

POWER10 Processor Chip

Technology and Packaging:

- 602mm² 7nm Samsung (18B devices)
- 18 layer metal stack, enhanced device
- Single-chip or Dual-chip sockets

Computational Capabilities:

- Up to 15 SMT8 Cores (2 MB L2 Cache / core)
(Up to 120 simultaneous hardware threads)
- Up to 120 MB L3 cache (low latency NUCA mgmt)
- 3x energy efficiency relative to POWER9
- Enterprise thread strength optimizations
- AI and security focused ISA additions
- 2x general, 4x matrix SIMD relative to POWER9
- EA-tagged L1 cache, 4x MMU relative to POWER9

Open Memory Interface:

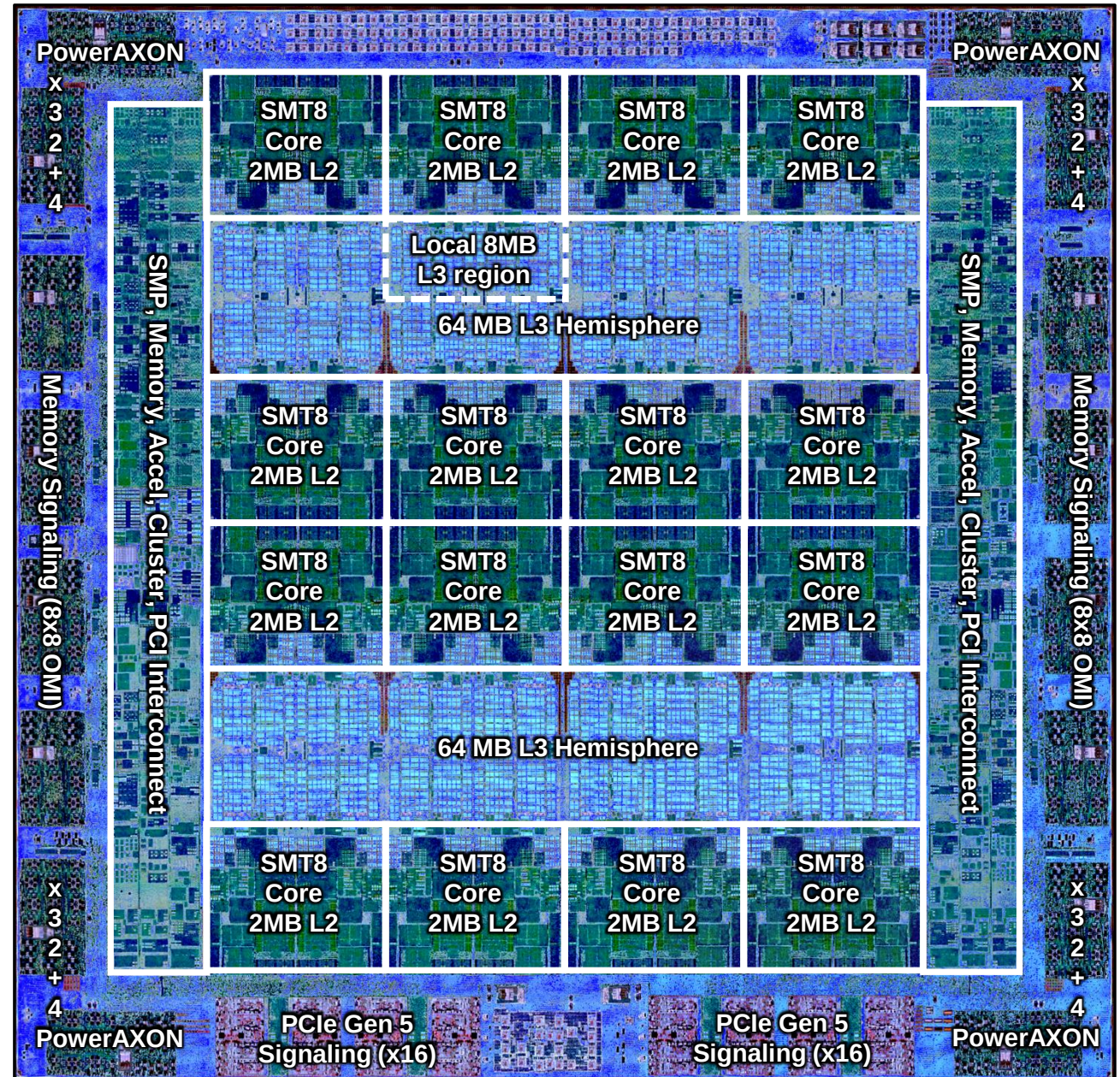
- 16 x8 at up to 32 GT/s (1 TB/s)
- Technology agnostic support: near/main/storage tiers
- Minimal (< 10ns latency) add vs DDR direct attach

PowerAXON Interface:

- 16 x8 at up to 32 GT/s (1 TB/s)
- SMP interconnect for up to 16 sockets
- OpenCAPI attach for memory, accelerators, I/O
- Integrated clustering (memory semantics)

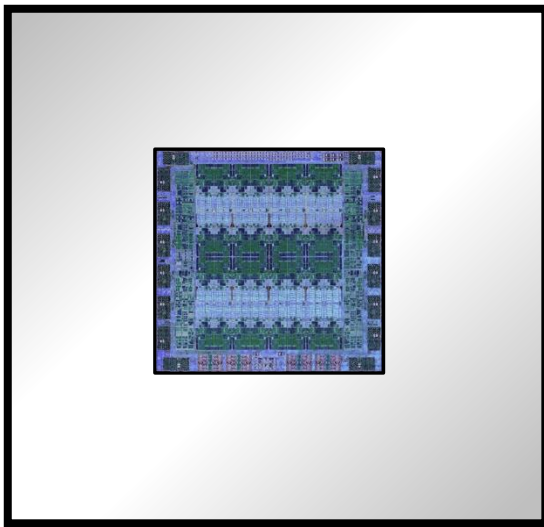
PCIe Gen 5 Interface:

- x64 / DCM at up to 32 GT/s



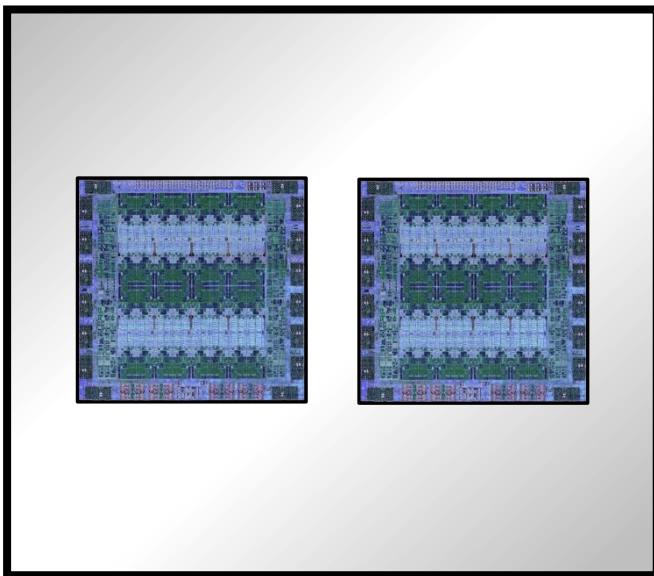
Die Photo courtesy of Samsung Foundry

Socket Composability: **SCM & DCM**



Single-Chip Module Focus:

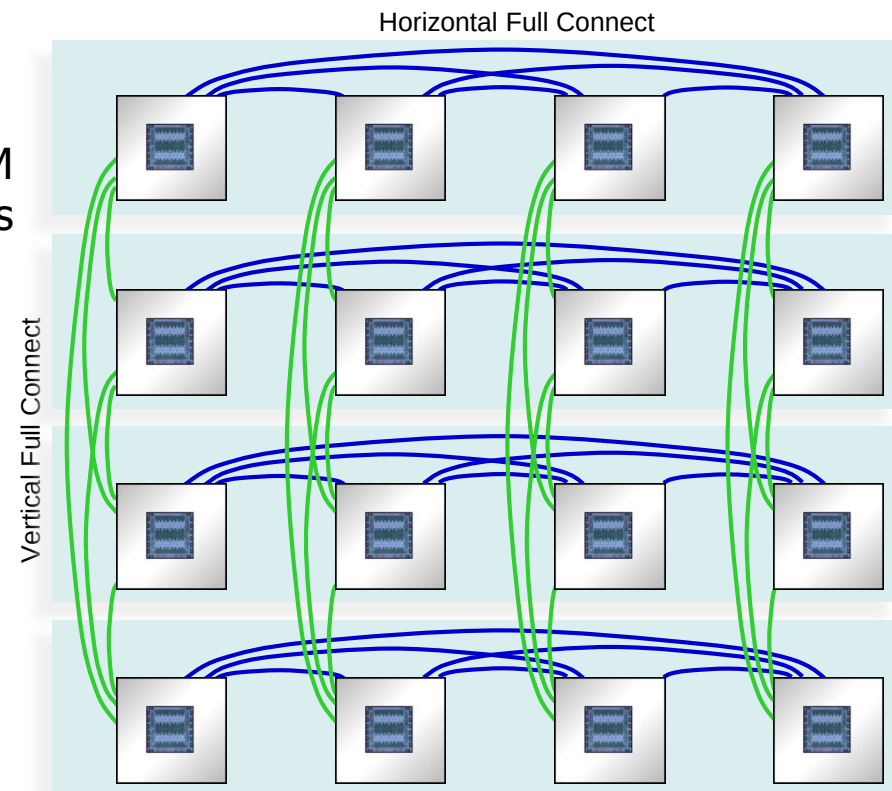
- 602mm² 7nm (18B devices)
- **Core/thread Strength**
 - Up to 15 SMT8 Cores (4+ GHz)
- **Capacity & Bandwidth / Compute**
 - Memory: x128 @ 32 GT/s
 - SMP/Cluster/Accel: x128 @ 32 GT/s
 - I/O: x32 PCIe G5
- **System Scale (Broad Range)**
 - 1 to 16 sockets



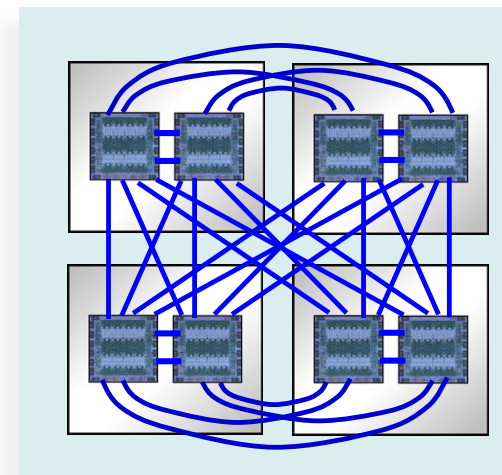
Dual-Chip Module Focus:

- 1204mm² 7nm (36B devices)
- **Throughput / Socket**
 - Up to 30 SMT8 Cores (3.5+ GHz)
- **Compute & I/O Density**
 - Memory: x128 @ 32 GT/s
 - SMP/Cluster/Accel: x192 @ 32 GT/s
 - I/O: x64 PCIe G5
- 1 to 4 sockets

Up to
16 SCM
Sockets



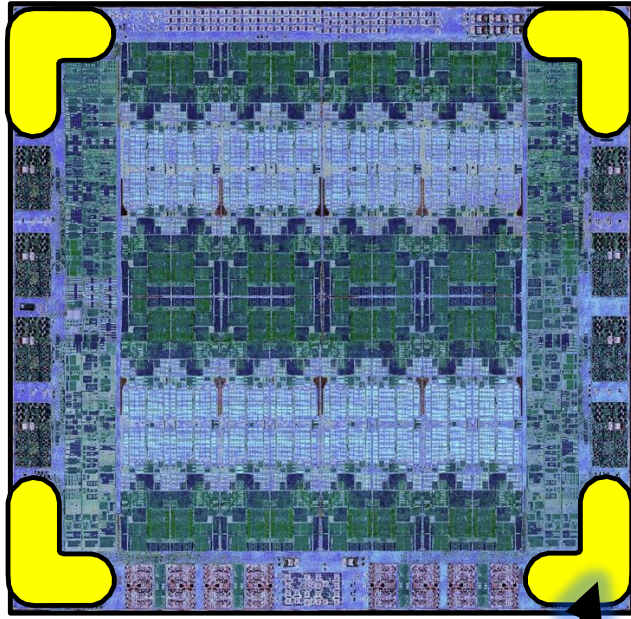
Up to
4 DCM
Sockets



IBM POWER10

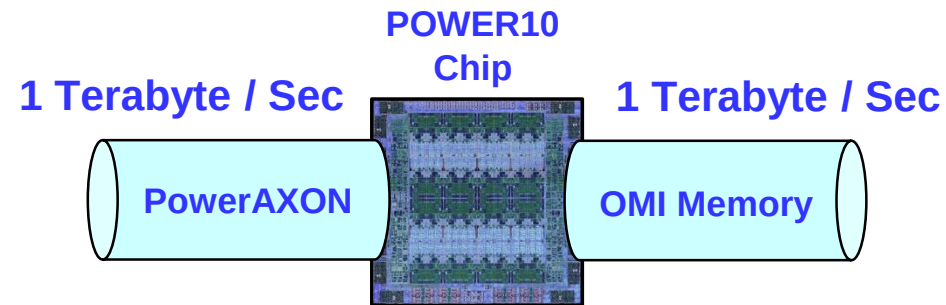
(Multi-socket configurations show processor capability only, and do not imply system product offerings)

System Composability: **PowerAXON & Open Memory Interfaces**

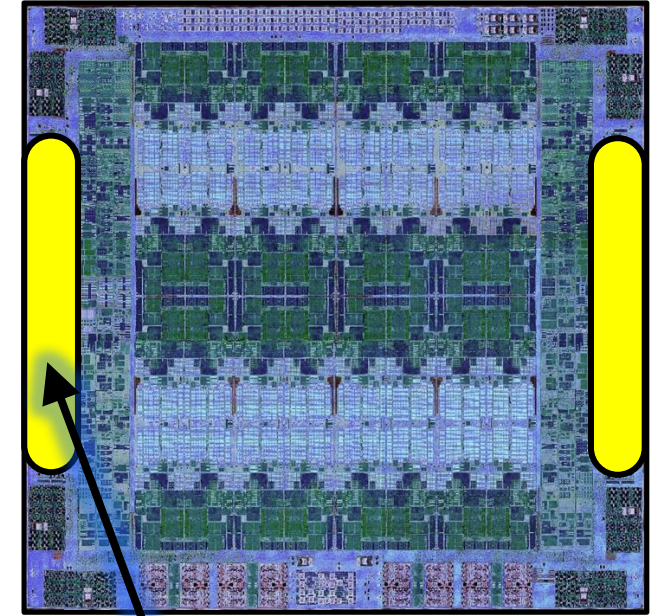


PowerAXON corner
4x8 @ 32 GT/s

Multi-protocol
“Swiss-army-knife”
Flexible / Modular Interfaces



Built on best-of-breed
Low Power, Low Latency,
High Bandwidth
Signaling Technology



OMI edge
8x8 @ 32 GT/s

6x bandwidth / mm²
compared to DDR4
signaling

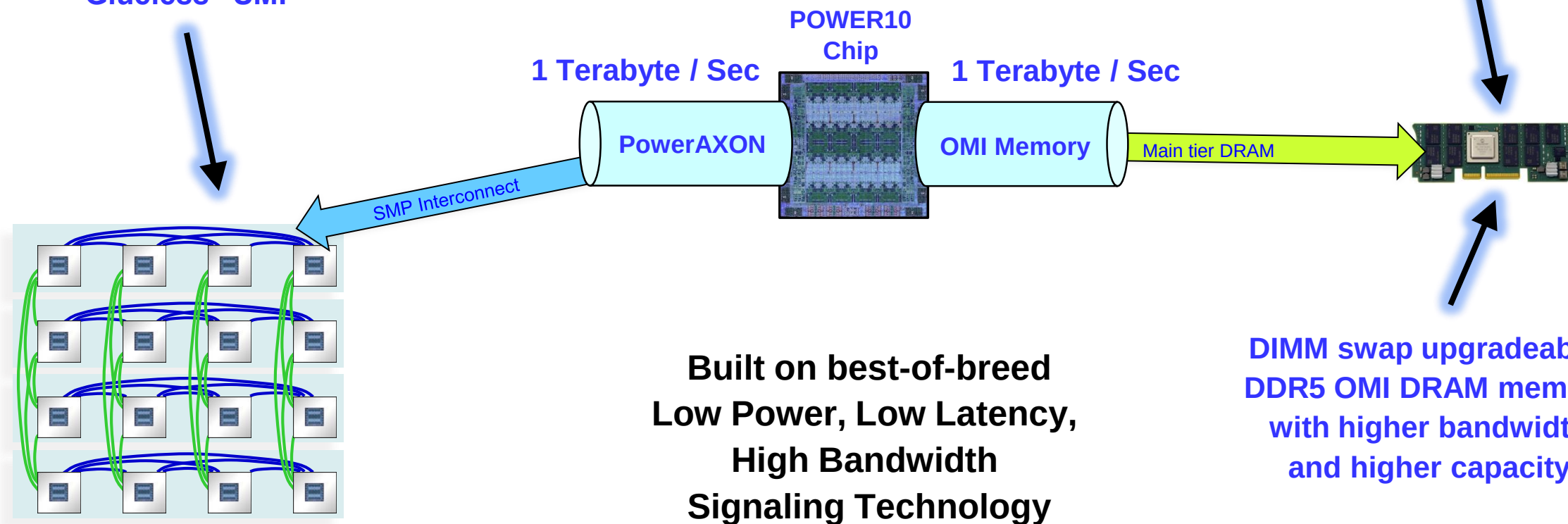
System Enterprise Scale and Bandwidth: SMP & Main Memory

Multi-protocol
“Swiss-army-knife”
Flexible / Modular Interfaces

Initial Offering:
Up to 4 TB / socket
OMI DRAM memory
410 GB/s peak bandwidth
(MicroChip DDR4 buffer)
< 10ns latency adder

Build up to 16 SCM socket
Robustly Scalable
High Bisection Bandwidth
“Glueless” SMP

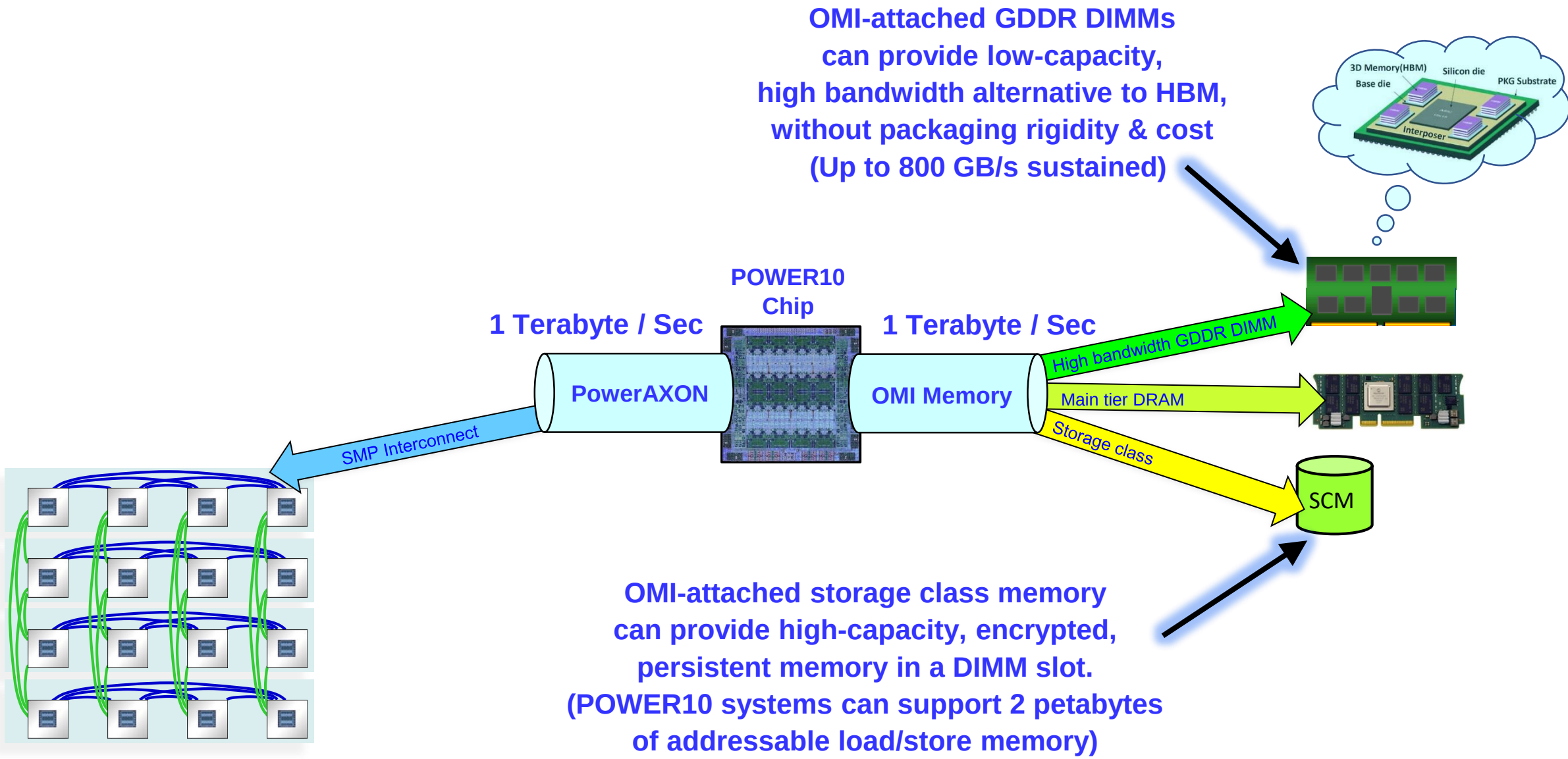
Allocate the bandwidth
however you need to use it



DIMM swap upgradeable:
DDR5 OMI DRAM memory
with higher bandwidth
and higher capacity

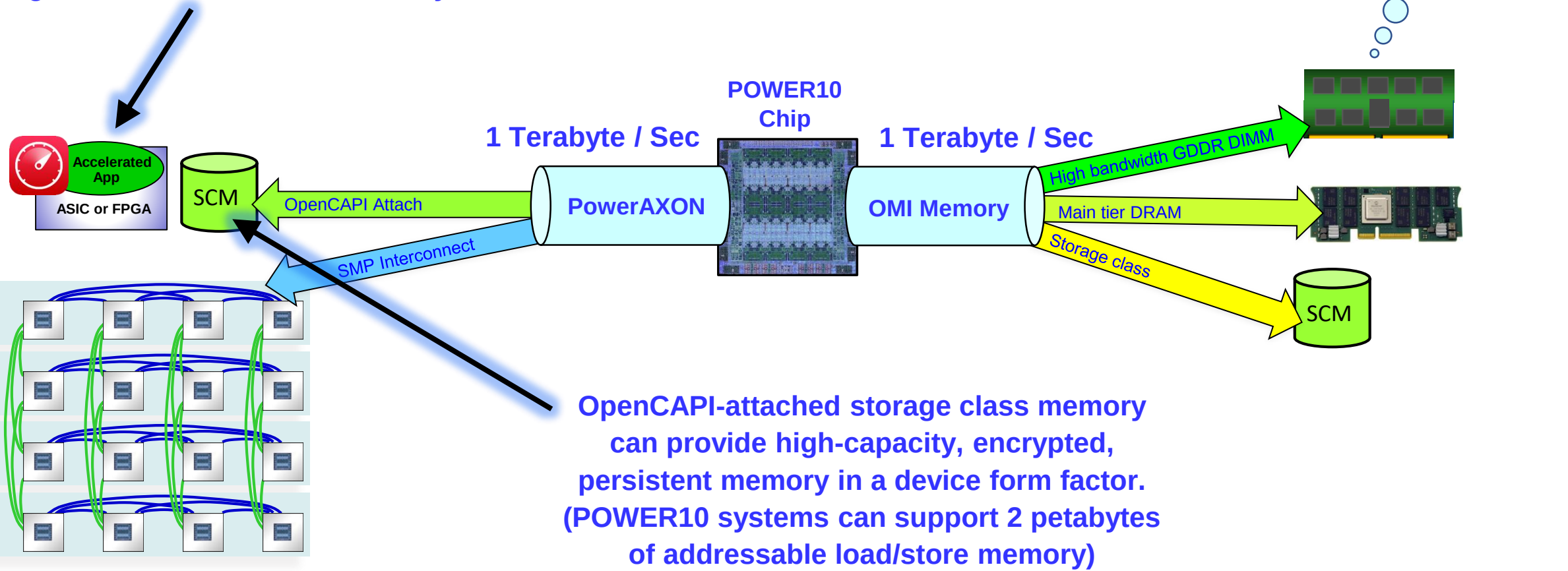
Built on best-of-breed
Low Power, Low Latency,
High Bandwidth
Signaling Technology

Data Plane Bandwidth and Capacity: Open Memory Interfaces



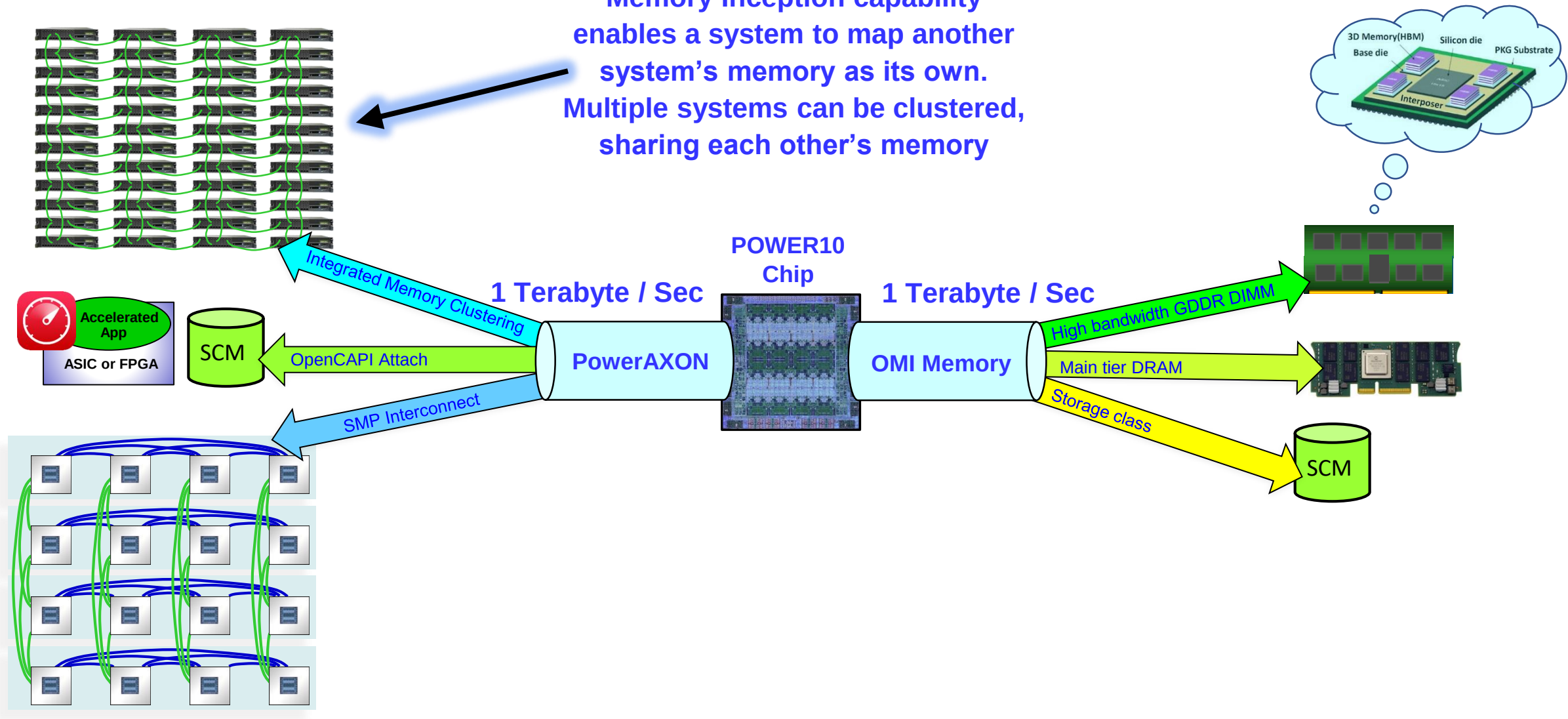
System Heterogeneity and Data Plane Capacity: OpenCAPI

OpenCAPI attaches FPGA
or ASIC-based Accelerators
to POWER10 host with
High Bandwidth and Low Latency



Pod Composability: PowerAXON Memory Clustering

Memory Inception capability enables a system to map another system's memory as its own. Multiple systems can be clustered, sharing each other's memory



Memory Clustering: Distributed Memory Disaggregation and Sharing

Use case: Share load/store memory amongst directly connected neighbors within Pod

Unlike other schemes, memory can be used:

- As low latency local memory
- As NUMA latency remote memory

Example: Pod = 8 systems each with 8TB

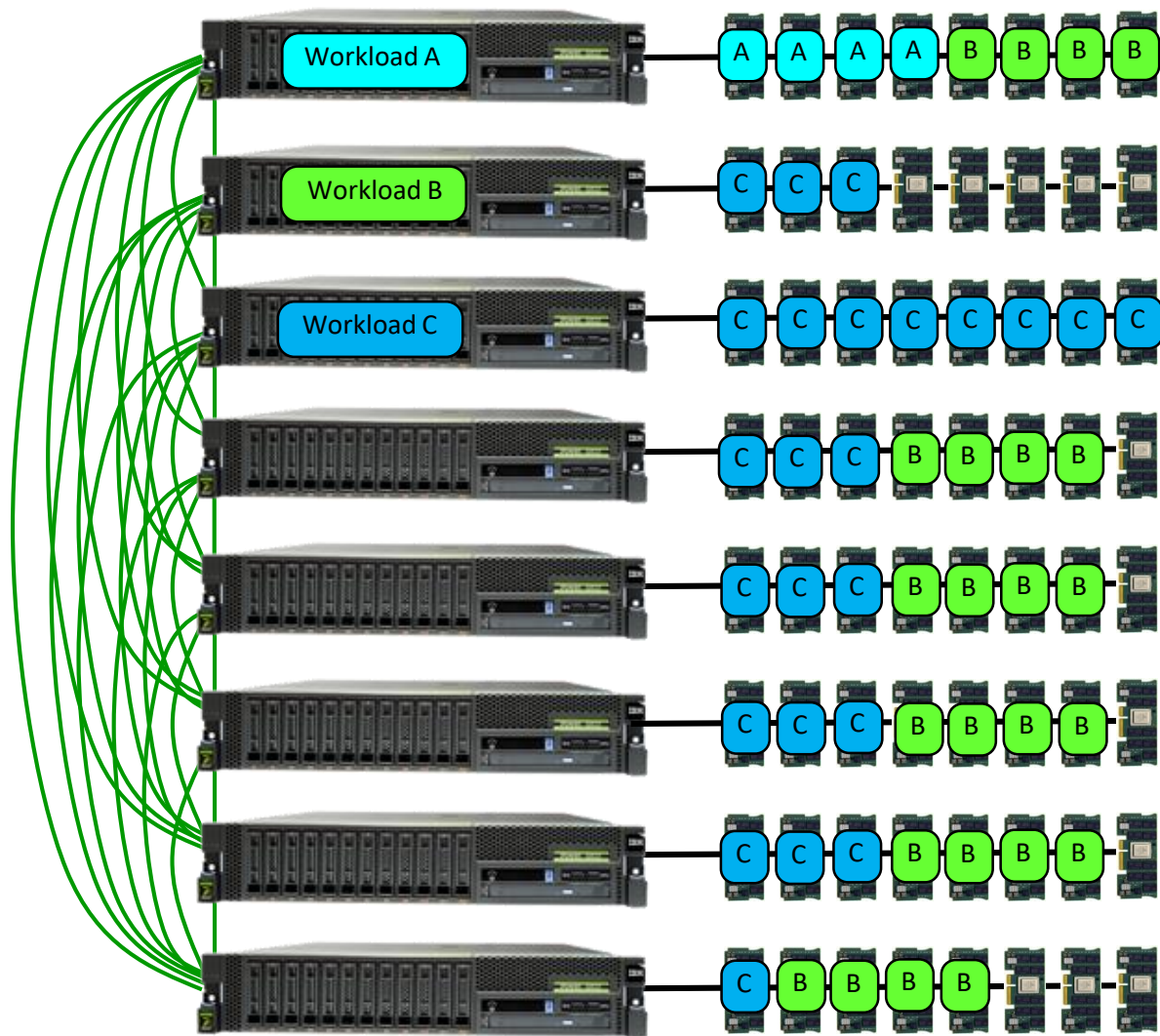
Workload A Rqmt: 4 TB low latency

Workload B Rqmt: 24 TB relaxed latency

Workload C Rqmt: 8 TB low latency plus
16TB relaxed latency

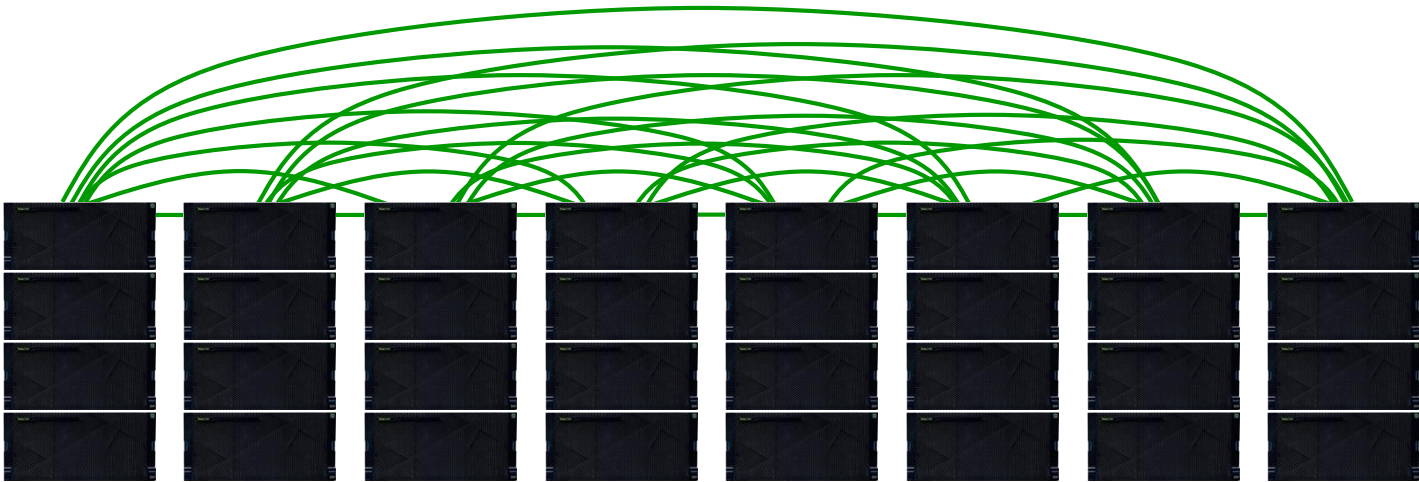
All Rqmts met by configuration shown

POWER10 2 Petabyte memory size enables
much larger configurations

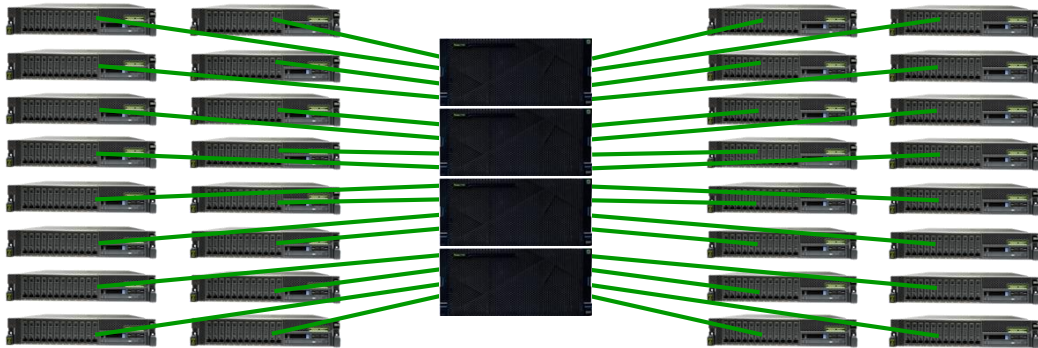


Memory Clustering: Enterprise-Scale Memory Sharing

Pod of Large Enterprise Systems
Distributed Sharing at Petabyte Scale



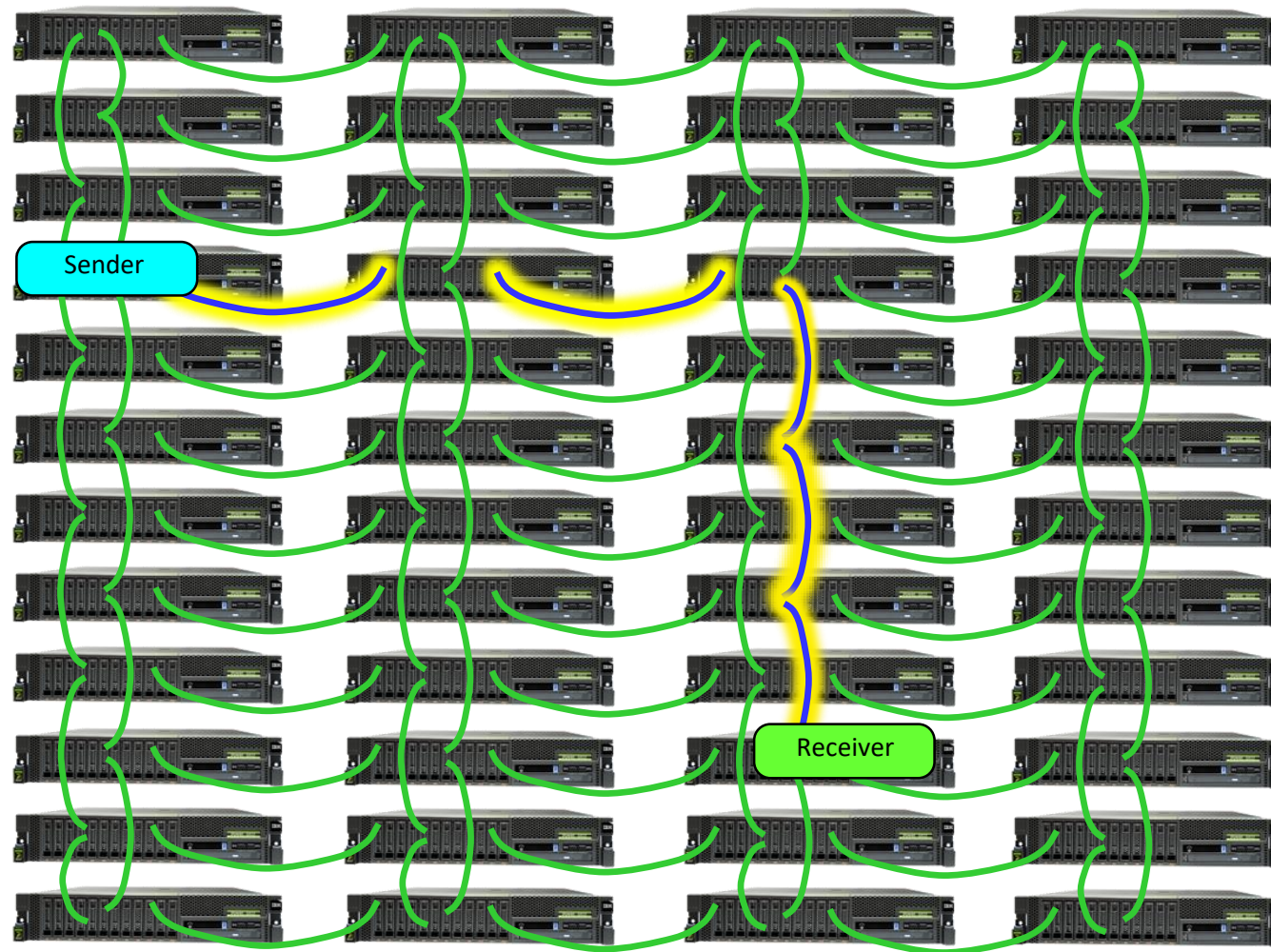
Or Hub-and-spoke with memory server
and memory-less compute nodes



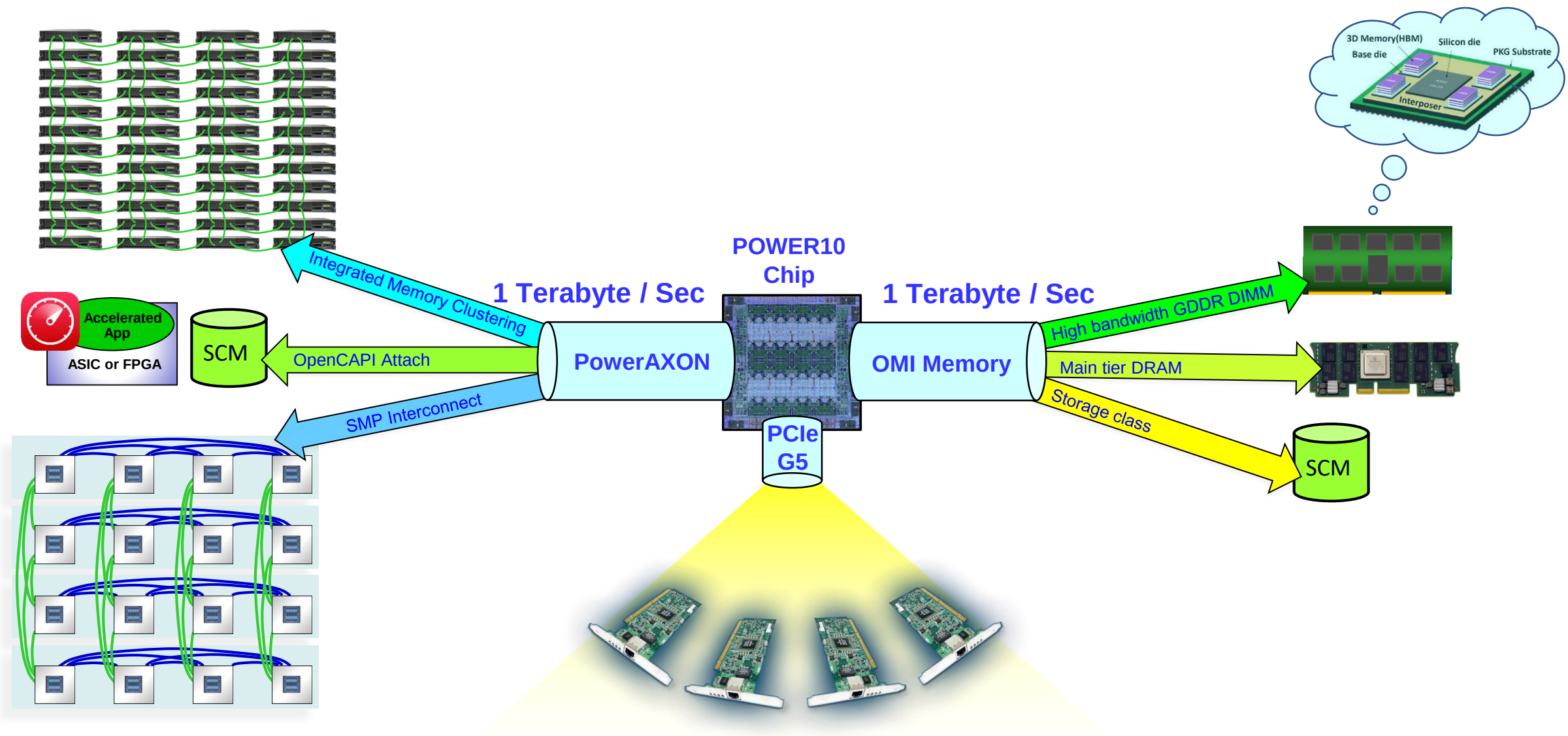
Memory Clustering: Pod-level Clustering

Use case: Low latency, high bandwidth
messaging scaling to 1000's of nodes

Leverage 2 Petabyte addressability to create
memory window into each destination for
messaging mailboxes

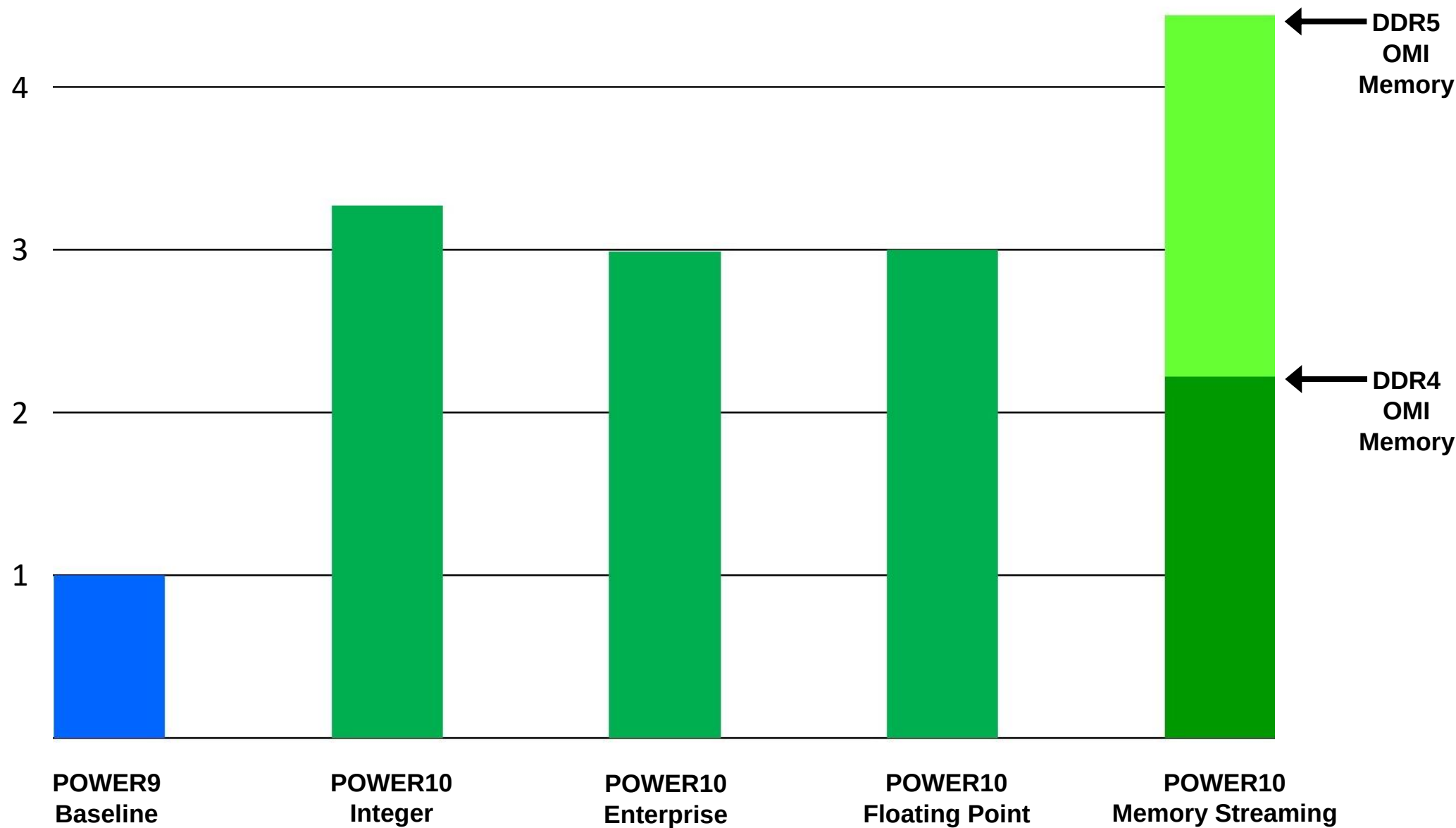


System Composability: PCIe Gen 5 Industry I/O Attach



(PowerAXON and OMI Memory configurations show processor capability only, and do not imply system product offerings)

POWER10 General Purpose Socket Performance Gains



(Performance assessments based upon pre-silicon engineering analysis of POWER10 dual-socket server offering vs POWER9 dual-socket server offering)

Powerful Core = Enterprise Strength + AI Infused

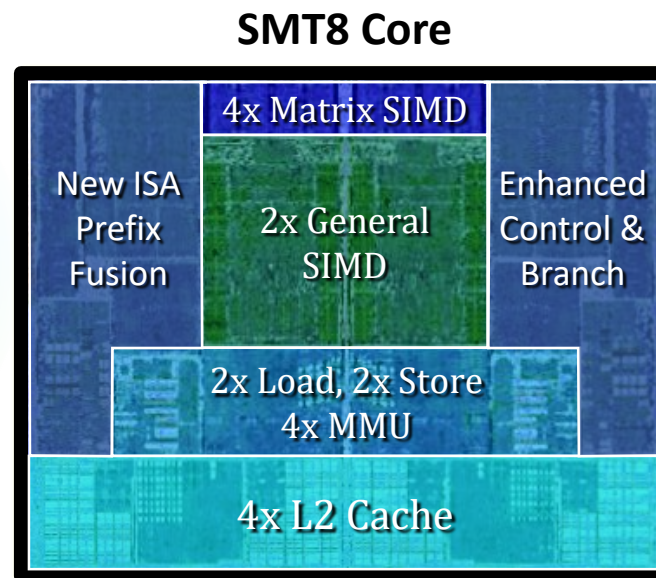
New Enterprise Micro-architecture

- Flexibility
 - Up to 8 threads per core / 240 per socket
- Optimized for performance and efficiency
 - +30% avg. core performance*
 - +20% avg. single thread performance*
 - 2.6x core performance efficiency* (3x @ socket)

AI Infused

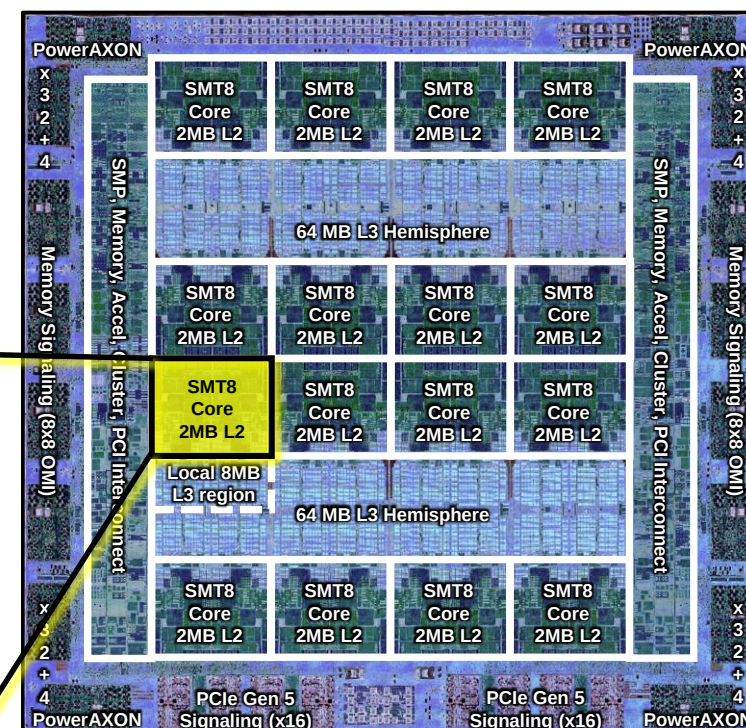
- 4x matrix SIMD acceleration*
- 2x bandwidth & general SIMD*
- 4x L2 cache capacity with improved thread isolation*
- New ISA with AI data-types

* versus POWER9



1-2 POWER10 chips per socket

- Up to **30 SMT8 Cores**
- Up to **60 SMT4 Cores**



Powerful Core : Enterprise Flexibility

Multiple World-class Software Stacks

Resilience and full stack integrity

- PowerVM, KVM
- AIX, IBMi, Linux on Power, OpenShift

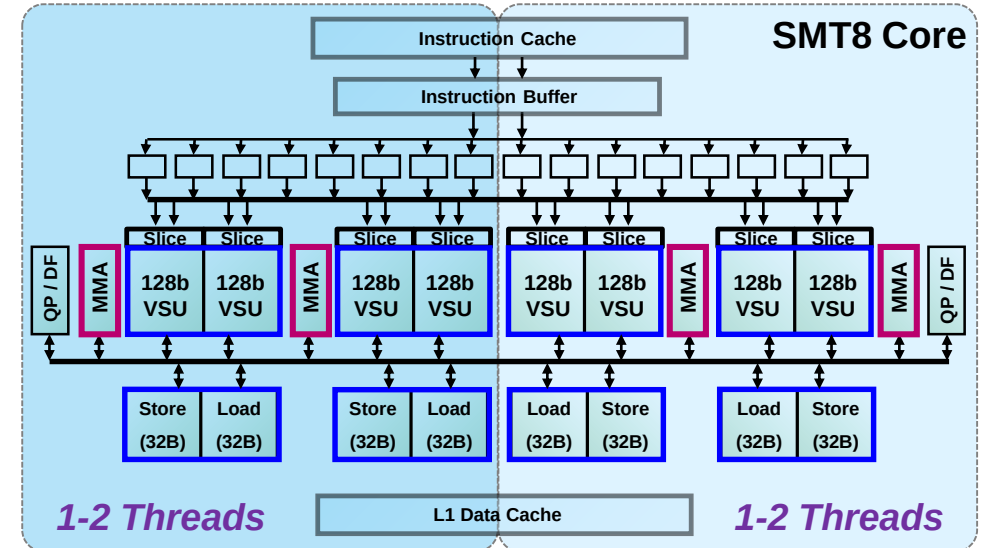


Partition flexibility and security

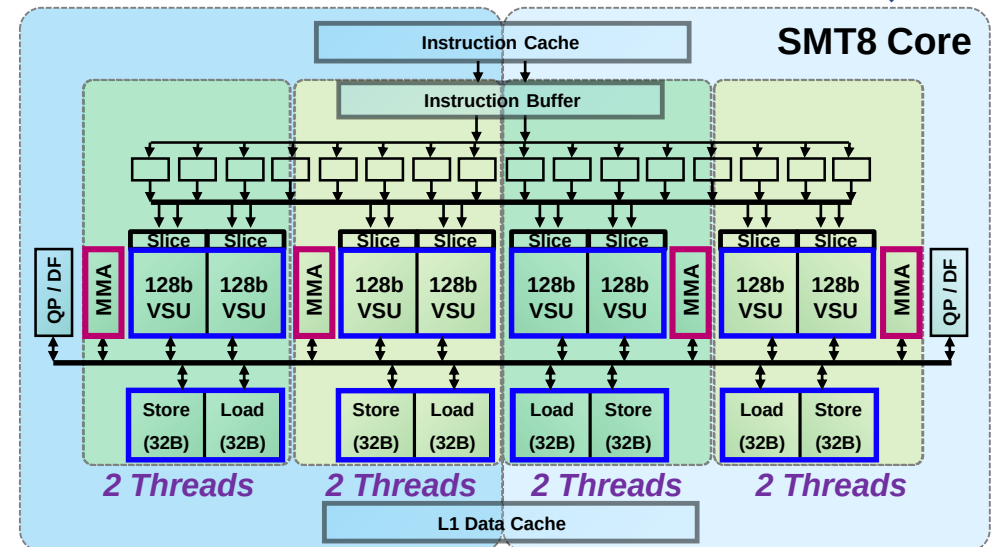
- Full-core level LPAR
- Thread-based LPAR scheduling
- **NEW:** With PowerVM Hypervisor
 - Nested KVM + PowerVM
 - Hardware assisted container/VM isolation



Hardware Based Workload Balance



Automatic Thread Resource Balancing



8 Threads Active

IBM POWER10

Powerful Architecture : AI Infused and Future Ready

POWER10 implements Power ISA v3.1

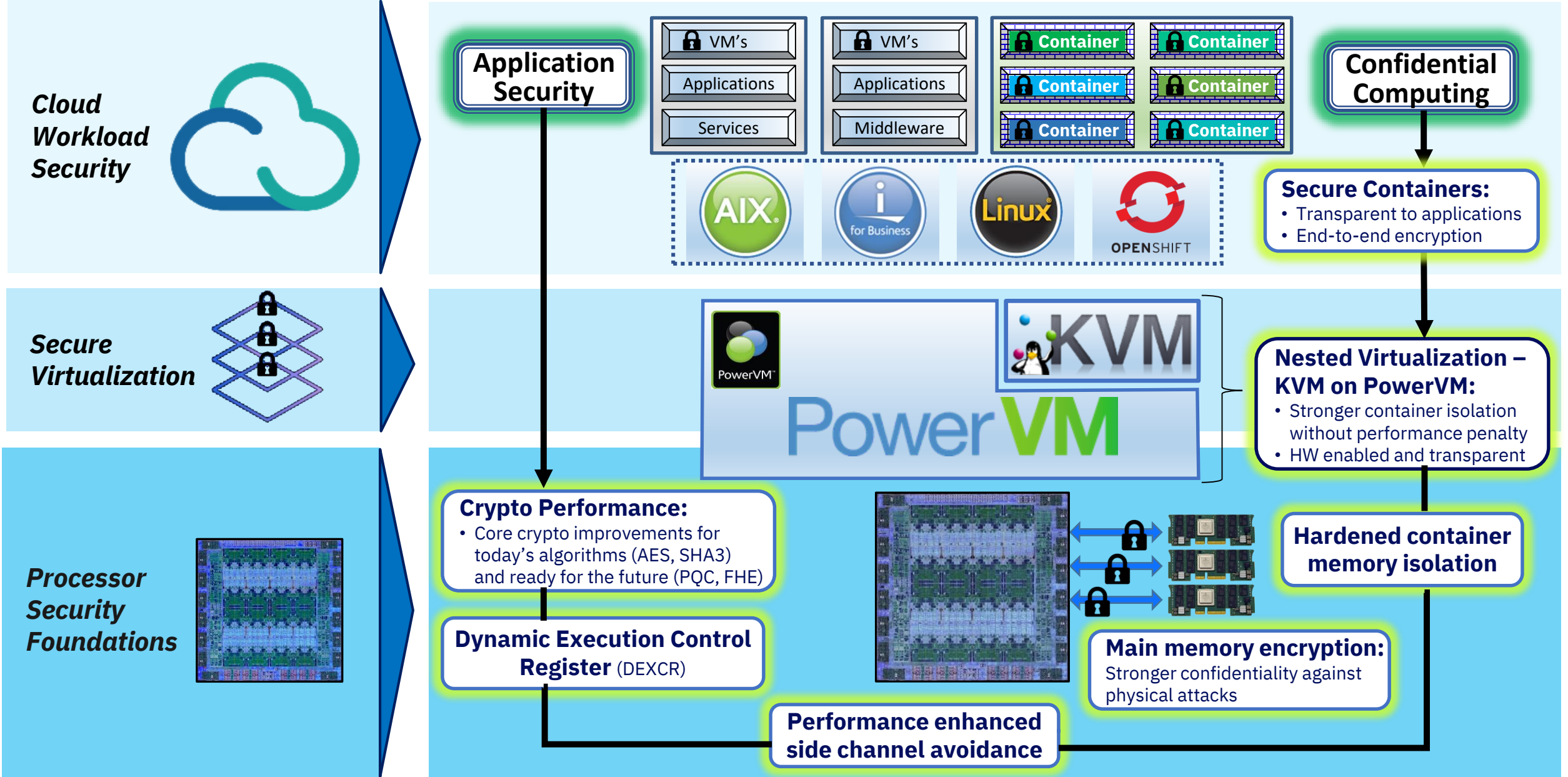
- v3.1 was the latest open Power ISA contributed to the OpenPOWER Foundation:
Royalty free and inclusive of patents for compliant designs



POWER10 Architecture – Feature Highlights

Prefix Architecture	Greatly expanded opcode space, pc-relative addressing, MMA masking, etc.	RISC friendly 8B instructions including modified and new opcode forms.
New Instructions and Datatypes	New Scalar instructions for control flow, and operation symmetry	Set Boolean extensions; quad-precision extensions; 128b integer extensions; test LSB by byte; byte reverse GPR; int mul/div modulo; string isolate/clear; pause, wait-reserve.
	New SIMD instructions for AI, throughput and data manipulation	32-byte load/store vector-pair; MMA (matrix math assist) with reduced precision; bfloat-16 converts; permute variations: extract, insert, splat, blend; compress/expand assist; mask generation; bit manipulation.
Advanced System Features and Ease of Use	Storage management	Persistent memory barrier / flush; store sync; translation extensions.
	Debug	PMU sampling, filtering; debug watchpoints; tracing.
	Hot/Cold page tracking	Recording for memory management.
	Copy/Paste extensions	Memory movement; continued on-chip acceleration: Gzip, 842 compression, AES/SHA.
Advanced EnergyScale	Adaptive power management	Additional performance boost across the operating range.
Security for Cloud	Transparent isolation and security for enterprise cloud workloads	Nested virtualization with KVM on PowerVM; secure containers; main memory encryption; dynamic execution control; secure PMU.

Security : End-to-End for the Enterprise Cloud

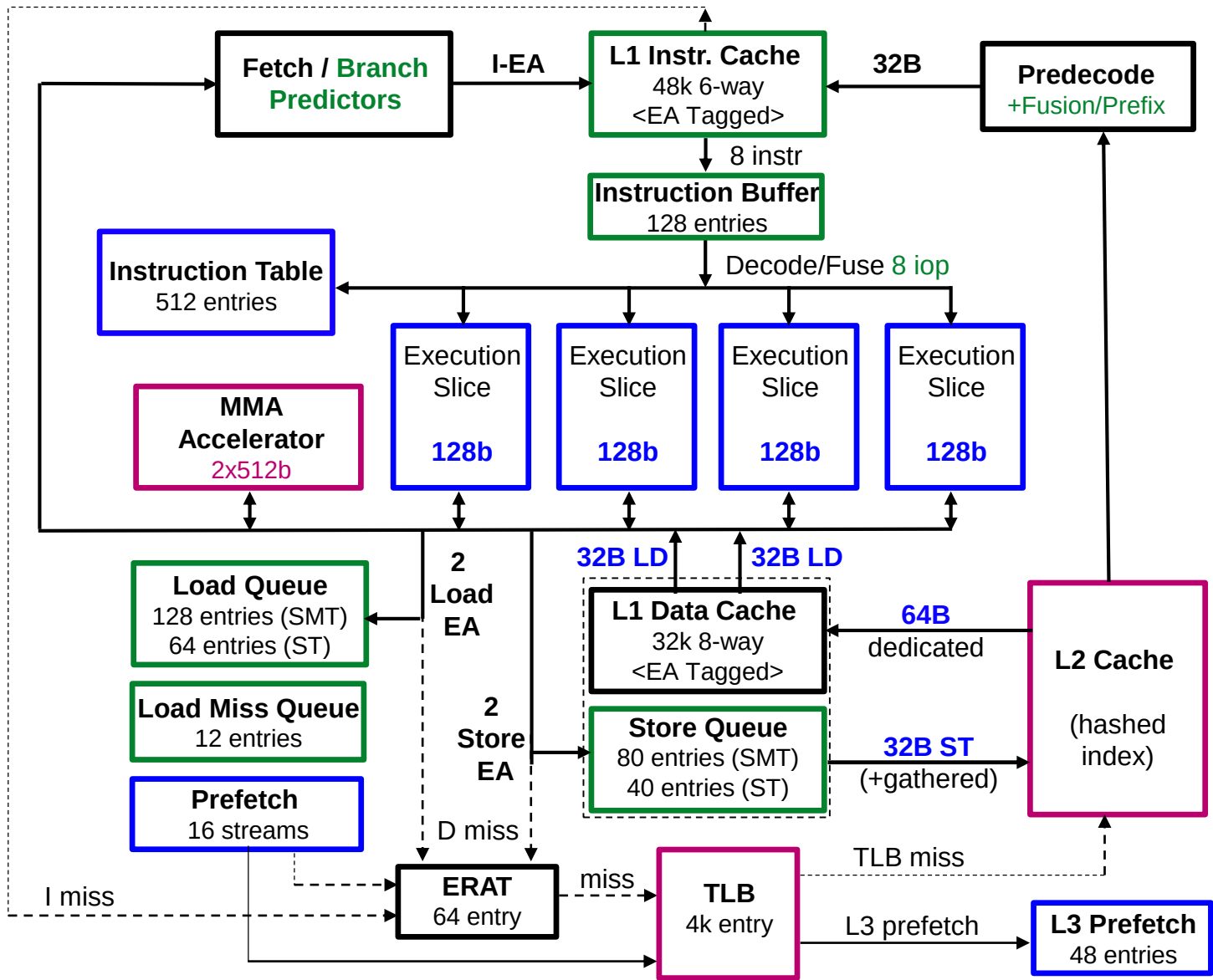


Powerful Core : Enterprise Strength

P
E
R
F
O
R
M
A
N
C
E

W
A
T
T

P10 Core Micro-architecture
½ SMT8 Core Resources Shown = SMT4 Core Equivalent



Capacity vs. POWER9: Improved >= 2x = 4x

IBM POWER10

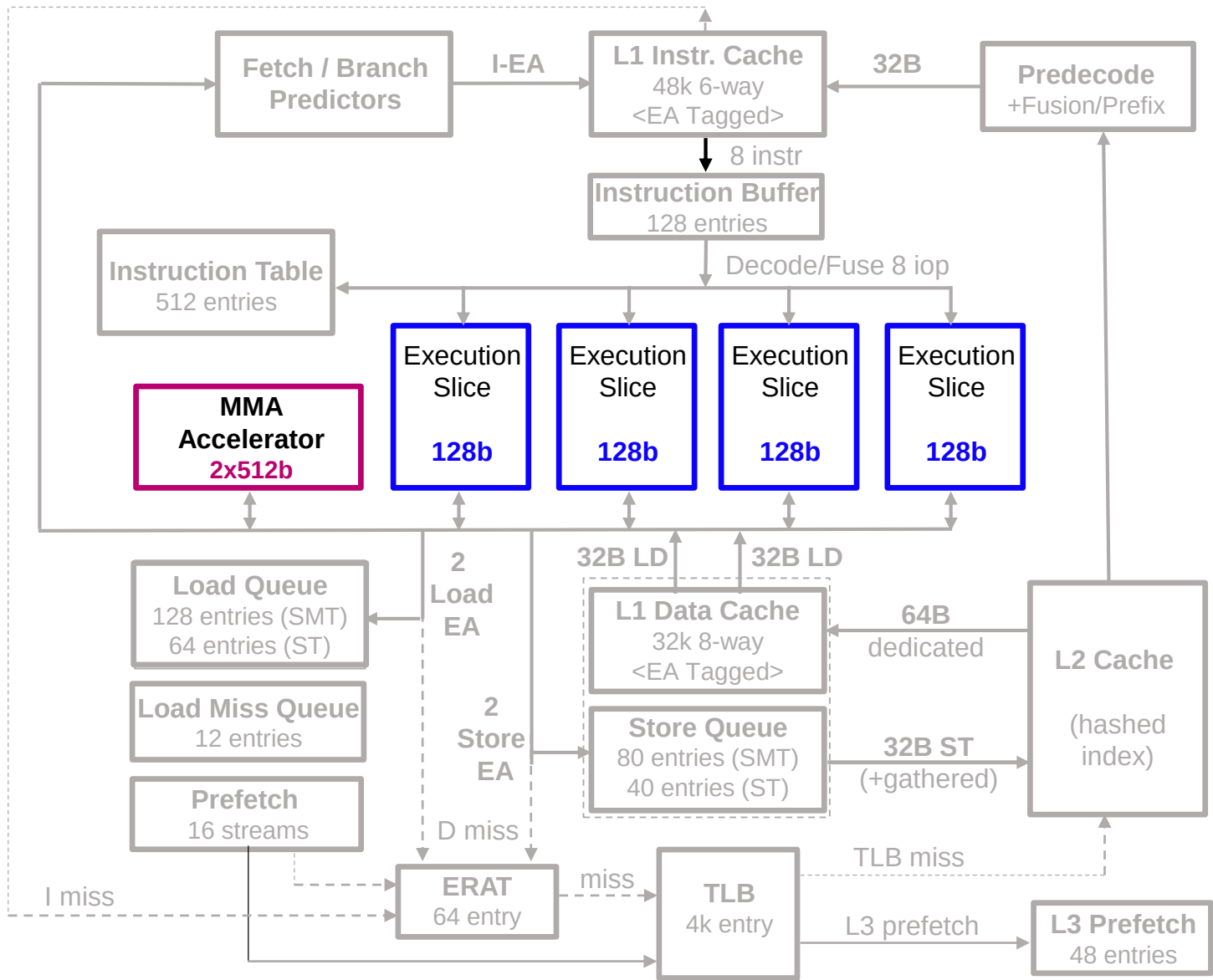
Powerful Core : Enterprise Strength

P
E
R
F
O
R
M
A
N
C
E

- Double SIMD + Inference acceleration
 - 2x SIMD, 4x MMA, 4x AES/SHA

W
A
T
T

P10 Core Micro-architecture
½ SMT8 Core Resources Shown = SMT4 Core Equivalent



Capacity vs. POWER9: Improved >= 2x = 4x

IBM POWER10

Powerful Core : Enterprise Strength

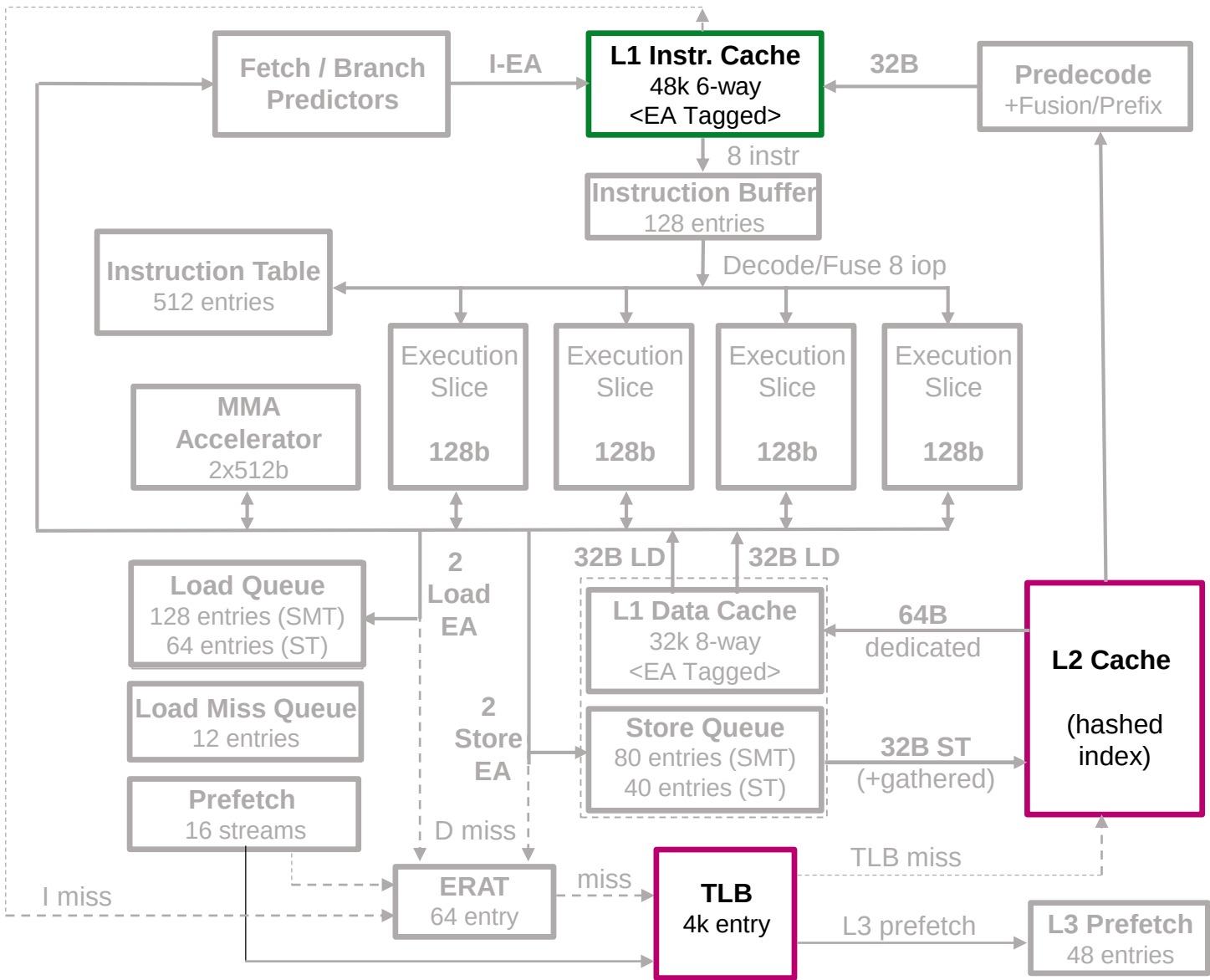
P
E
R
F
O
R
M
A
N
C
E

- **Double SIMD + Inference acceleration**
 - 2x SIMD, 4x MMA, 4x AES/SHA
- **Larger working-sets**
 - 1.5x L1-Instruction cache, 4x L2, 4x TLB

W
A
T
T

P10 Core Micro-architecture

½ SMT8 Core Resources Shown = SMT4 Core Equivalent



Capacity vs. POWER9:

Improved

>= 2x

= 4x

IBM POWER10

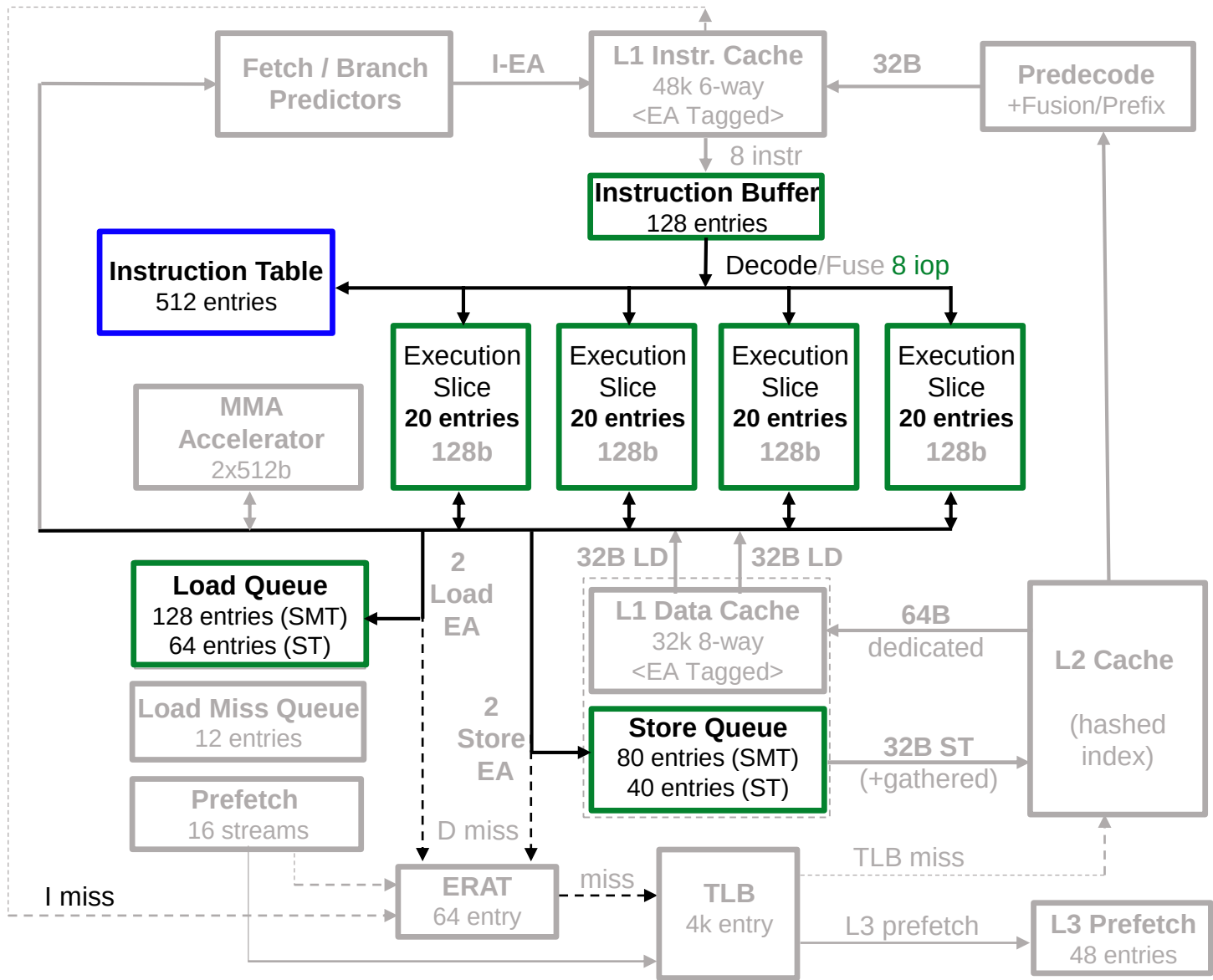
Powerful Core : Enterprise Strength

P
E
R
F
O
R
M
A
N
C
E

- Double SIMD + Inference acceleration
 - 2x SIMD, 4x MMA, 4x AES/SHA
- Larger working-sets
 - 1.5x L1-Instruction cache, 4x L2, 4x TLB
- Deeper/wider instruction windows

W
A
T
T

P10 Core Micro-architecture
½ SMT8 Core Resources Shown = SMT4 Core Equivalent



Capacity vs. POWER9: Improved >= 2x = 4x

IBM POWER10

Powerful Core : Enterprise Strength

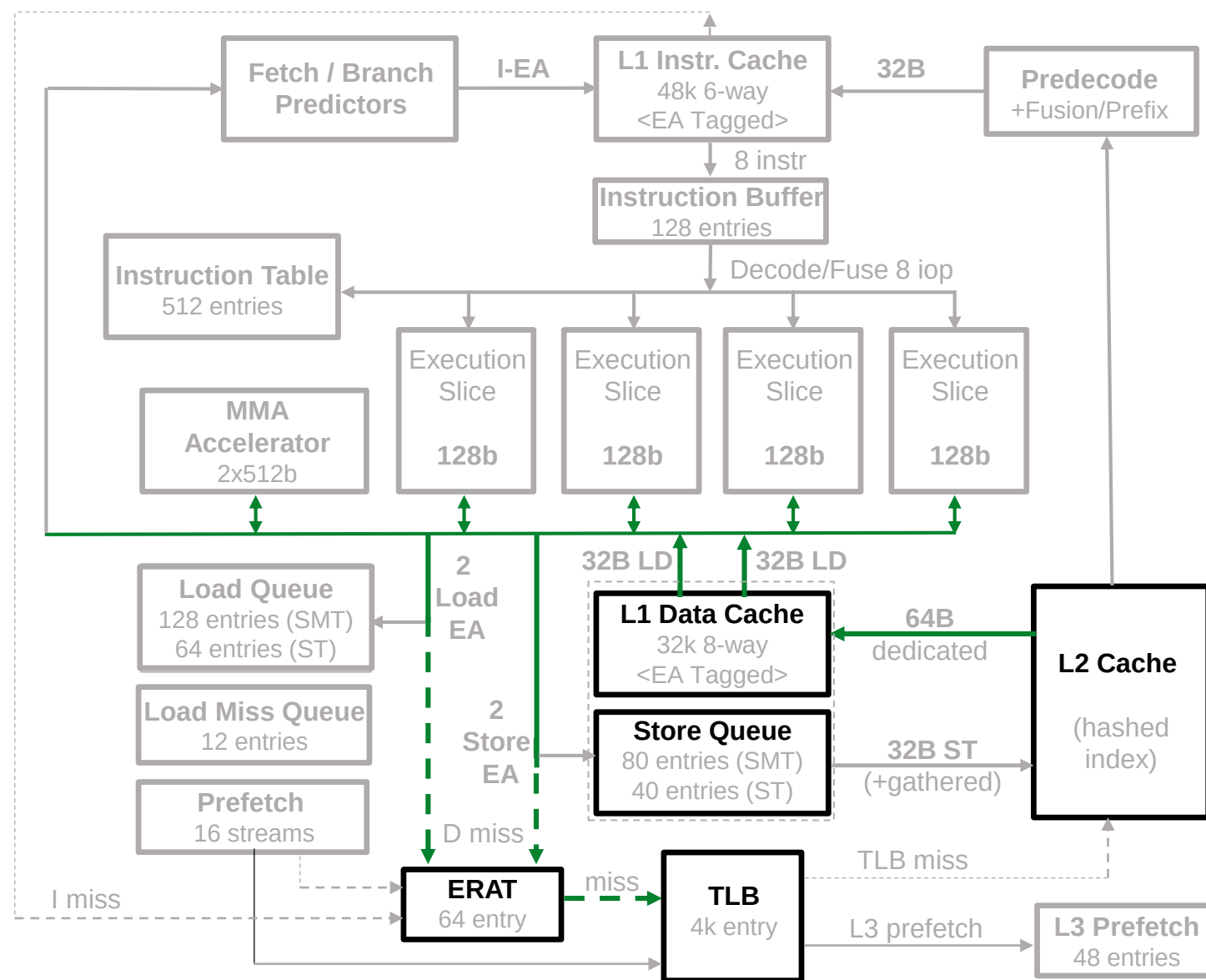
PERFORMANCE

- **Double SIMD + Inference acceleration**
 - 2x SIMD, 4x MMA, 4x AES/SHA
- **Larger working-sets**
 - 1.5x L1-Instruction cache, 4x L2, 4x TLB
- **Deeper/wider instruction windows**
- **Data latency (cycles)**
 - L2 13.5 (minus 2), L3 27.5 (minus 8)
 - L1-D cache 4 +0 for Store forward (minus 2)
 - TLB access +8.5 (minus 7)

WATT

P10 Core Micro-architecture

½ SMT8 Core Resources Shown = SMT4 Core Equivalent



Capacity vs. POWER9:

Improved

>= 2x

= 4x

IBM POWER10

Powerful Core : Enterprise Strength

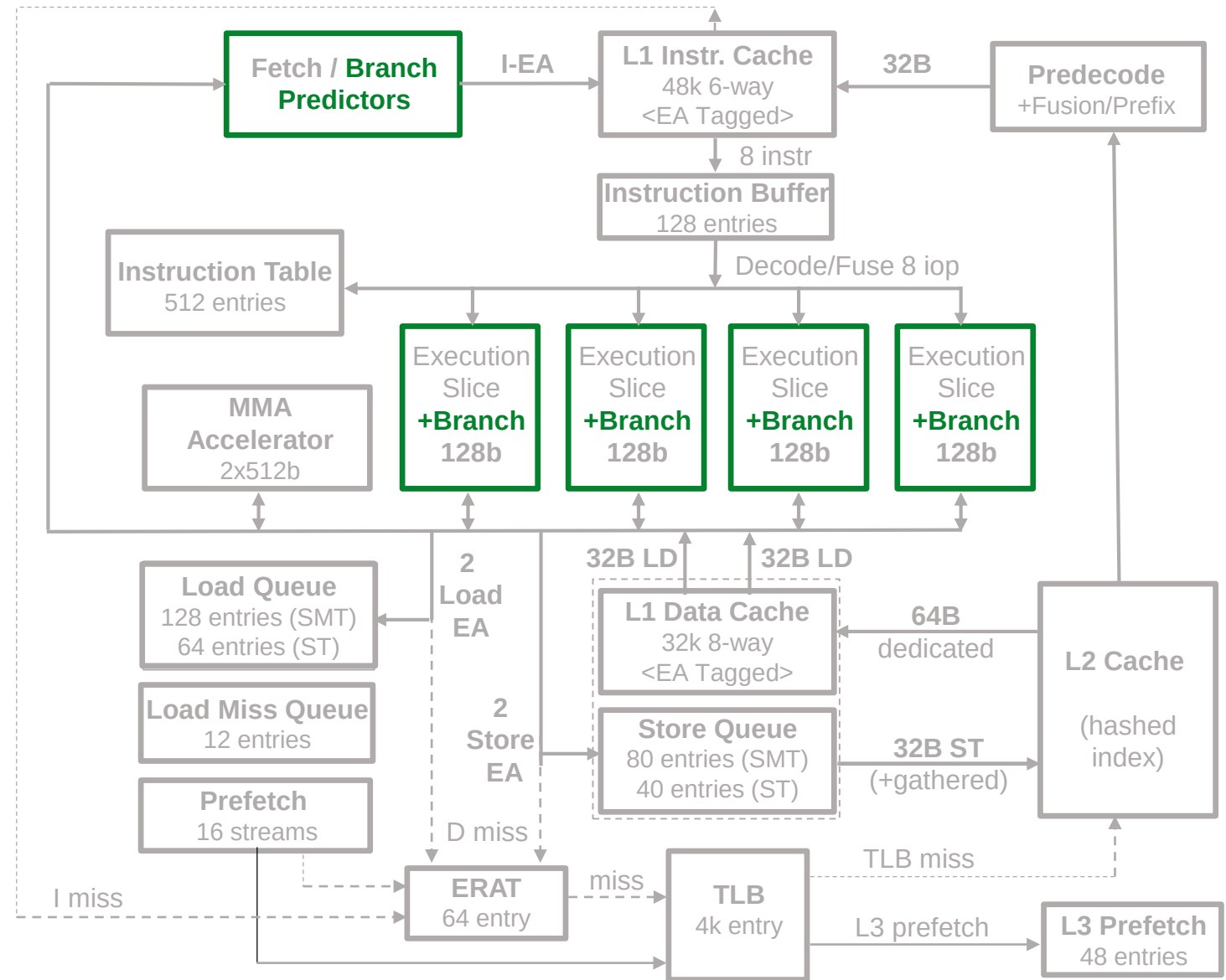
PERFORMANCE

- **Double SIMD + Inference acceleration**
 - 2x SIMD, 4x MMA, 4x AES/SHA
- **Larger working-sets**
 - 1.5x L1-Instruction cache, 4x L2, 4x TLB
- **Deeper/wider instruction windows**
- **Data latency (cycles)**
 - L2 13.5 (minus 2), L3 27.5 (minus 8)
 - L1-D cache 4 +0 for Store forward (minus 2)
 - TLB access +8.5 (minus 7)
- **Branch**
 - Target registers with GPR in main regfile
 - New predictors: target and direction, 2x BHT

WATT

P10 Core Micro-architecture

½ SMT8 Core Resources Shown = SMT4 Core Equivalent



Capacity vs. POWER9:

Improved

>= 2x

= 4x

IBM POWER10

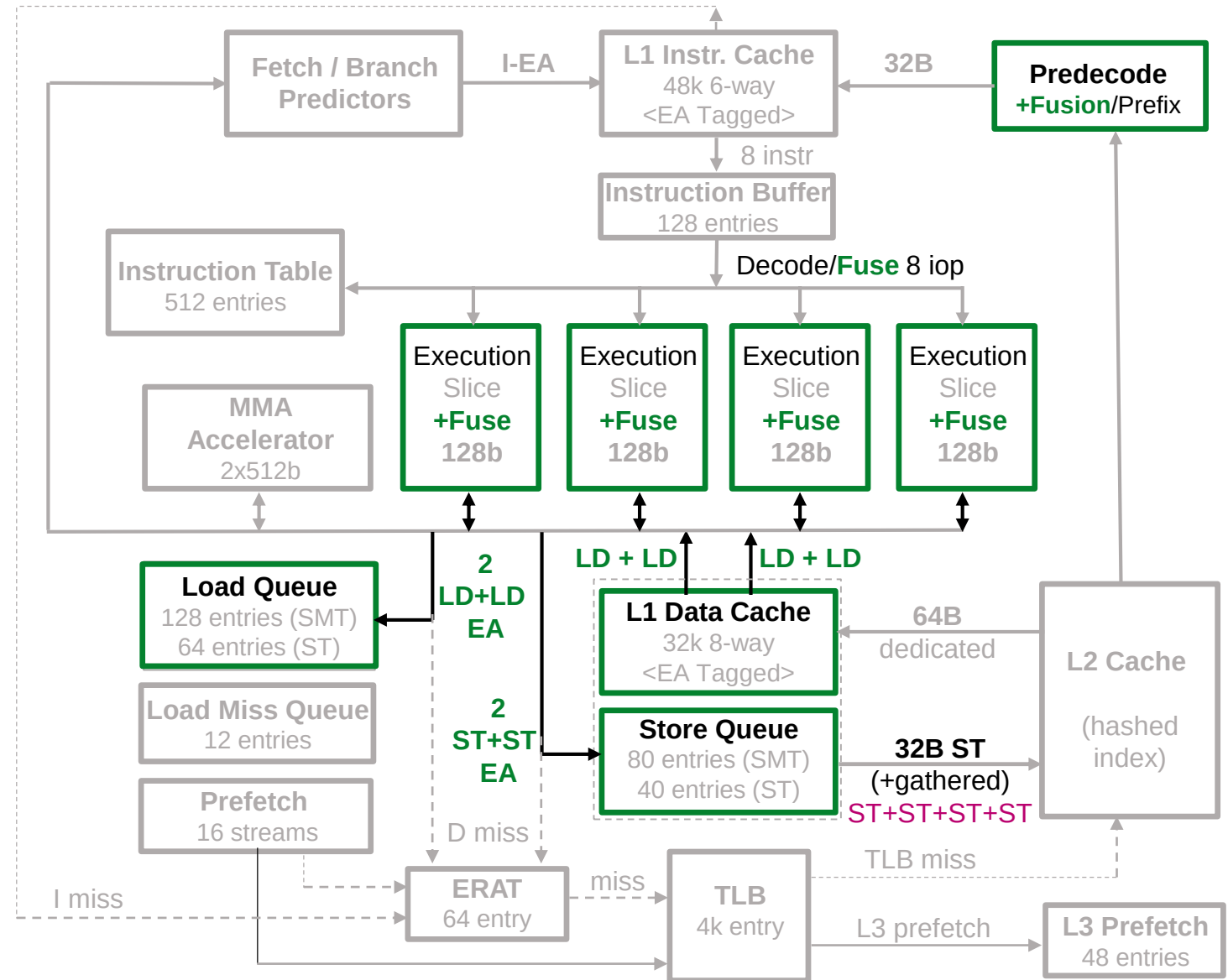
Powerful Core : Enterprise Strength

PERFORMANCE

- **Double SIMD + Inference acceleration**
 - 2x SIMD, 4x MMA, 4x AES/SHA
- **Larger working-sets**
 - 1.5x L1-Instruction cache, 4x L2, 4x TLB
- **Deeper/wider instruction windows**
- **Data latency (cycles)**
 - L2 13.5 (minus 2), L3 27.5 (minus 8)
 - L1-D cache 4 +0 for Store forward (minus 2)
 - TLB access +8.5 (minus 7)
- **Branch**
 - Target registers with GPR in main regfile
 - New predictors: target and direction, 2x BHT
- **Fusion**
 - Fixed, SIMD, other: merge and back to back
 - Load, store : consecutive storage

WATT

P10 Core Micro-architecture
 $\frac{1}{2}$ SMT8 Core Resources Shown = SMT4 Core Equivalent



Capacity vs. POWER9:

Improved

>= 2x

= 4x

IBM POWER10

PERFORMANCE

- W
-
- A
-
- T
-
- T

- P10 Core Micro-architecture**
 $\frac{1}{2}$ SMT8 Core Resources Shown = SMT4 Core Equivalent

The diagram illustrates the P10 Core Micro-architecture, showing the flow of instructions and data through various components:

 - Fetch / Branch Predictors** send **I-EA** to the **L1 Instr. Cache** (48k 6-way <EA Tagged>).
 - The **L1 Instr. Cache** sends **8 instr** to the **Instruction Buffer** (128 entries).
 - The **Instruction Buffer** sends **Decode/Fuse 8 iop** to the **Execution Slices** (4 slices, each 128b).
 - The **Execution Slices** send **32B LD** to the **L1 Data Cache** (32k 8-way <EA Tagged>).
 - The **L1 Data Cache** sends **64B dedicated** to the **L2 Cache** (hashed index).
 - The **L1 Data Cache** sends **32B ST (+gathered)** to the **L2 Cache**.
 - The **L2 Cache** sends **L3 prefetch** to the **L3 Prefetch** (48 entries).
 - The **L3 Prefetch** sends **L3 prefetch** to the **TLB** (4k entry).
 - The **TLB** sends **miss** to the **ERAT** (64 entry).
 - The **ERAT** sends **D miss** to the **Load Queue** (128 entries (SMT), 64 entries (ST)).
 - The **Load Queue** sends **2 Load EA** to the **Execution Slices**.
 - The **Load Queue** sends **2 Store EA** to the **Store Queue** (80 entries (SMT), 40 entries (ST)).
 - The **Store Queue** sends **32B LD** to the **L1 Data Cache**.
 - The **Store Queue** sends **miss** to the **TLB**.
 - The **TLB** sends **TLB miss** to the **L2 Cache**.
 - The **TLB** sends **miss** to the **ERAT**.
 - The **ERAT** sends **I miss** to the **Fetch / Branch Predictors**.
 - The **Fetch / Branch Predictors** send **I-EA** to the **L1 Instr. Cache**.
 - The **L1 Instr. Cache** sends **32B** to the **Predecode +Fusion/Prefix** block.
 - The **Predecode +Fusion/Prefix** block sends **8 instr** to the **Instruction Buffer**.
 - The **Instruction Buffer** sends **Decode/Fuse 8 iop** to the **Execution Slices**.
 - The **Execution Slices** send **32B LD** to the **L1 Data Cache**.
 - The **L1 Data Cache** sends **64B dedicated** to the **L2 Cache**.
 - The **L1 Data Cache** sends **32B ST (+gathered)** to the **L2 Cache**.
 - The **L2 Cache** sends **L3 prefetch** to the **L3 Prefetch**.
 - The **L3 Prefetch** sends **L3 prefetch** to the **TLB**.
 - The **TLB** sends **miss** to the **ERAT**.
 - The **ERAT** sends **D miss** to the **Load Queue**.
 - The **Load Queue** sends **2 Load EA** to the **Execution Slices**.
 - The **Load Queue** sends **2 Store EA** to the **Store Queue**.
 - The **Store Queue** sends **32B LD** to the **L1 Data Cache**.
 - The **Store Queue** sends **miss** to the **TLB**.
 - The **TLB** sends **TLB miss** to the **L2 Cache**.
 - The **TLB** sends **miss** to the **ERAT**.
 - The **ERAT** sends **I miss** to the **Fetch / Branch Predictors**.

PERFORMANCE

- W
A
T
T**

- P10 Core Micro-architecture**
 $\frac{1}{2}$ SMT8 Core Resources Shown = SMT4 Core Equivalent

The diagram illustrates the P10 core micro-architecture, showing the flow of instructions and data between various components. Key components and their interactions include:

 - Fetch / Branch Predictors**: Connected to the **L1 Instr. Cache** via an **I-EA** (Instruction-Effective Address) path.
 - L1 Instr. Cache**: 48k 6-way <EA Tagged>. It feeds into the **Instruction Buffer** (128 entries) via an **8 instr** path.
 - Predecode +Fusion/Prefix**: Connected to the **L1 Instr. Cache** via a **32B** path.
 - Instruction Buffer**: Feeds into the **Execution Slices** via a **Decode/Fuse 8 iop** path.
 - Execution Slices**: Four slices, each containing **+STF 128b** (Store Forwarding Buffer). They are connected to the **L1 Data Cache** via **32B LD** (Load) and **32B ST** (Store) paths.
 - L1 Data Cache**: 32k 8-way <EA Tagged>. It is connected to the **L2 Cache** via a **64B dedicated** path.
 - L2 Cache**: (hashed index). It feeds back into the **Predecode +Fusion/Prefix** block via a **32B ST (+gathered)** path.
 - Store Queue**: 80 entries (SMT) / 40 entries (ST). It is connected to the **L1 Data Cache** via a **32B ST (+gathered)** path.
 - Load Queue**: 128 entries (SMT) / 64 entries (ST). It is connected to the **Execution Slices** via a **2 Load EA** path.
 - Load Miss Queue**: 12 entries. It is connected to the **Execution Slices** via a **2 Store EA** path.
 - Prefetch**: 16 streams. It is connected to the **Execution Slices** via a **D miss** path.
 - ERAT** (Execution Return Address Table): 64 entry. It is connected to the **Execution Slices** via a **D miss** path.
 - TLB** (Translation Lookaside Buffer): 4k entry. It is connected to the **Execution Slices** via a **D miss** path.
 - L3 Prefetch**: 48 entries. It is connected to the **TLB** via an **L3 prefetch** path.

Improvements over POWER9 are highlighted in green:

 - +Fusion/Prefix** in the Predecode block.
 - +STF** (Store Forwarding) in the Execution Slices.
 - +gathered** in the Store Queue.

At the bottom, a comparison states: **Watt vs. POWER9: Improved**.

IBM POWER10

PERFORMANCE

- W
A
T
T

- P10 Core Micro-architecture**
 ½ SMT8 Core Resources Shown = SMT4 Core Equivalent

The diagram illustrates the P10 Core Micro-architecture. Key components and their interactions include:

 - Fetch / Branch Predictors**: Send **I-EA** to the **L1 Instr. Cache**.
 - L1 Instr. Cache**: 48k 6-way <EA Tagged>. Receives 32B from the **Predecode +Fusion/Prefix** block and sends 8 instr to the **Instruction Buffer**.
 - Instruction Buffer**: 128 entries. Sends **Decode/Fuse 8 iop** to the **Execution Slices**.
 - Execution Slices**: Four slices, each 128b. They interact with the **L1 Data Cache** via **32B LD** and **32B ST**.
 - L1 Data Cache**: 32k 8-way <EA Tagged>. Receives 64B dedicated from the **L2 Cache** and sends 32B ST (+gathered) to the **L2 Cache**.
 - L2 Cache**: (hashed index). Receives TLB miss from the **TLB** and sends L3 prefetch to the **L3 Prefetch** block.
 - L3 Prefetch**: 48 entries.
 - Load Queue**: 128 entries (SMT), 64 entries (ST). Receives 2 Load EA from the **L1 Data Cache**.
 - Load Miss Queue**: 12 entries.
 - Prefetch**: 16 streams.
 - ERAT**: 64 entry. Receives I miss from the **Fetch / Branch Predictors** and D miss from the **L1 Data Cache**. Sends miss to the **TLB**.
 - TLB**: 4k entry. Receives miss from the **ERAT** and sends TLB miss to the **L2 Cache**.
 - Store Queue**: 80 entries (SMT), 40 entries (ST). Receives 2 Store EA from the **L1 Data Cache**.
 - MMA Accelerator**: 2x512b. Interacts with the **Execution Slices**.
 - Instruction Table**: 512 entries. Receives I-EA from the **Fetch / Branch Predictors**.
 - Predecode +Fusion/Prefix**: Sends 32B to the **L1 Instr. Cache** and receives L3 prefetch from the **L3 Prefetch** block.

PERFORMANCE

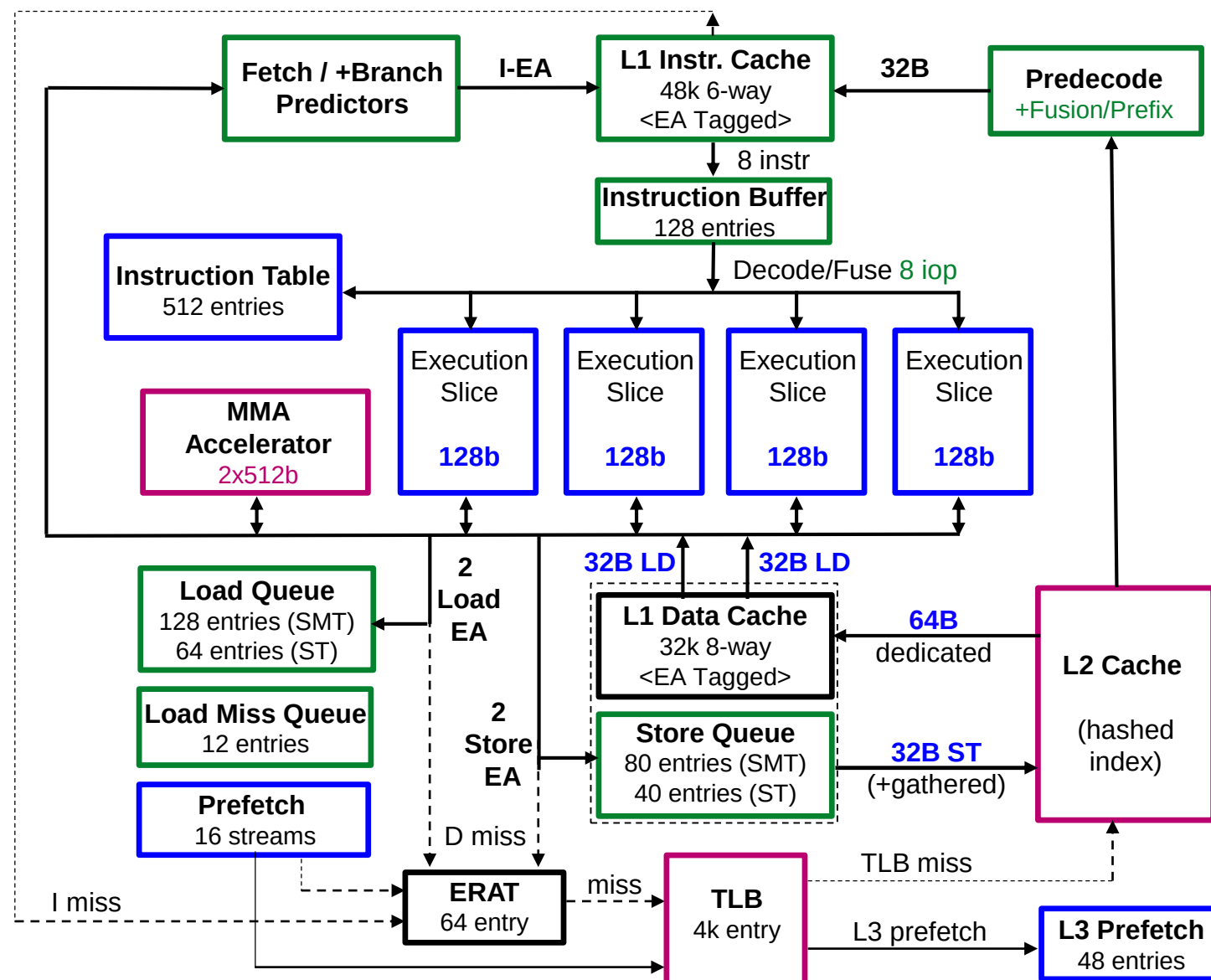
- W
A
T
T**

- ### Capacity vs. POWER9:

$$\geq 2x$$

IBM POWER10

¹/₂ SMT8 Core Resources Shown = SMT4 Core Equivalent



Powerful Core : Strength & Efficiency

P
E
R
F
O
R
M
A
N
C
E

- **Double SIMD + Inference acceleration**
 - 2x SIMD, 4x MMA, 4x AES/SHA
- **Larger working-sets**
 - 1.5x L1-Instruction cache, 4x L2, 4x TLB
- **Deeper/wider instruction windows**
- **Data latency (cycles)**
 - L2 13.5 (minus 2), L3 27.5 (minus 8)
 - L1-D cache 4 +0 for Store forward (minus 2)
 - TLB access +8.5 (minus 7)
- **Branch**
 - Target registers with GPR in main regfile
 - New predictors: target and direction, 2x BHT
- **Fusion**
 - Fixed, SIMD, other: merge and back to back
 - Load, store : consecutive storage

W
A
T
T

- **Improved clock-gating**
- **Design & micro-arch efficiency**
- **Branch accuracy: less wasted work**
- **Fusion / gather: less units of work**
- **Reduced ports / access**
 - Sliced target reg-file
 - Reduced read ports / entry
- **EA-tagged L1-D Cache & L1-I Cache**
 - CAM with cache-way/index
 - ERAT only on cache miss

1.3x

0.5x

= 2.6X performance / watt

POWER10 vs. POWER9 Core

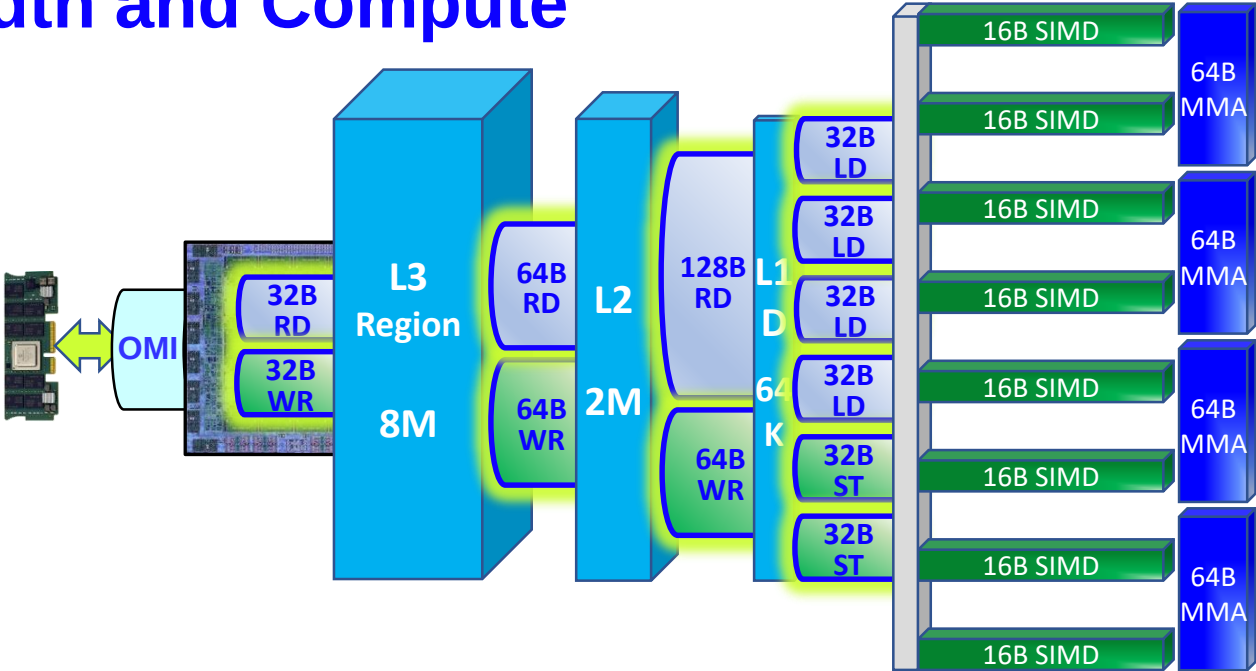
Powerful Core : AI Infused Bandwidth and Compute

2x Bytes from all sources
(OMI, L3, L2, L1 caches*)



BYTES

FLOP



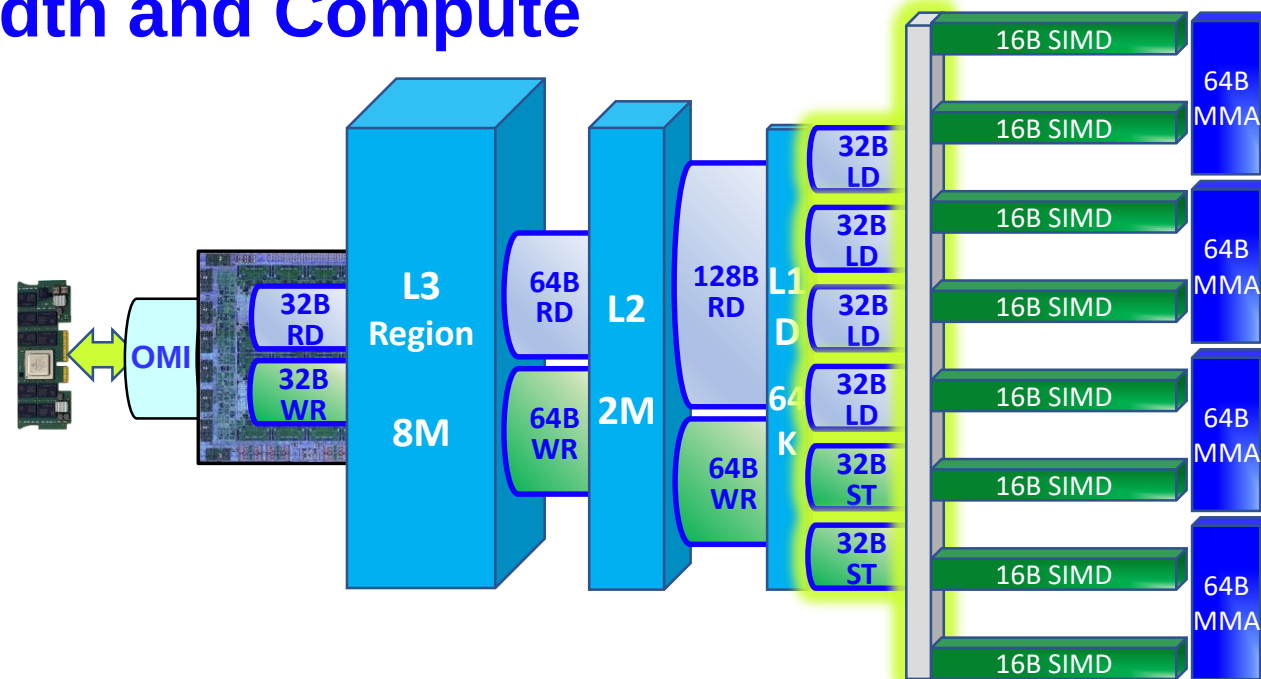
* versus POWER9

Powerful Core : AI Infused Bandwidth and Compute

2x Bytes from all sources

(OMI, L3, L2, L1 caches*)

- 4 32B loads, 2 32B stores per SMT8 Core
 - New ISA or fusion
 - Thread max 2 32B loads, 1 32B store



* versus POWER9

IBM POWER10

B
Y
T
E
S

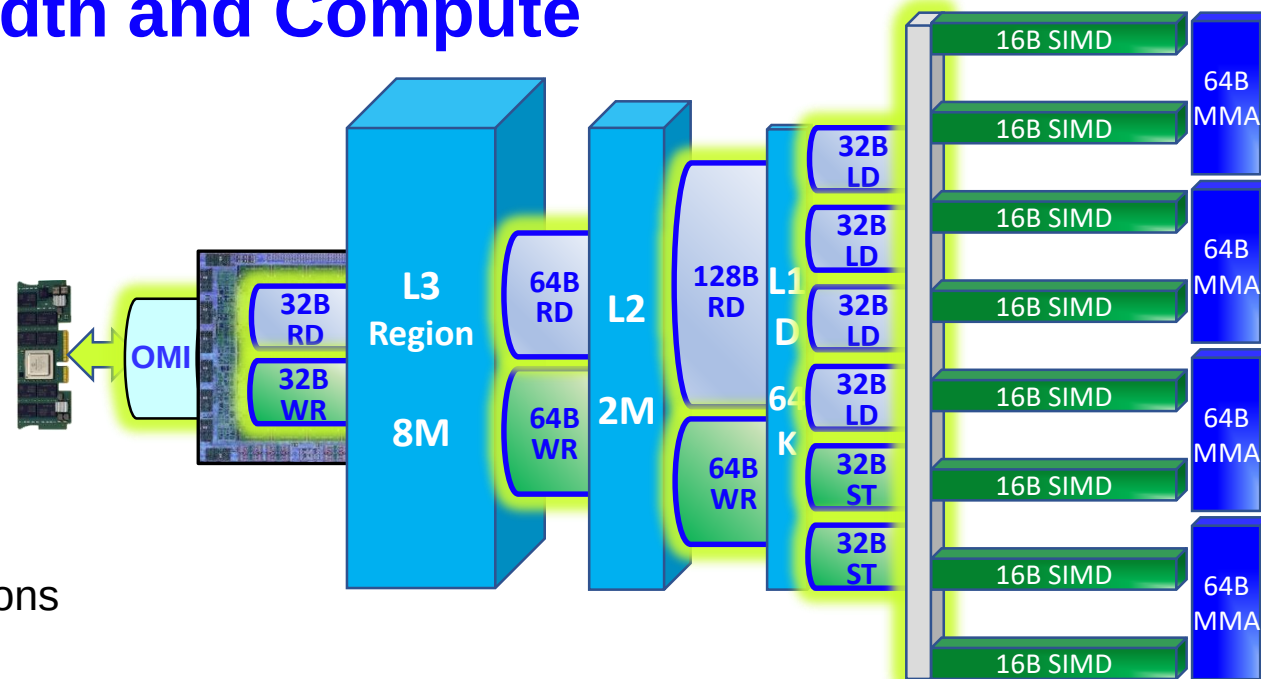
F
L
O
P

Powerful Core : AI Infused Bandwidth and Compute

2x Bytes from all sources

(OMI, L3, L2, L1 caches*)

- 4 32B loads, 2 32B stores per SMT8 Core
 - New ISA or fusion
 - Thread max 2 32B loads, 1 32B store
- OMI Memory to one Core
 - 256 GB/s peak, 120 GB/s sustained
 - With 3x L3 prefetch and memory prefetch extensions



* versus POWER9

IBM POWER10

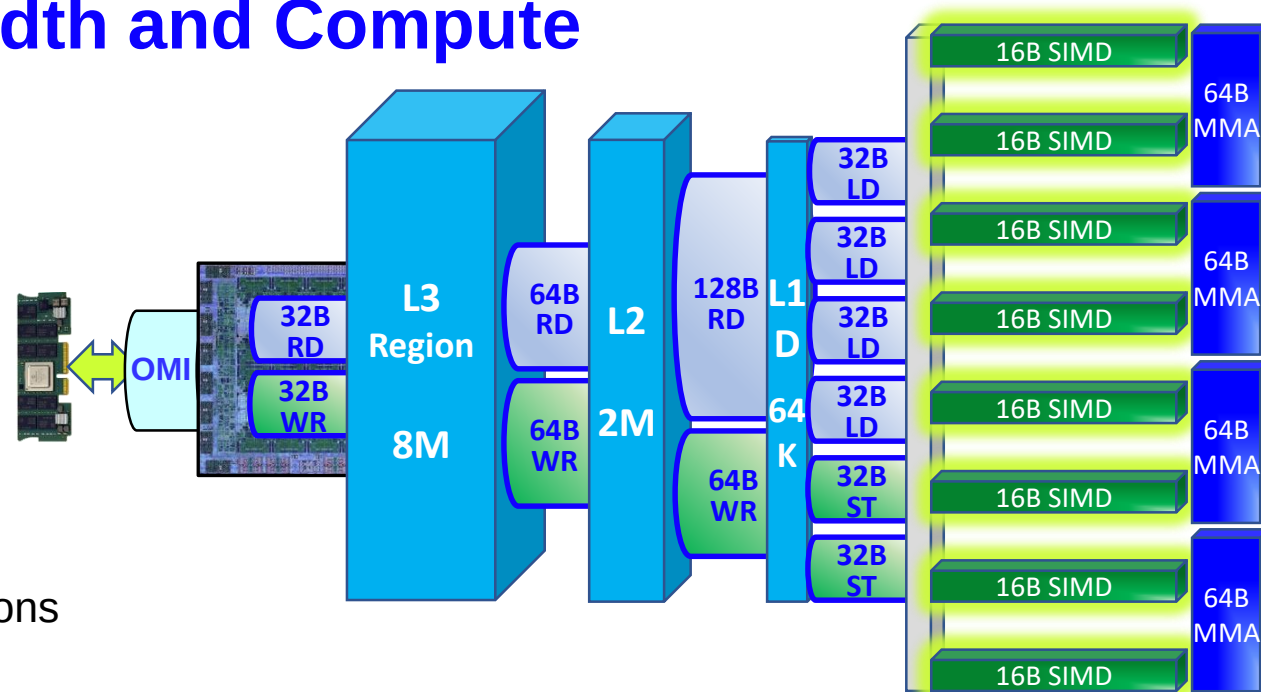
Powerful Core : AI Infused Bandwidth and Compute

B
Y
T
E
S

2x Bytes from all sources

(OMI, L3, L2, L1 caches*)

- 4 32B loads, 2 32B stores per SMT8 Core
 - New ISA or fusion
 - Thread max 2 32B loads, 1 32B store
- OMI Memory to one Core
 - 256 GB/s peak, 120 GB/s sustained
 - With 3x L3 prefetch and memory prefetch extensions



F
L
O
P

2x Bandwidth matched SIMD*

- 8 independent SIMD engines per Core
 - Fixed, float, permute

* versus POWER9

IBM POWER10

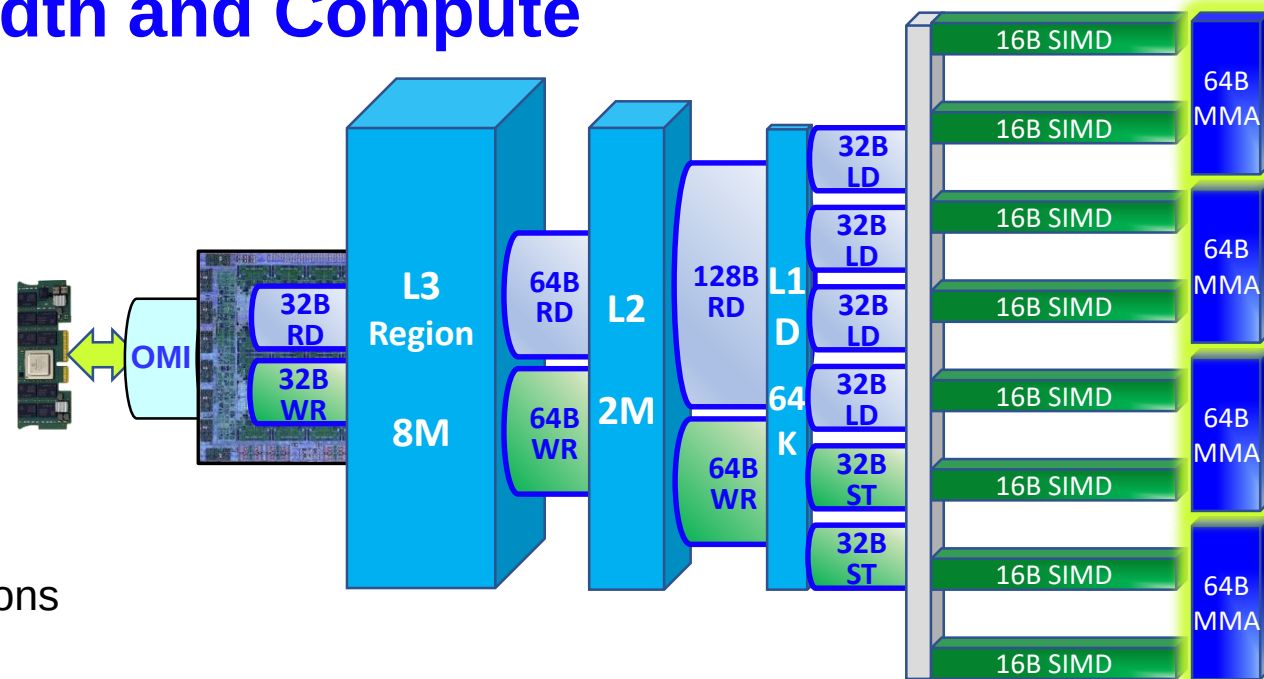
Powerful Core : AI Infused Bandwidth and Compute

BYTES

2x Bytes from all sources

(OMI, L3, L2, L1 caches*)

- 4 32B loads, 2 32B stores per SMT8 Core
 - New ISA or fusion
 - Thread max 2 32B loads, 1 32B store
- OMI Memory to one Core
 - 256 GB/s peak, 120 GB/s sustained
 - With 3x L3 prefetch and memory prefetch extensions



FLOP

2x Bandwidth matched SIMD*

- 8 independent SIMD engines per Core
 - Fixed, float, permute

4-32x Matrix Math Acceleration*

- 4 512b engines per core = 2048b results / cycle
 - Matrix math outer products: $A \leftarrow \{\pm\}A \{\pm\}XY^T$
 - Double, Single, Reduced precision

Rank	Operand Type (X,Y)			Accumulator	Peak [FL]OPS / cycle		
k	Type	X	Y ^T	A	Instruction	Thread	SMT8 Core
1	Float-64 DP	4×1	1×2	4×2 (Fp-64)	16	32	64
	Float-32 SP	4×1	1×4		32	64	128
2	Float-16 HP	4×2	2×4	4×4 (Fp-32)	64	128	256
	Bfloat-16 HP	4×2	2×4				
	Int-16	4×2	2×4				
4	Int-8	4×4	4×4	4×4 (Int-32)	128	256	512
8	Int 4	4×8	8×4		256	512	1024

* versus POWER9

IBM POWER10

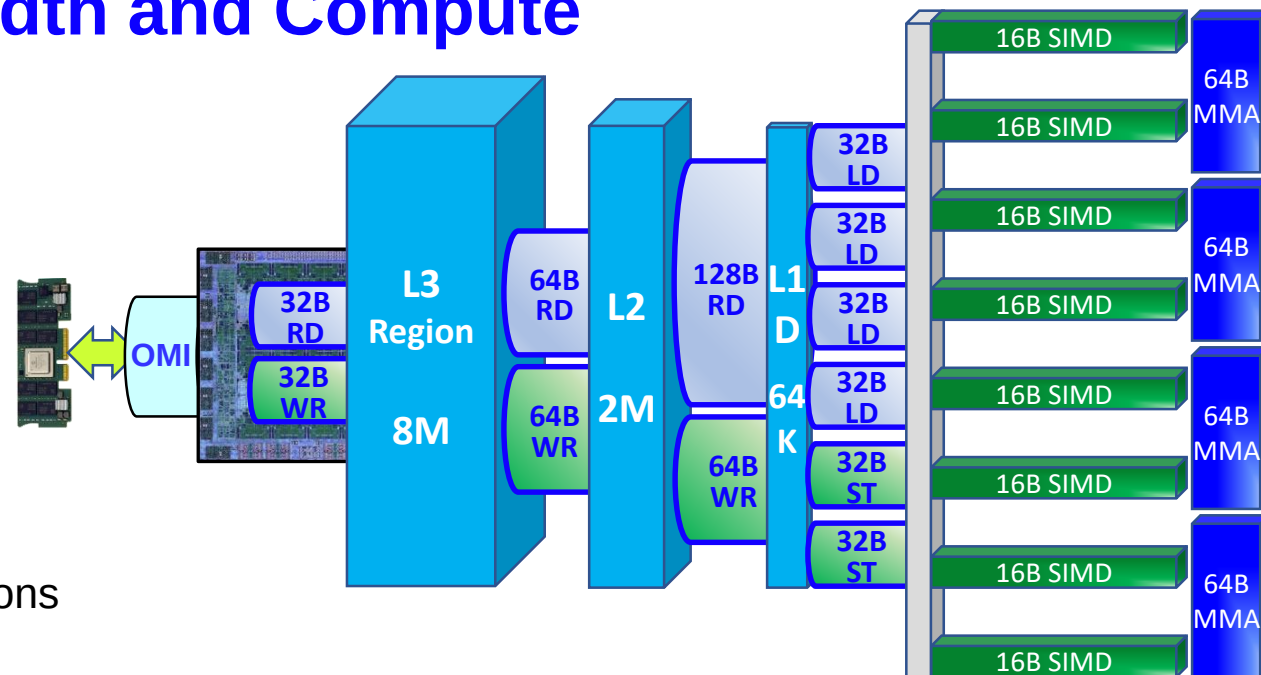
Powerful Core : AI Infused Bandwidth and Compute

BYTES

2x Bytes from all sources

(OMI, L3, L2, L1 caches*)

- 4 32B loads, 2 32B stores per SMT8 Core
 - New ISA or fusion
 - Thread max 2 32B loads, 1 32B store
- OMI Memory to one Core
 - 256 GB/s peak, 120 GB/s sustained
 - With 3x L3 prefetch and memory prefetch extensions



FLP

2x Bandwidth matched SIMD*

- 8 independent SIMD engines per Core
 - Fixed, float, permute

4-32x Matrix Math Acceleration*

- 4 512b engines per core = 2048b results / cycle
 - Matrix math outer products: $A \leftarrow \{\pm\}A \{\pm\}XY^T$
 - Double, Single, Reduced precision

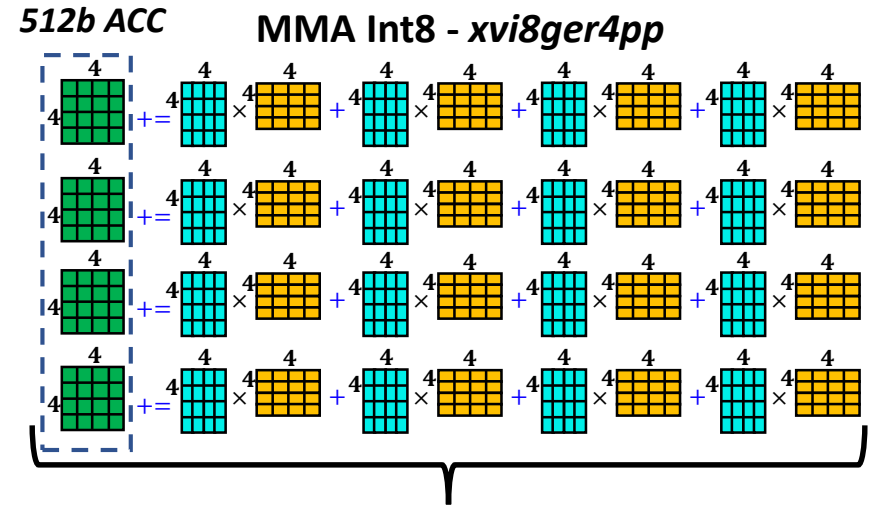
Rank	Operand Type (X,Y)			Accumulator	Peak [FL]OPS / cycle		
k	Type	X	Y ^T	A	Instruction	Thread	SMT8 Core
1	Float-64 DP	4×1	1×2	4×2 (Fp-64)	16	32	64
	Float-32 SP	4×1	1×4		32	64	128
2	Float-16 HP	4×2	2×4	4×4 (Fp-32)	64	128	256
	Bfloat-16 HP	4×2	2×4				
	Int-16	4×2	2×4				
4	Int-8	4×4	4×4	4×4 (Int-32)	128	256	512
8	Int 4	4×8	8×4		256	512	1024

* versus POWER9

IBM POWER10

AI Infused Core: Inference Acceleration

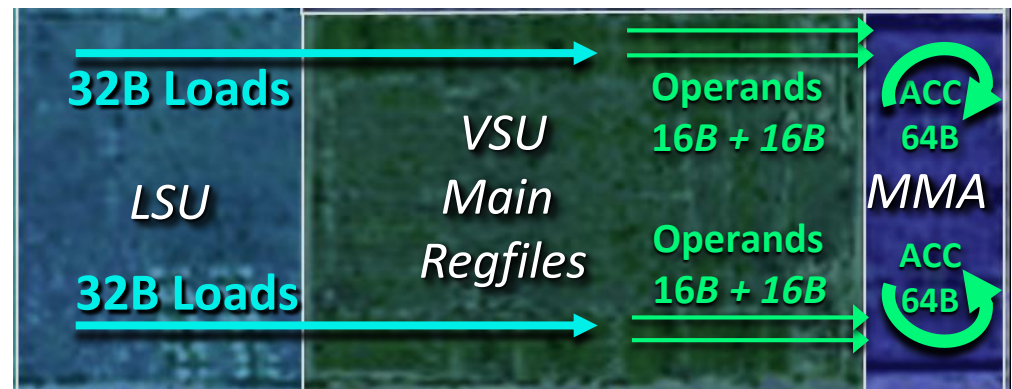
- 4x+ per core throughput
- 3x → 6x thread latency reduction (SP, int8)*
- **POWER10 Matrix Math Assist (MMA) instructions**
 - 8 512b architected Accumulator (ACC) Registers
 - 4 parallel units per SMT8 core
- **Consistent VSR 128b register architecture**
 - Minimal SW ecosystem disruption – no new register state
 - Application performance via updated library (OpenBLAS, etc.)
 - POWER10 aliases 512b ACC to 4 128b VSR's
 - Architecture allows redefinition of ACC
- **Dense-Math-Engine microarchitecture**
 - Built for data re-use algorithms
 - Includes separate physical register file (ACC)
 - 2x efficiency vs. traditional SIMD for MMA



4 per cycle per SMT8 core

Matrix Optimized / High Efficiency

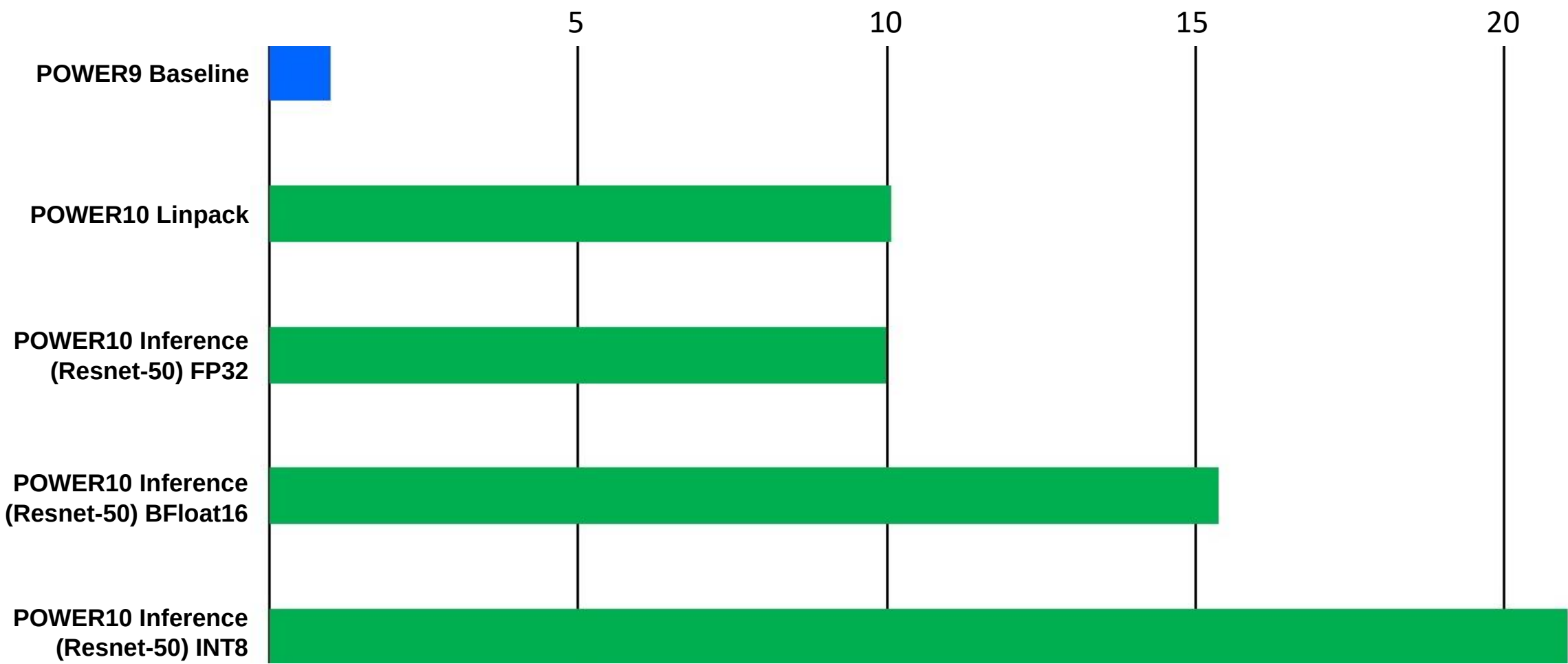
Result data remains local to compute



Inference Accelerator dataflow (2 per SMT8 core)

* versus POWER9

POWER10 SIMD / AI Socket Performance Gains



(Performance assessments based upon pre-silicon engineering analysis of POWER10 dual-socket server offering vs POWER9 dual-socket server offering)