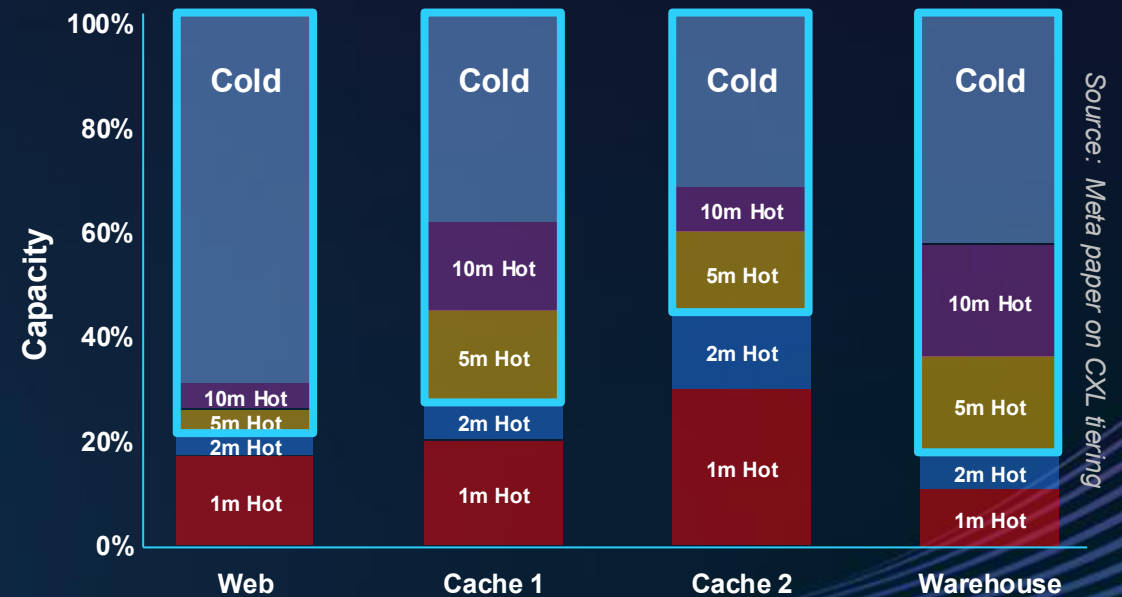


Stranded memory as a percentage
of total DRAM

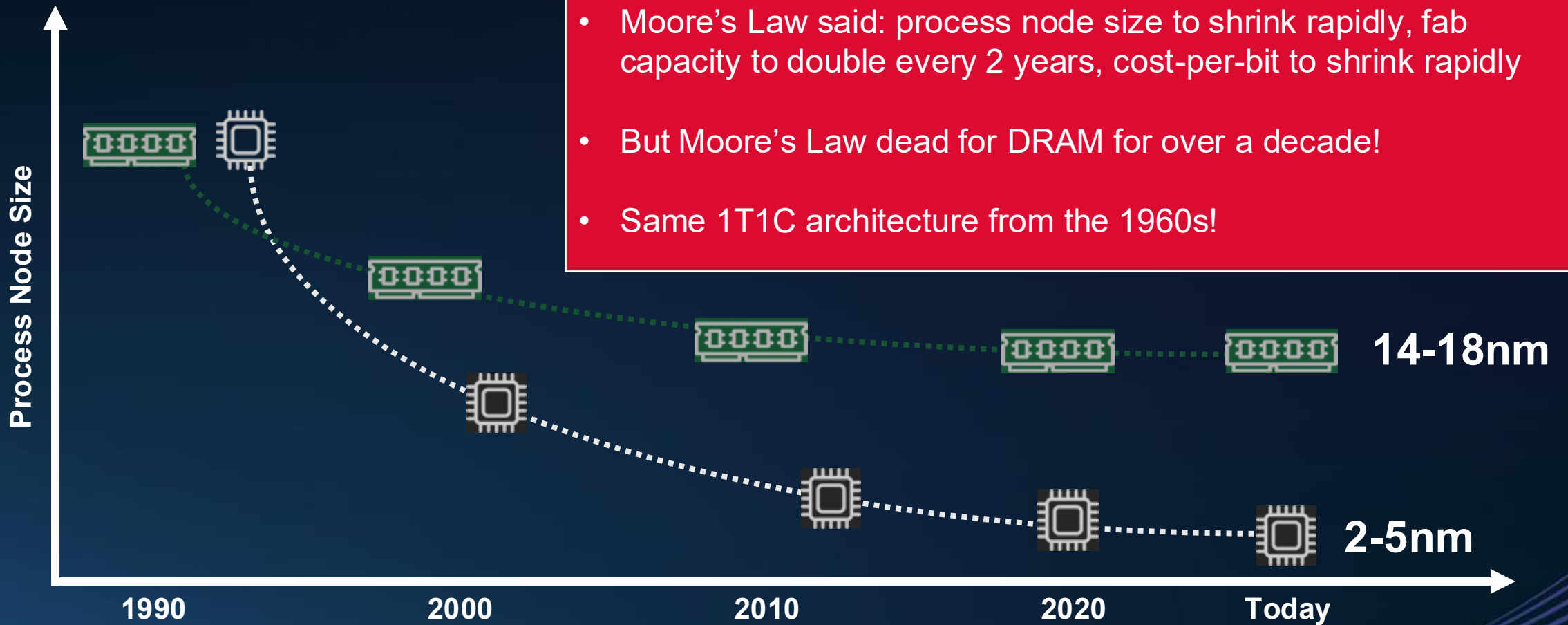
∞ Meta

50-80%

Percentage of memory pages that go cold
(representative of many environments!)



Why The DRAM Problem Isn't Going Away



How We Revolutionize Memory Cost and Capacity



Increase DRAM utilization

Eliminate overprovisioning
and stranded DRAM



No hardware or software changes

HW / architectural and
SW changes drive up cost



Bring flash into memory tier

Flash is 50X cheaper and
30X lower power vs DRAM

“Great idea, but it won’t work...”

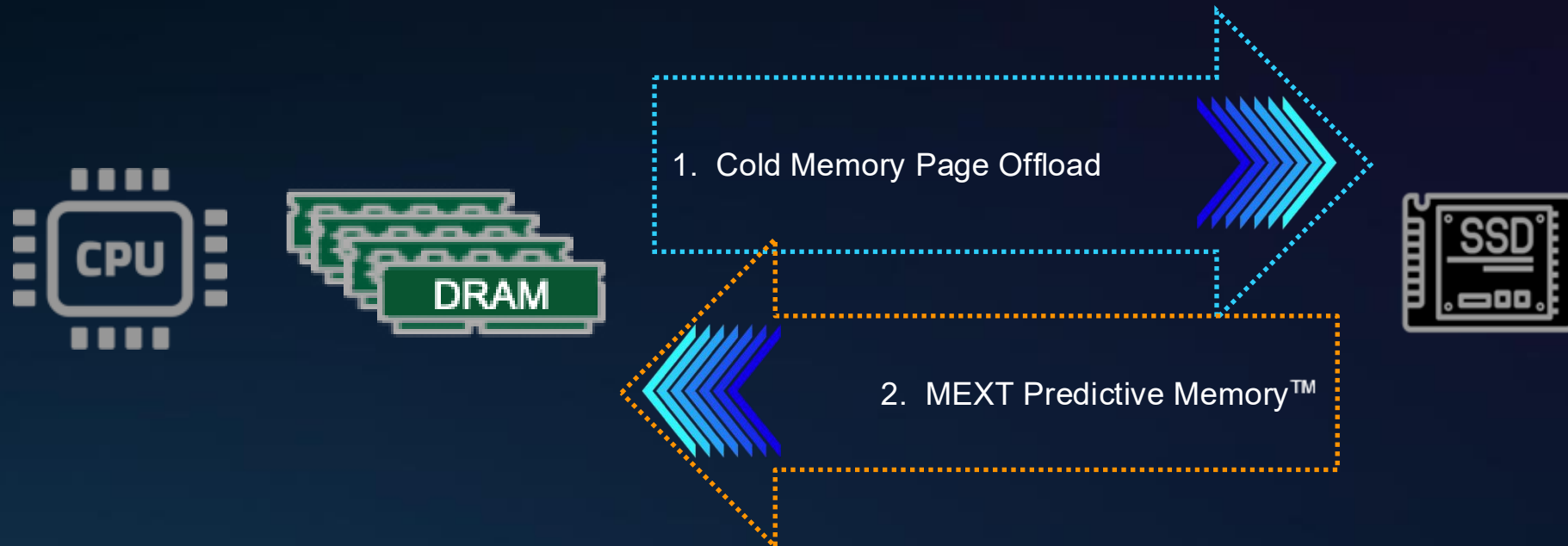
**Flash is 500X slower
than DRAM**

...until MEXT Predictive Memory™





MEXT Predictively & Transparently Tiers Memory



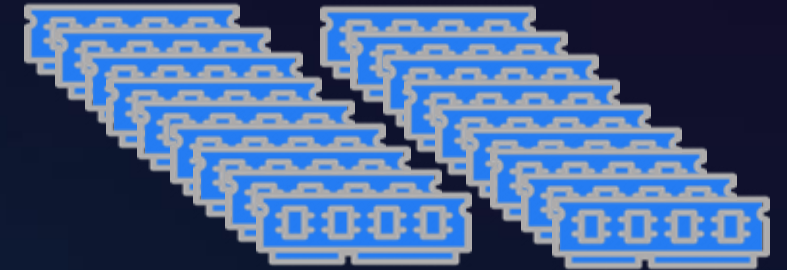
1. MEXT transparently moves cold pages into NVMe flash
2. MEXT AI engine predicts and pushes pages back before they are needed

MEXT Makes Flash Appear as Memory

Constrained memory



Expanded memory



MEXT Plug-In

- Lightweight, low-overhead Linux Driver

MEXT Predictive Memory™ Engine

- Ensemble of AI models enabling transparent auto-tiering of memory pages
- Continual training & inferencing using only a single CPU core

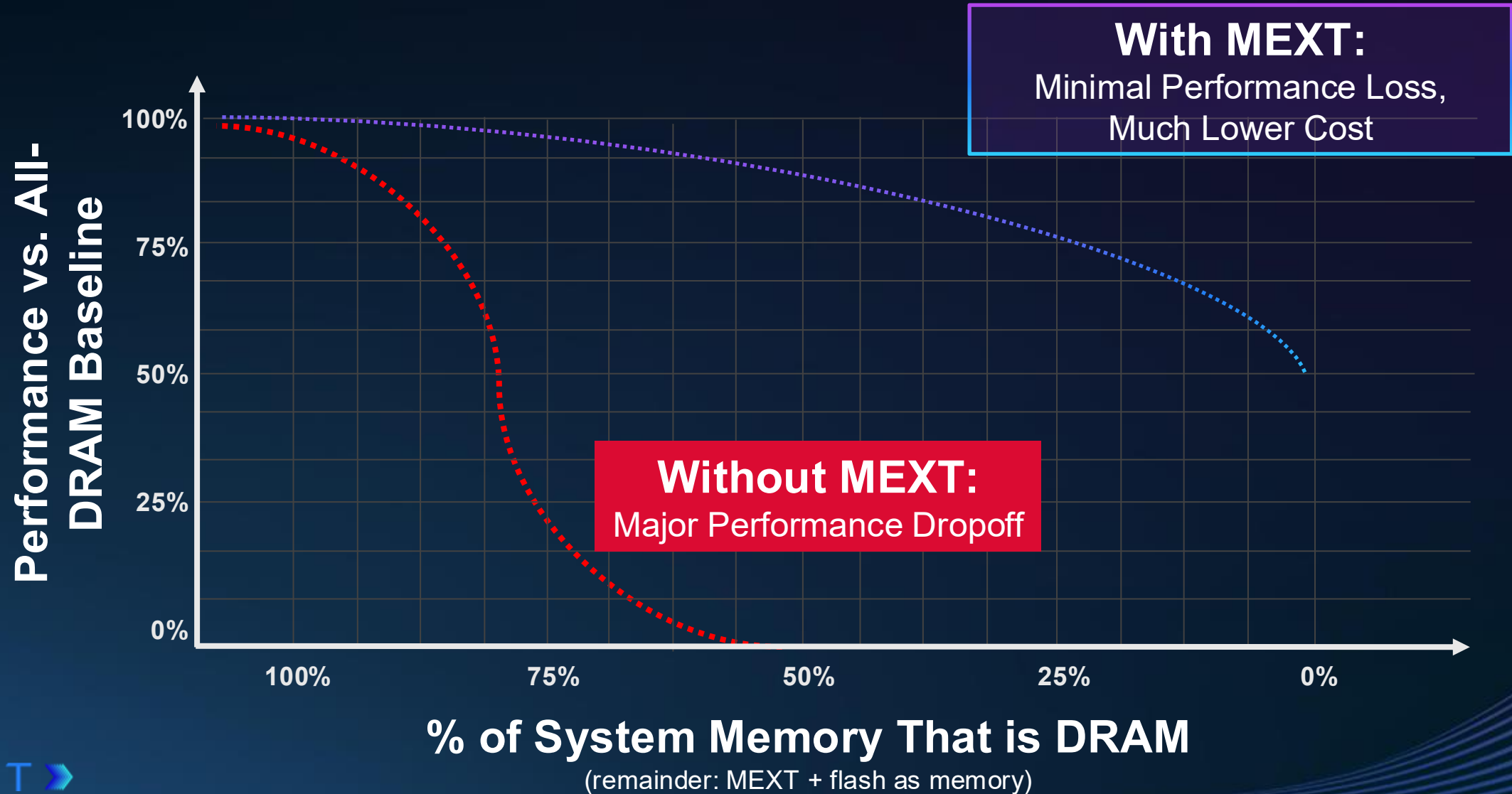
MEXT View™ Observability Tool

- Illustrates memory utilization and MEXT AI prediction accuracy using Open Telemetry (OTel)

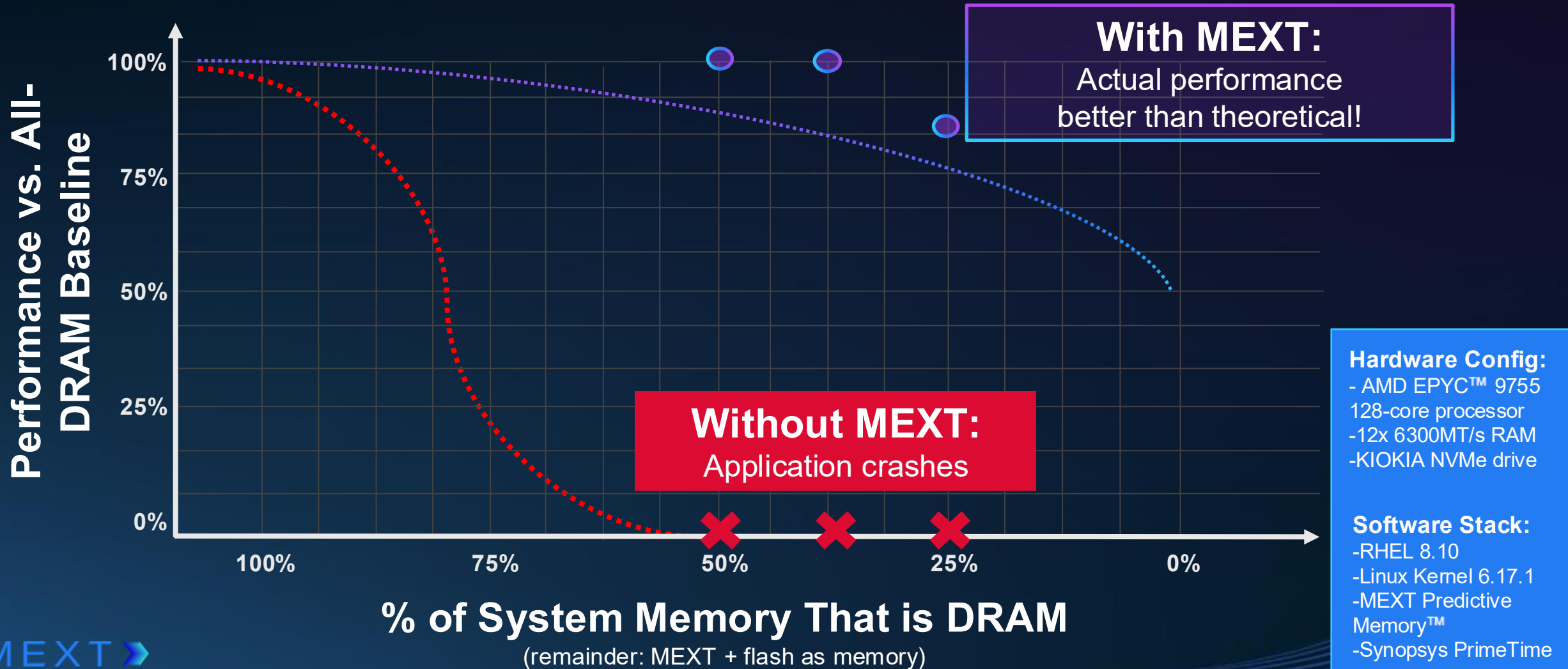
Revolutionizing Cost Savings On-Premises

Total Memory	Dell R6725	Dell R6725
	ALL DRAM	DRAM + MEXT Memory™
1.5TB	\$99K*	\$30K* 70% Savings
3TB	No Longer Available	\$68K Priceless
6TB	No Longer Available	\$111K Priceless

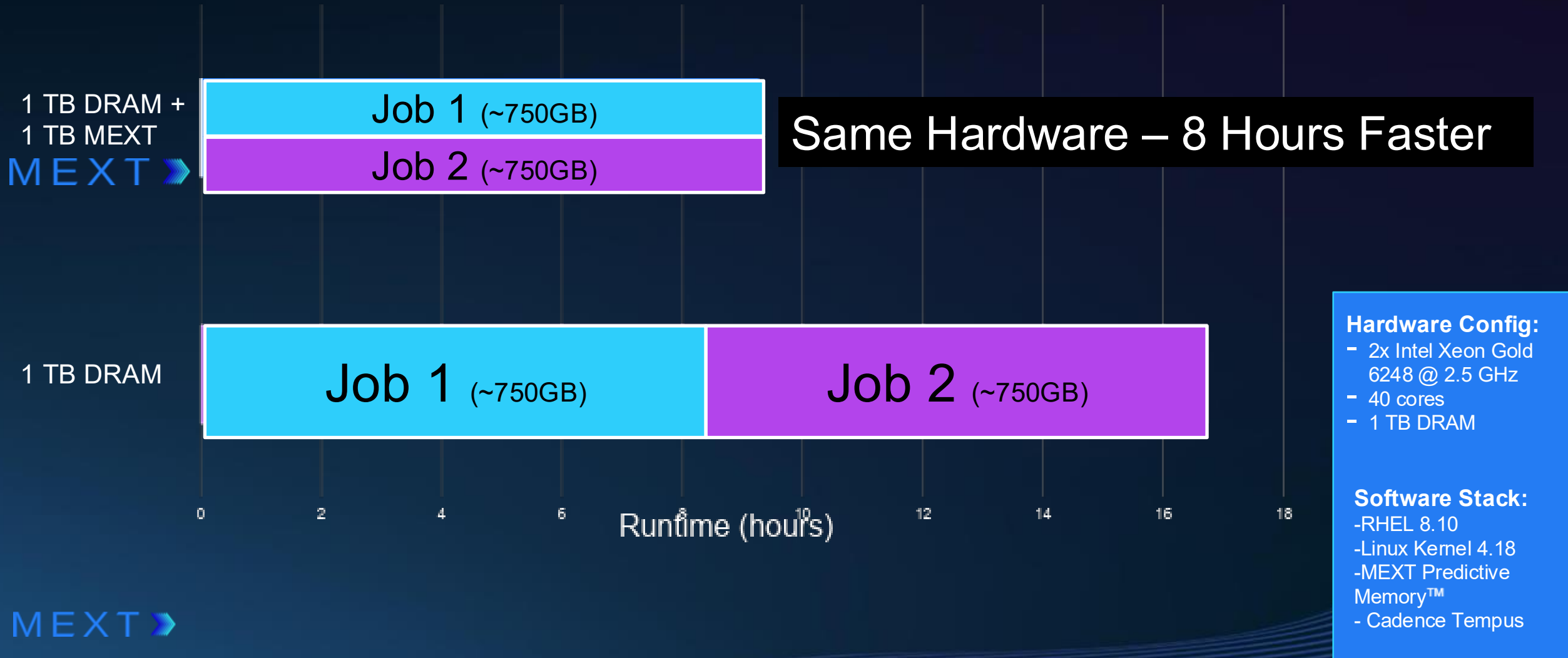
Theoretical Performance



Actual Performance: Synopsys PrimeTime (EDA)

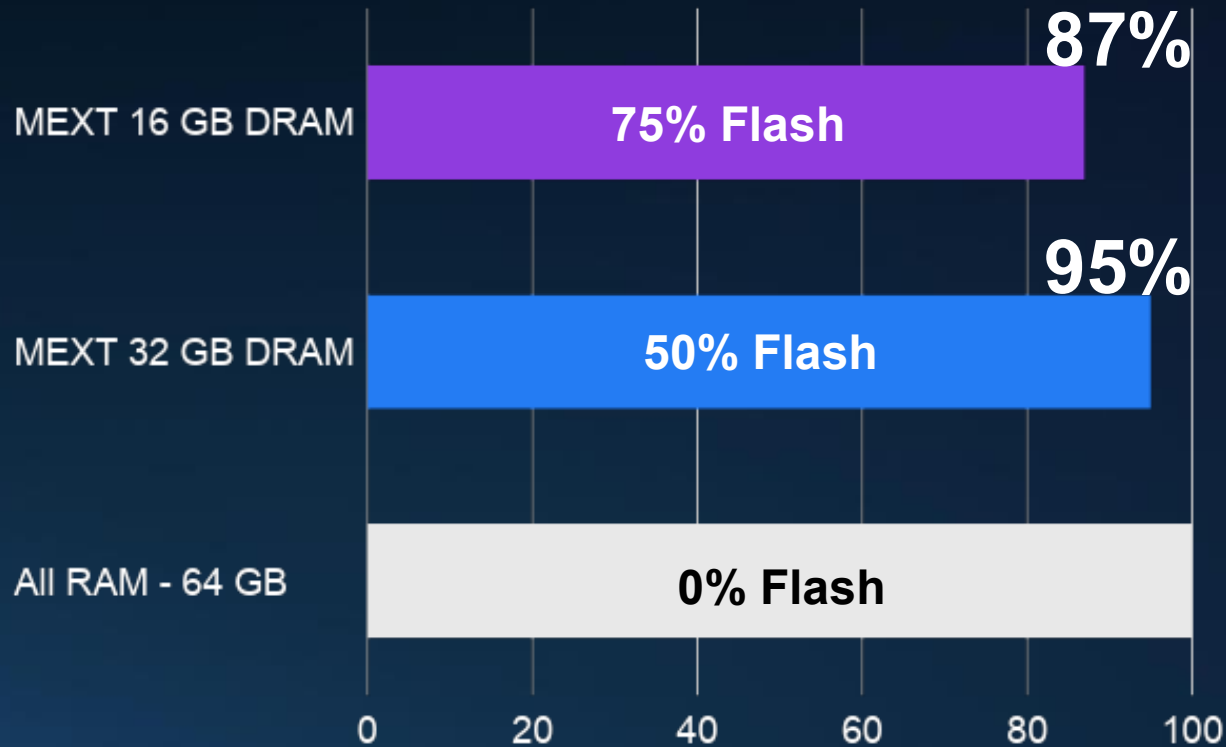


Cadence Tempus: 2 Jobs, 1.8x Faster

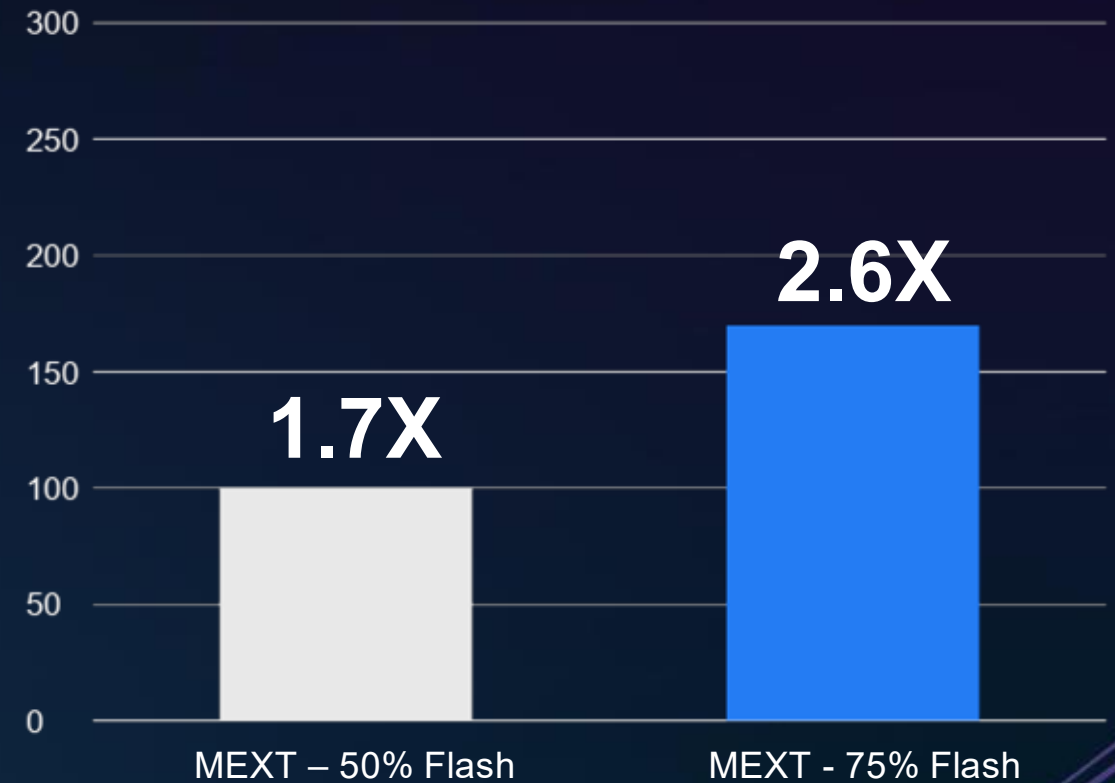


Improved Economics for Example: Neo4j

Neo4j (WDBench) – 64 GB Memory



Neo4j Performance / \$



MEXT Workload Examples

Analytics



Relational In-Memory Database



In-Memory Key/Caching Value



EDA



Graph Database



AI-Native Vector Database



AI Models



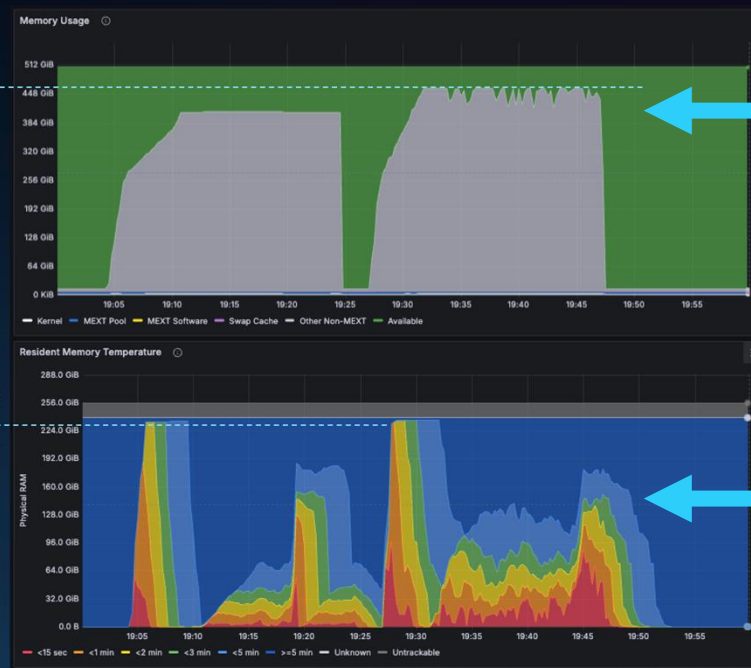
Animation Rendering



MEXT View™: Actual customer results

Significant Memory Usage (~450GB)

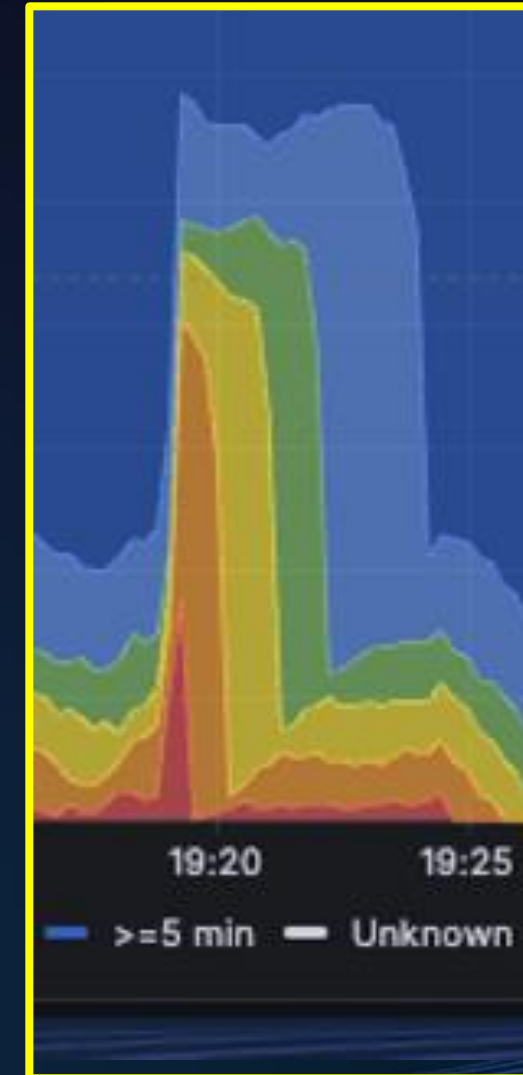
Memory is very active (~256GB in <1min)



512GB Memory
256GB DRAM + 256GB MEXT

256GB DRAM
Most is not used for “long” periods

MEXT View™: Actual customer results



> 5 minute
< 5 minute
< 3 minute
< 2 minute
< 1 minute
<15 seconds

In a Nutshell: MEXT Value

>2X Perf / \$

Delivering more bang for the buck
on-prem or in the cloud

Abundant Memory

“Software expands to fill available memory” (Wirth’s Law),
avoid sharding of large datasets for better performance

Lower Power

Basis in 30X lower-power
form factor (flash) vs DRAM

Server Life

More capacity on existing
servers, extending lifespan

Fast Deployment

Software-only, drops in,
any infra / configuration

Simplified SKUs

No lock-in to standard
memory sizes, scale easily



Fast Forward!