

# MXCore: Rethinking Matrix Multiplication Accelerators for Microscaling (MX) Arithmetic

Jayanth Jonnalagadda<sup>1\*</sup>

Paper ID: 336

Gamze İslamoğlu<sup>1\*</sup>, Arpan Suravi Prasad<sup>1</sup>, Angelo Garofalo<sup>2</sup>, Francesco Conti<sup>2</sup>, and Luca Benini<sup>1,2</sup>

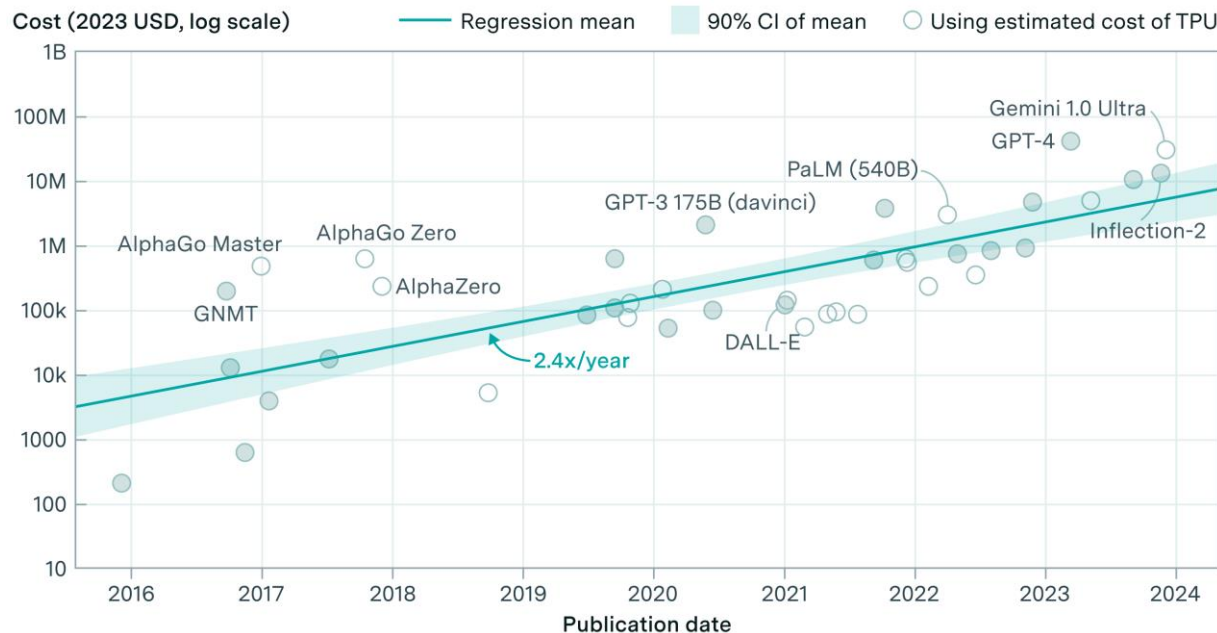
<sup>1</sup>Integrated Systems Laboratory (IIS), ETH Zürich, Switzerland

<sup>2</sup>DEI, University of Bologna, Italy

\*Equal Contribution

# Rapidly Rising Costs of Frontier AI

- Frontier AI models are rapidly growing.
  - Billions of parameters, peta-operation-level compute, and hundreds of gigabytes of memory.
  - Steep training cost increase of **2.4×** per year.
  - Expected to cross a billion dollars by **2027**.

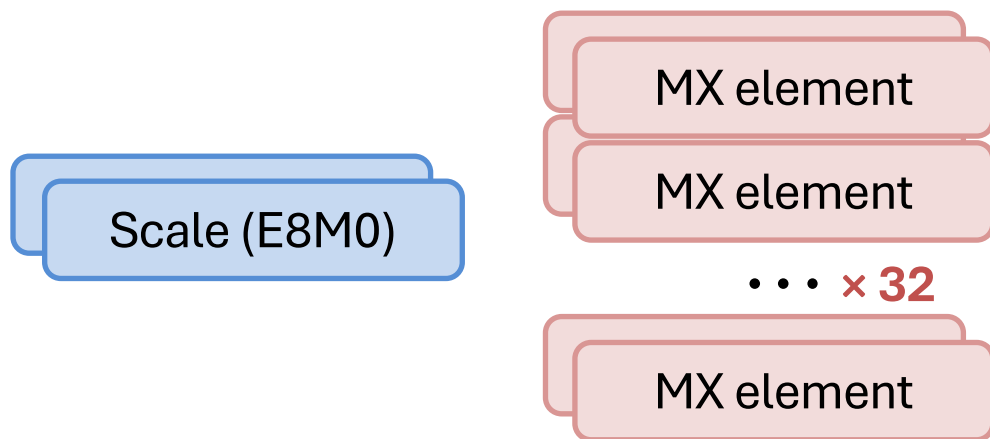


Ben Cottier, Robi Rahman, Loredana Fattorini, Nestor Maslej, Tamay Besiroglu, and David Owen. 2025. The rising costs of training frontier AI models. doi:10. 48550/arXiv.2405.21015.



# Low-Precision AI is Promising

- Microscaling (MX) Data Formats.
  - Open Compute Project (2023).
  - Narrow bit-width (4 – 8 bit) formats.
    - MXFP4, MXFP6, MXFP8, and INT8.
  - Block-level scaling.
    - Wide dynamic range and fine-grained control.
  - Accuracies close to full precision.



**OPEN**  
Compute  
Project®



**NVIDIA®**

**arm**

**AMD**

**Meta intel®**

**Microsoft**

**Qualcomm**

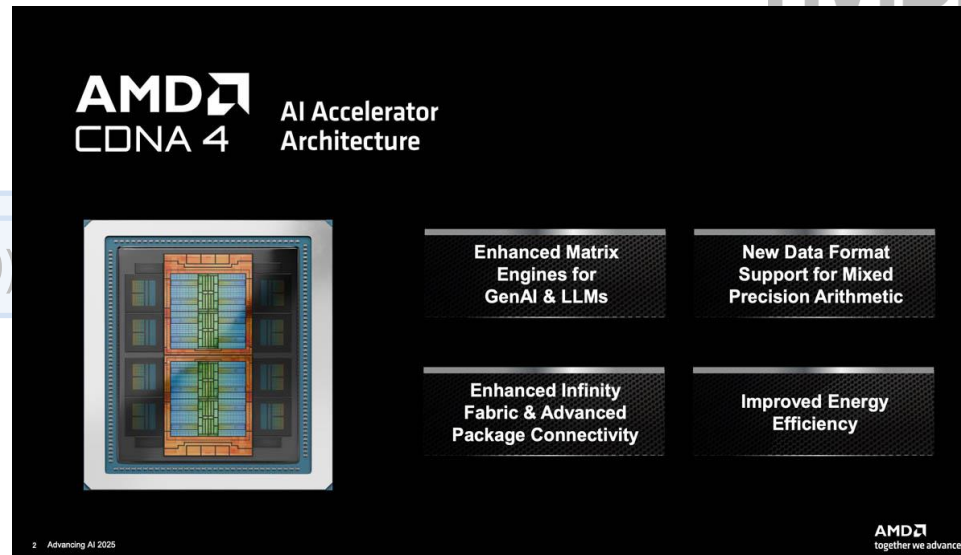
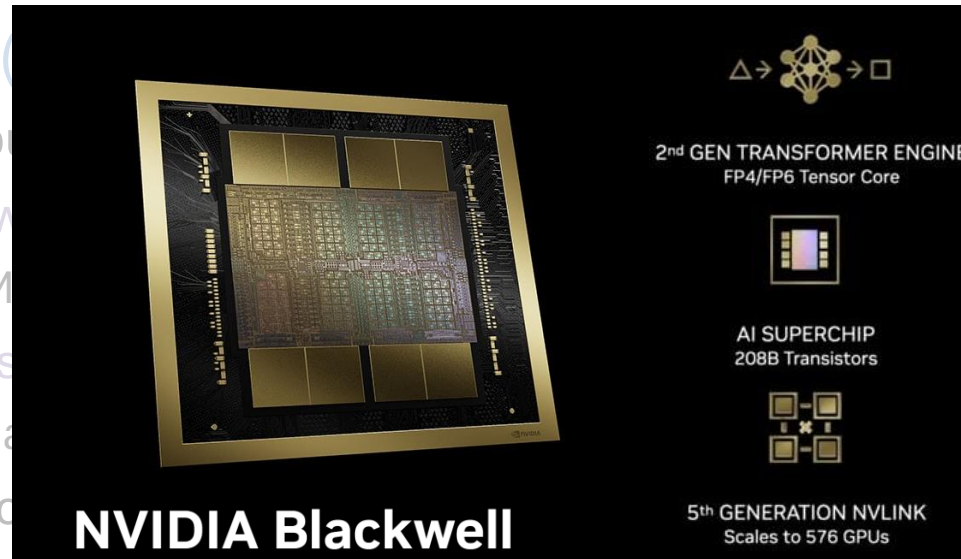
<https://www.opencompute.org/documents/ocp-microscaling-formats-mx-v1-0-spec-final-pdf>



# Industry-leading MX Architectures

- Microscaling (

- Open Compute
- Narrow bit-width
  - MXFP4, M
- Block-level s
- Wide dynamic range
- Accuracies c



Scale (E8M0)

2 Advancing AI 2025

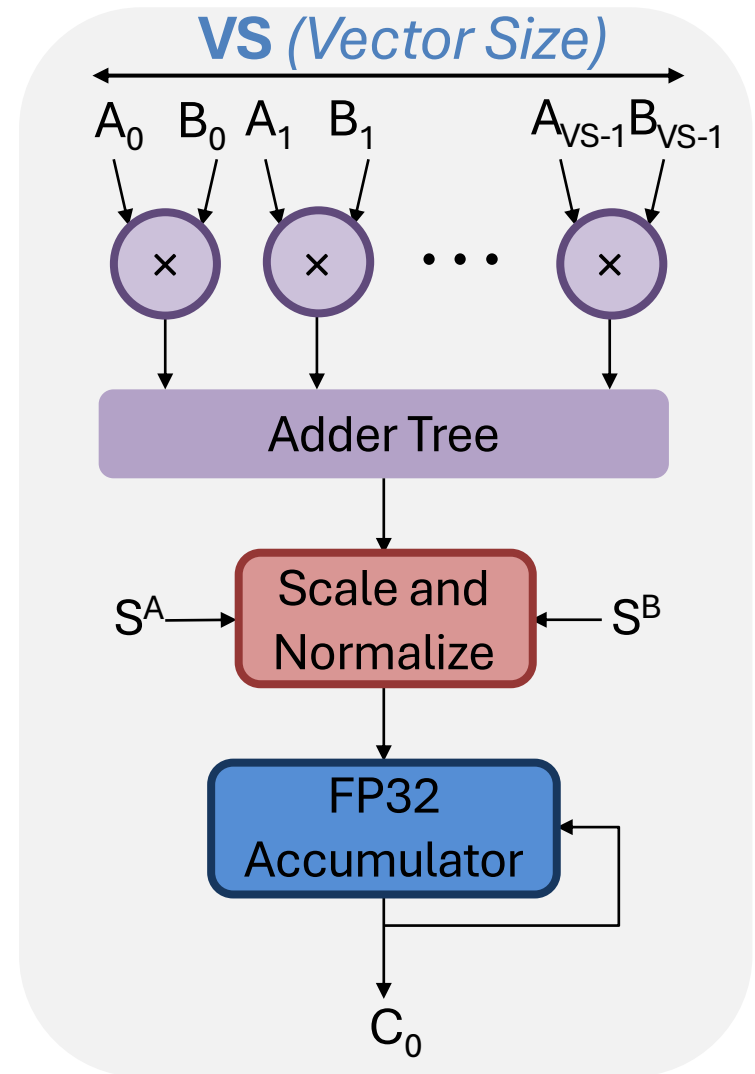
AMD  
together we advance...





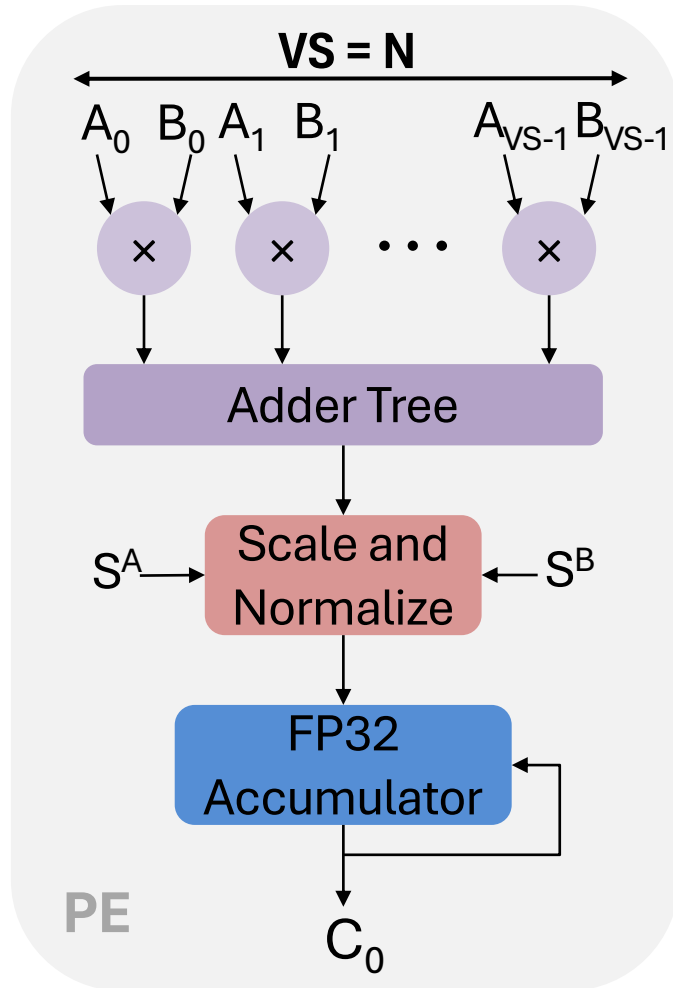
# Unique Aspects of MX Arithmetic

- Operand Precision Asymmetry.
  - Narrow multipliers.
    - 4 – 8 bit operands.
  - Wide accumulations.
    - Accumulation in FP32.
- Block-level Scaling.
  - Intra-block: directly accumulate.
  - Inter-block: scale before accumulate.

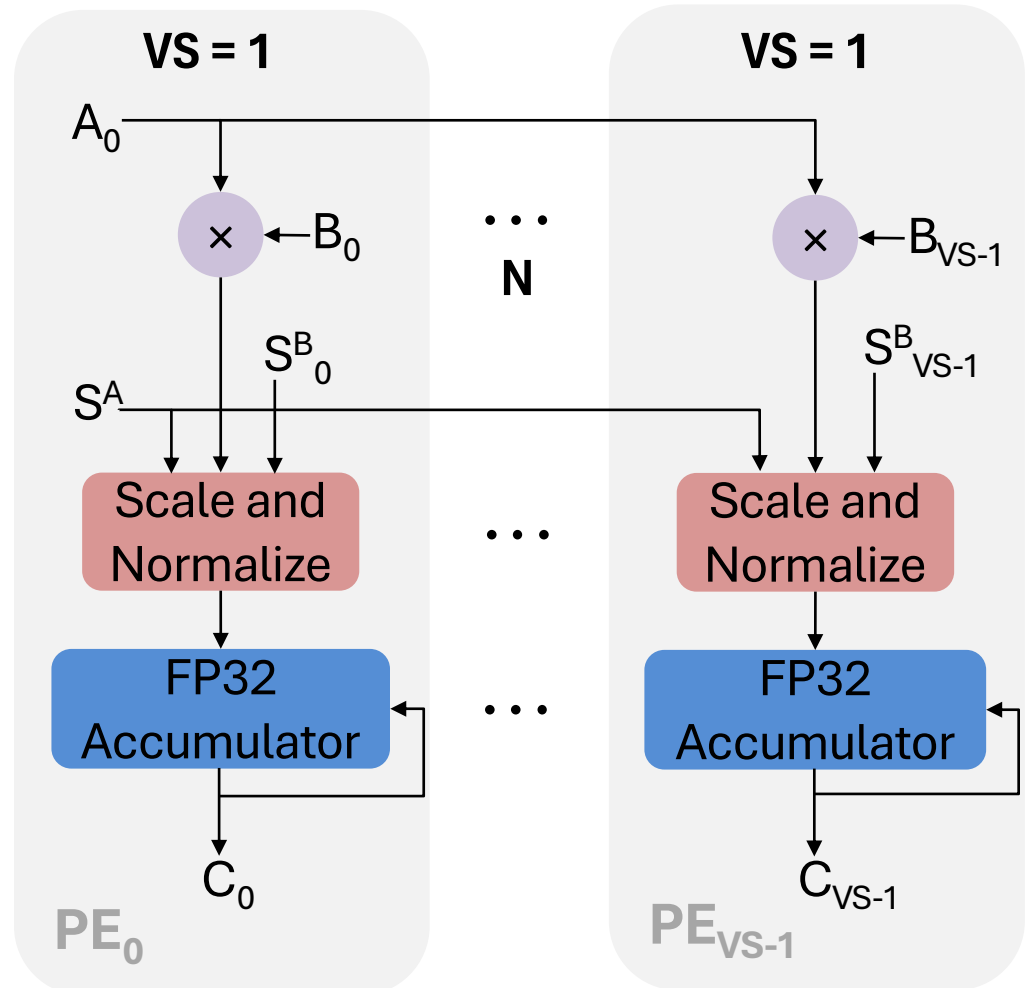


# Inner vs. Outer Product for MX

## Inner Product

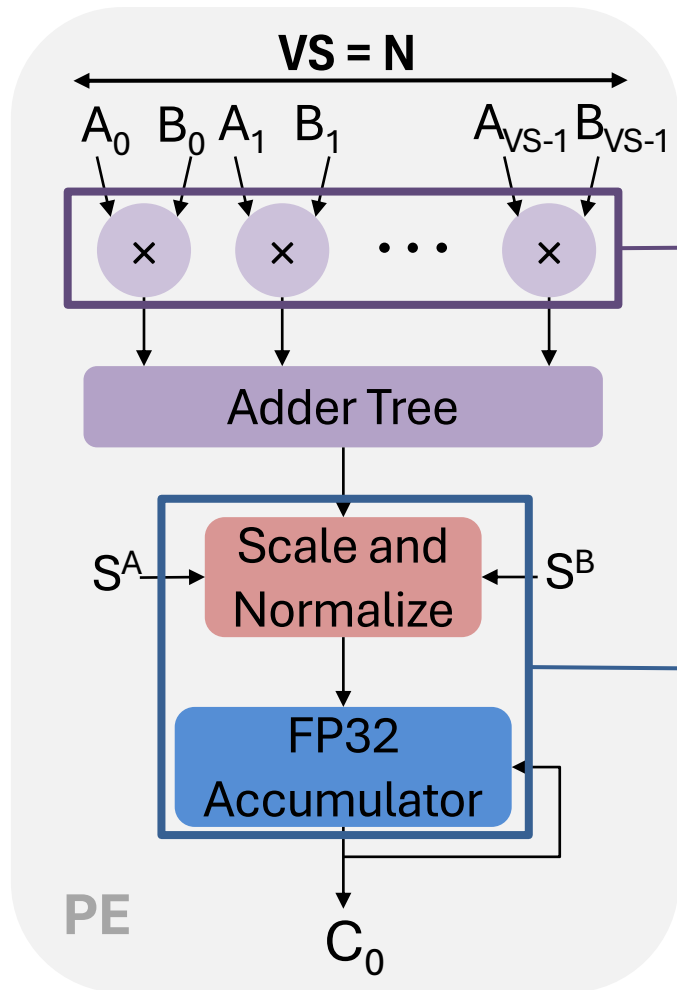


## Outer Product

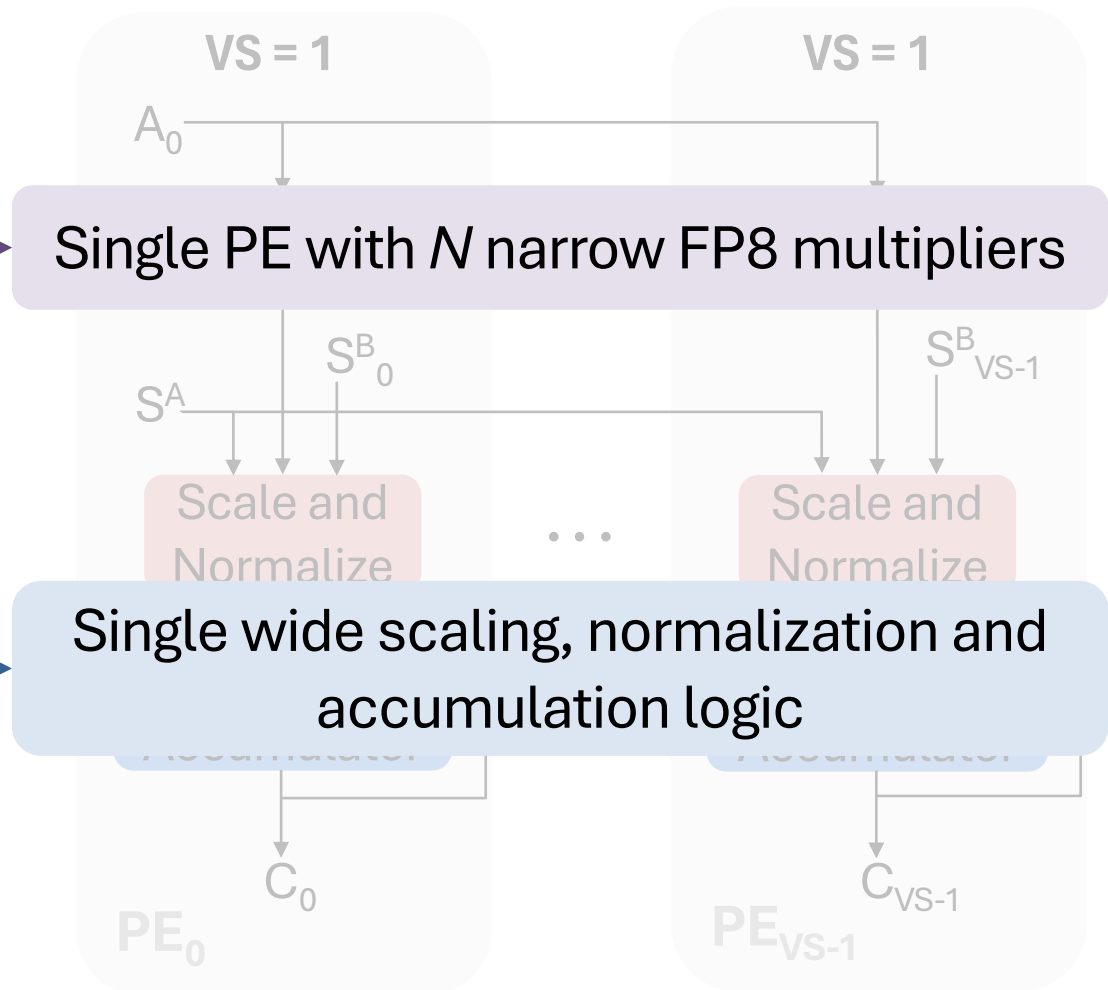


# Inner vs. Outer Product for MX

## Inner Product

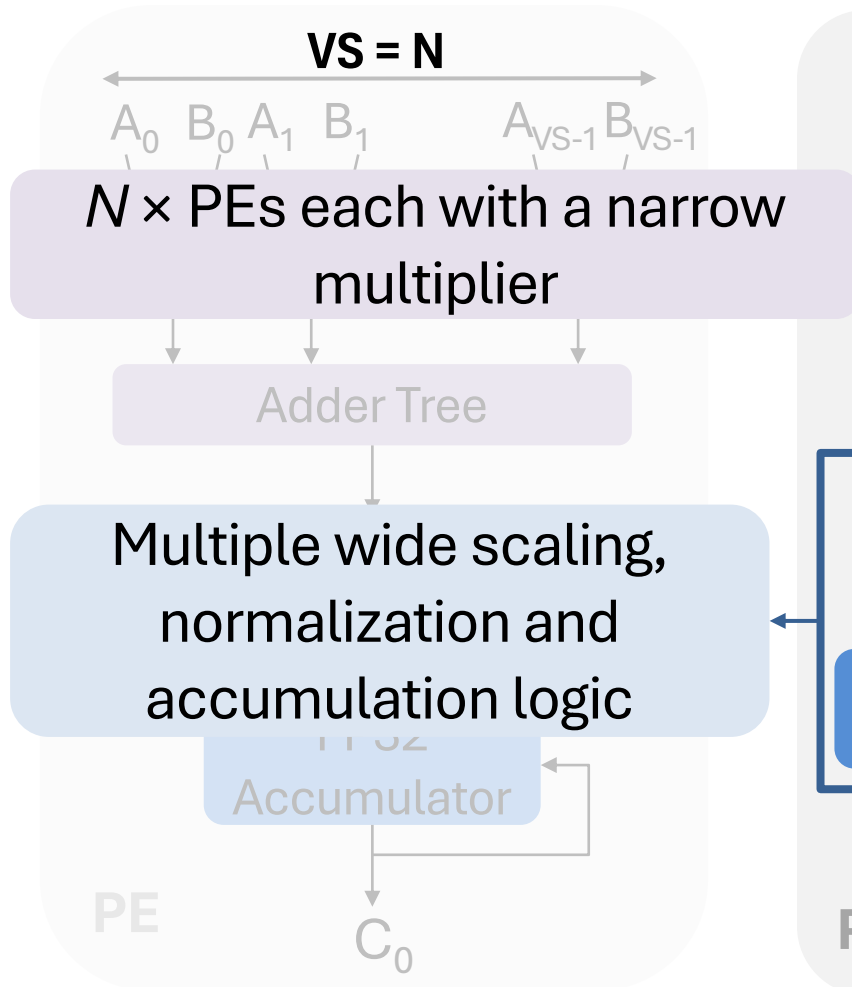


## Outer Product

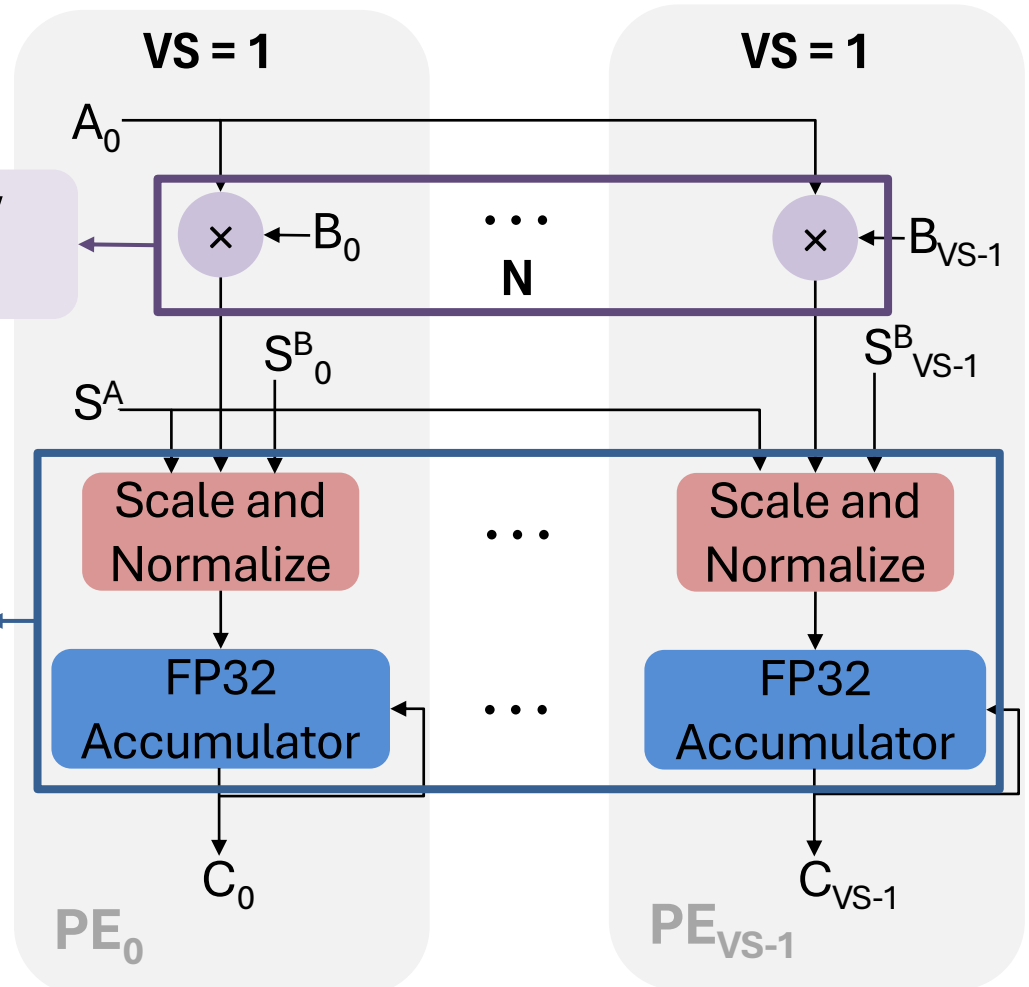


# Inner vs. Outer Product for MX

## Inner Product



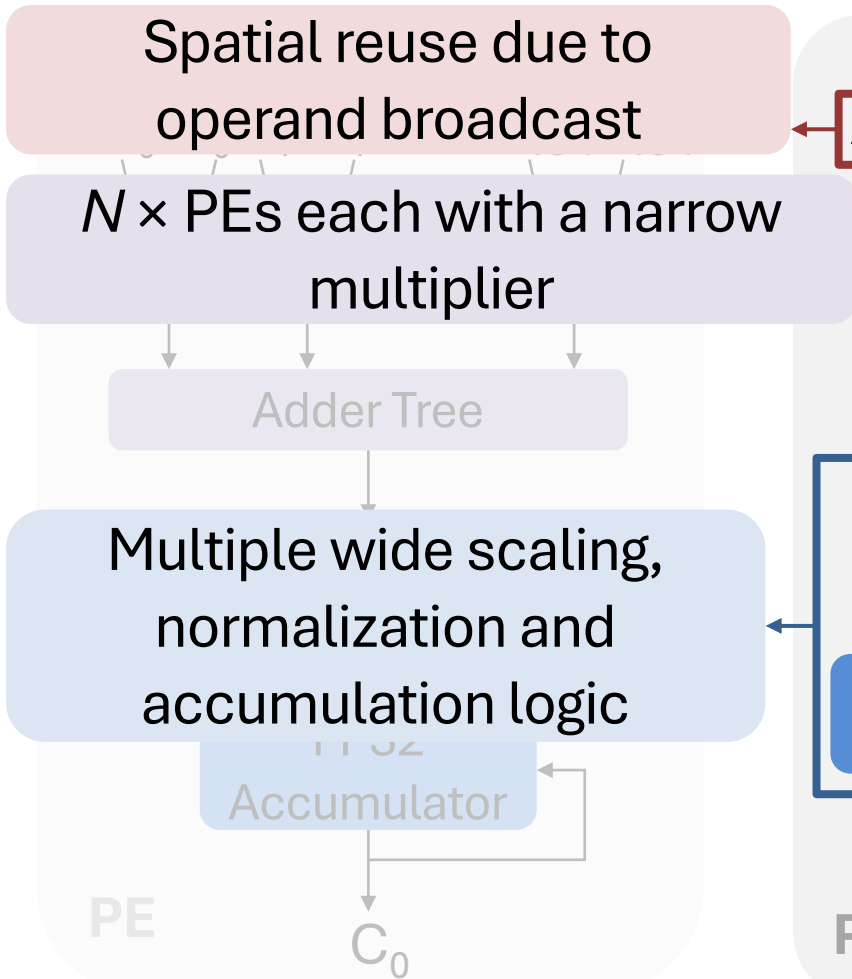
## Outer Product



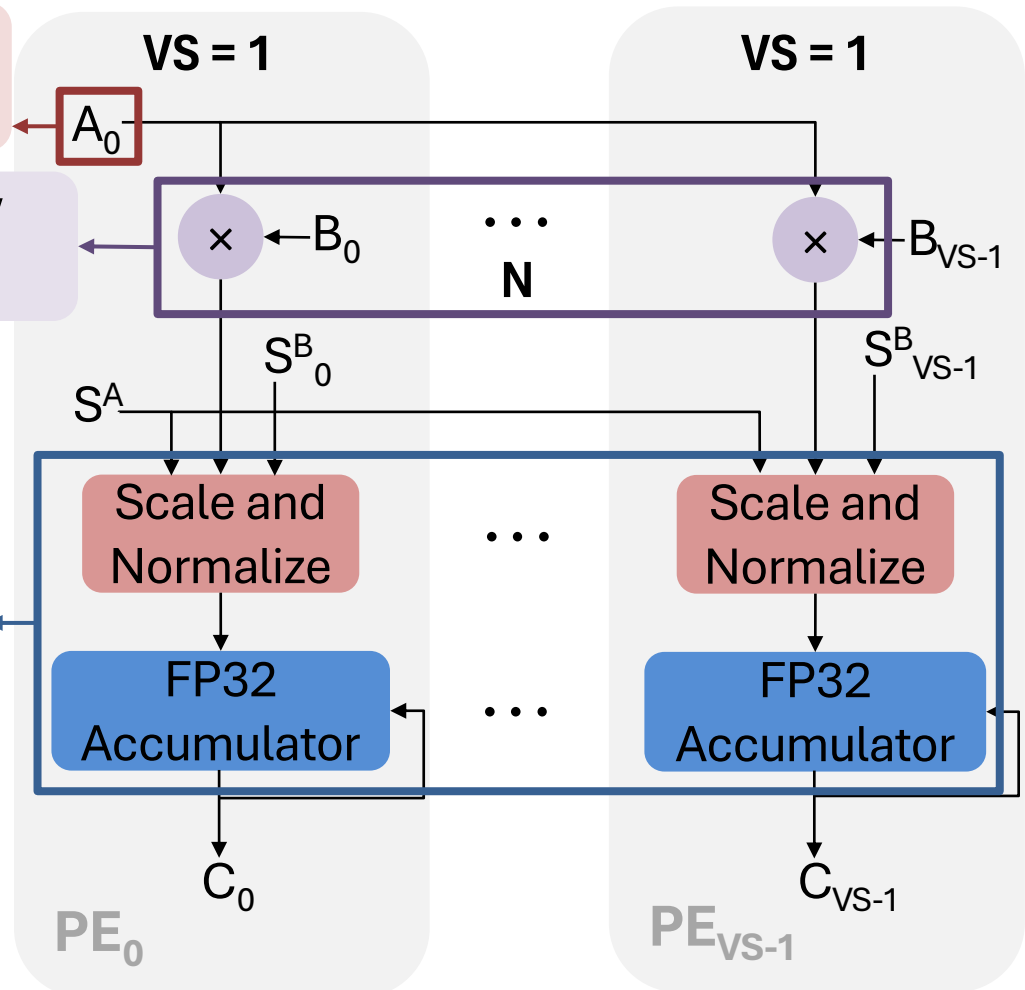


# Inner vs. Outer Product for MX

## Inner Product

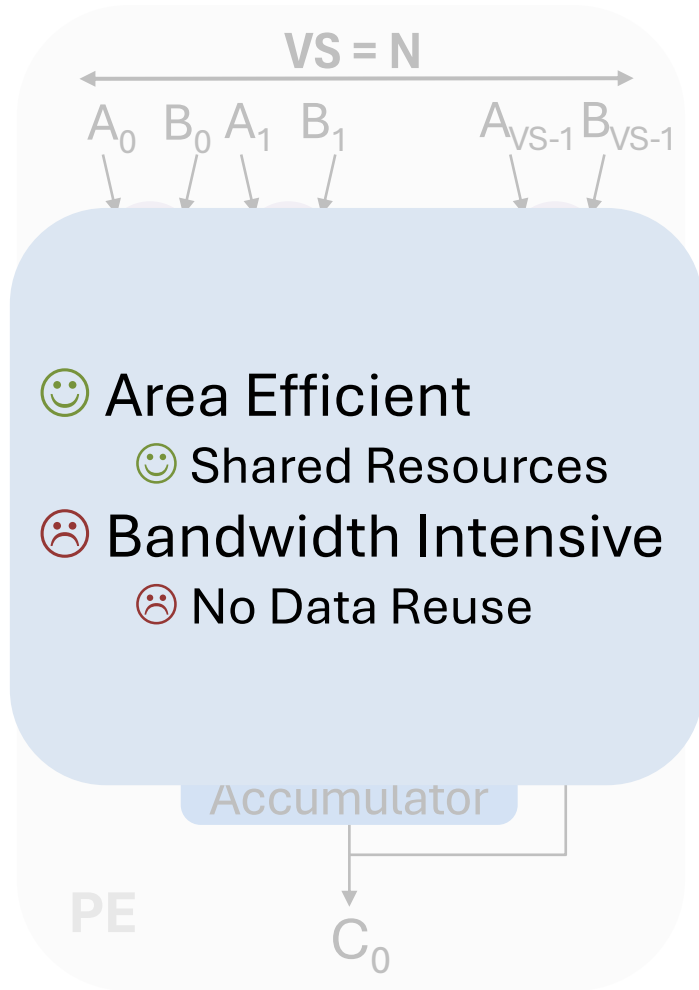


## Outer Product

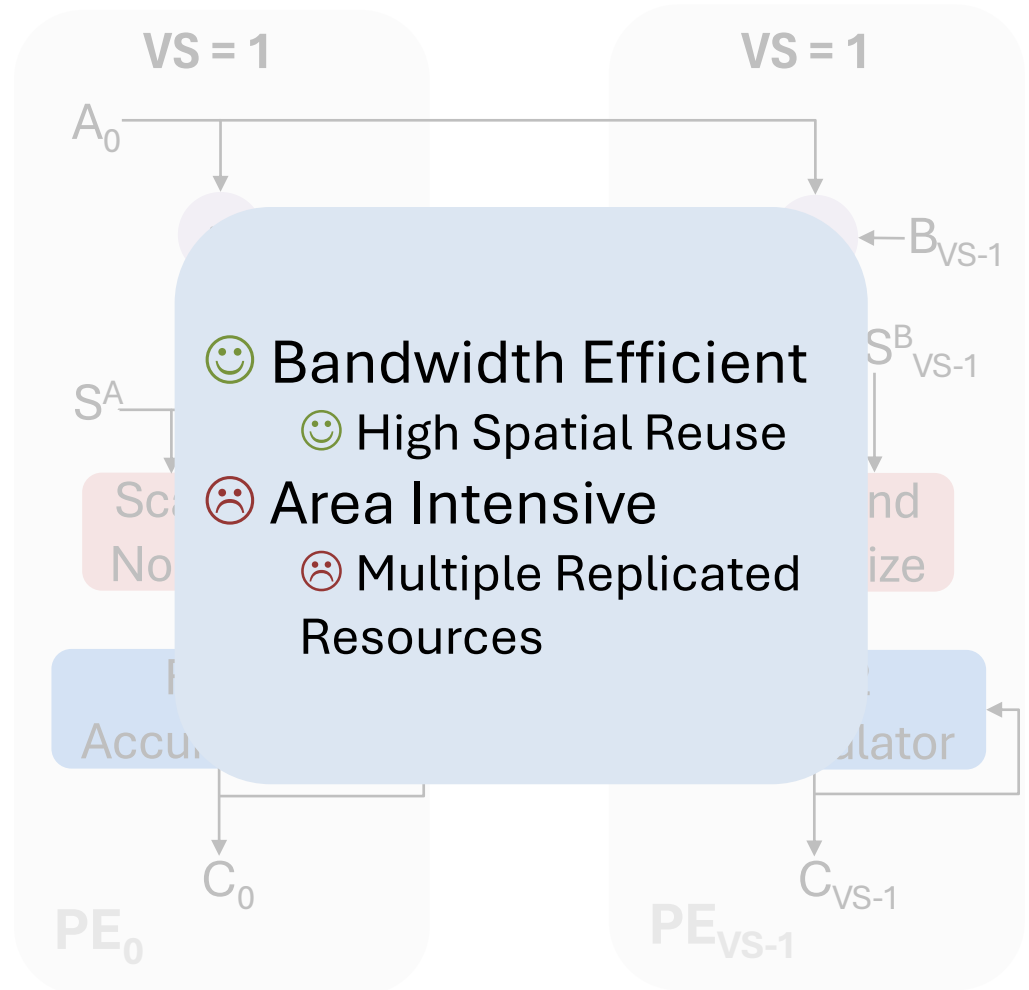


# Inner vs. Outer Product Trade-offs

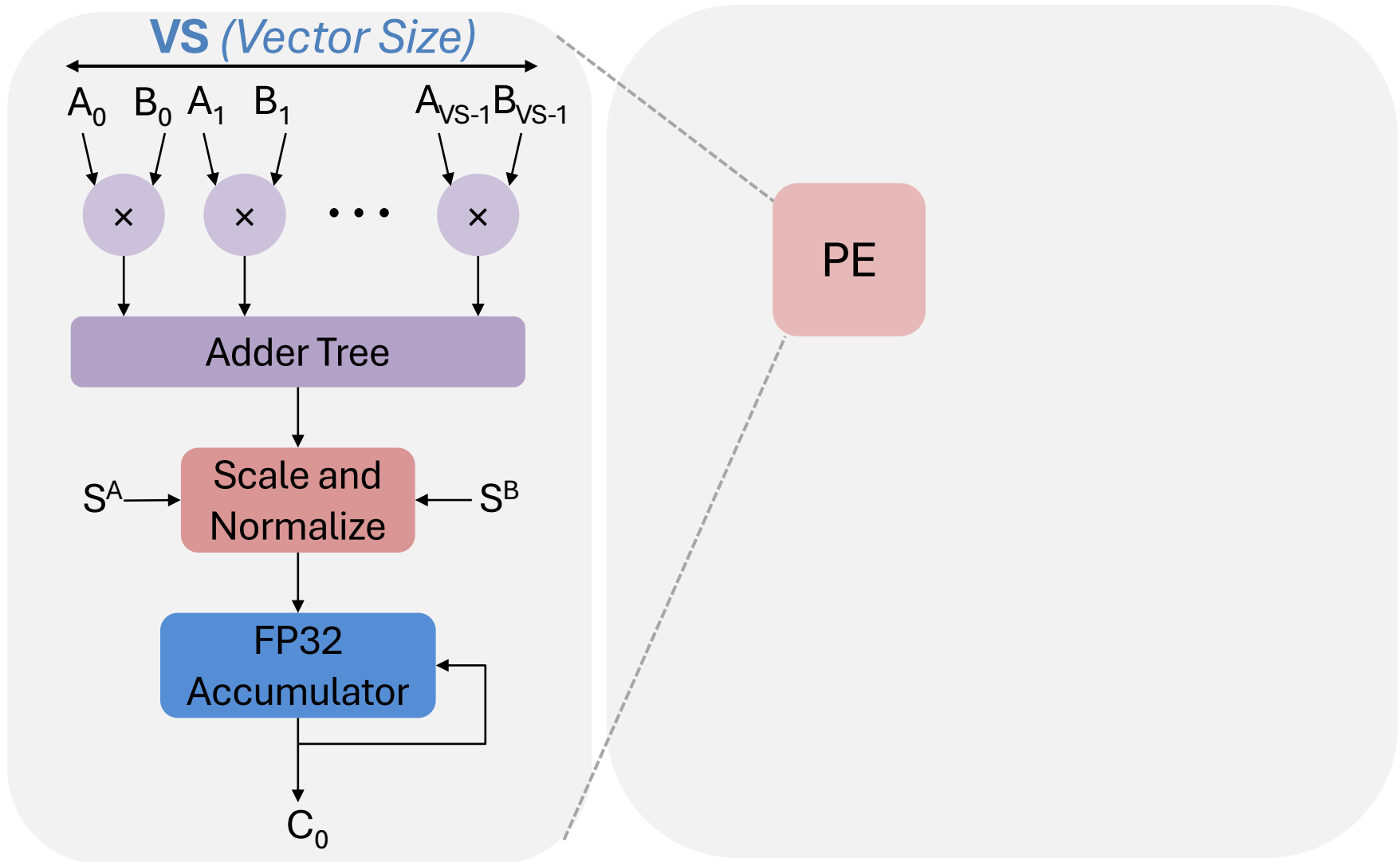
## Inner Product



## Outer Product

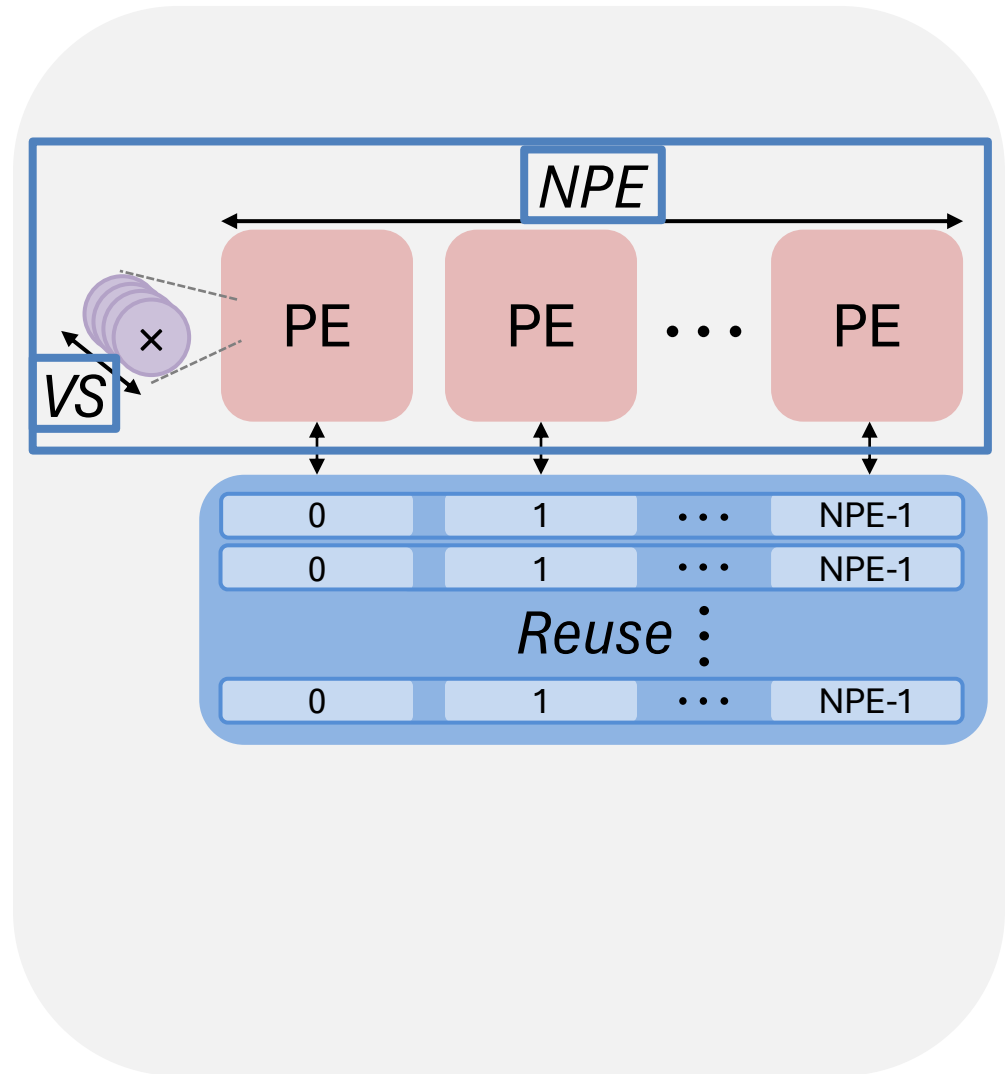


# Architectural Framework of MXCore

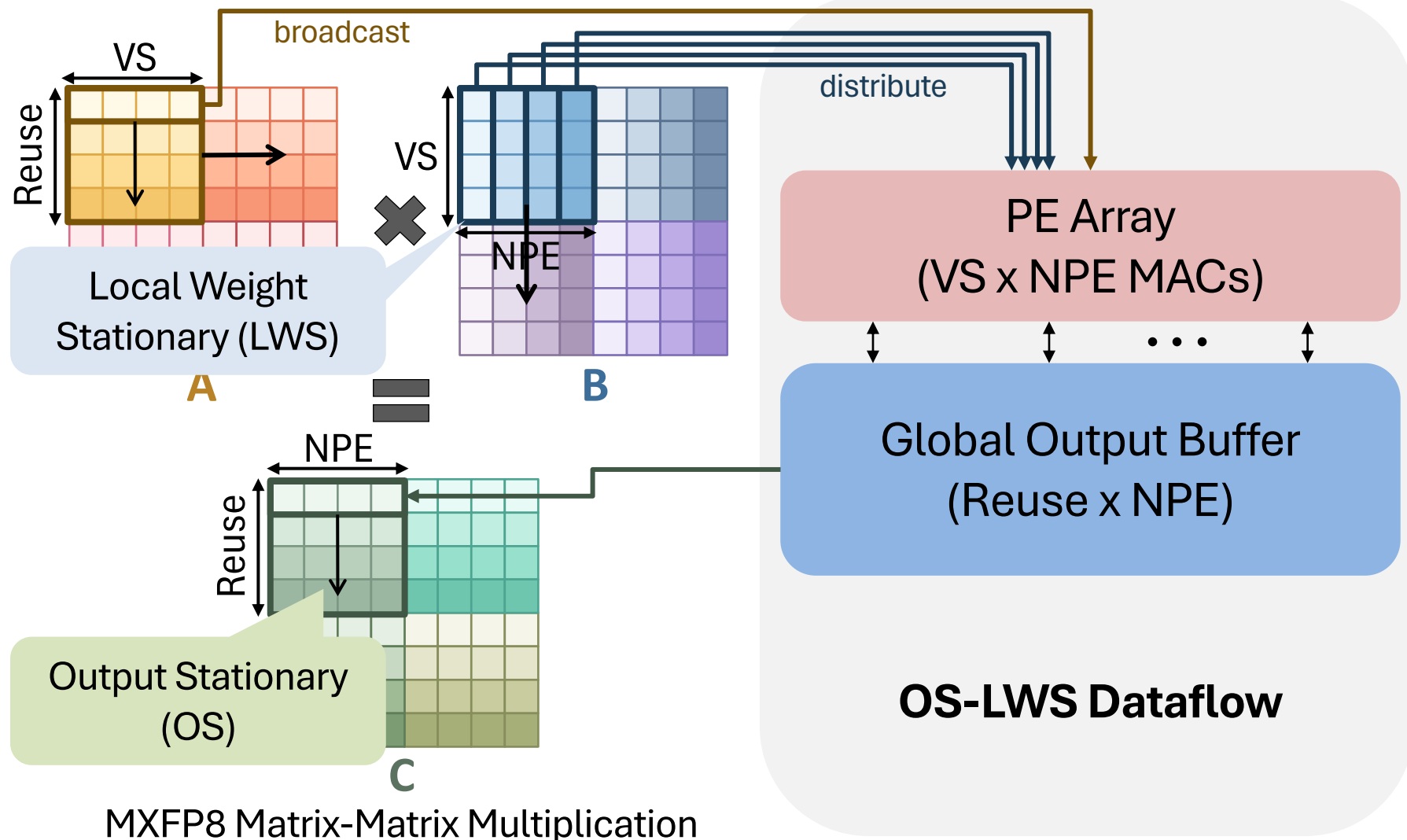


# Architectural Framework of MXCore

- VS (Vector Size)
  - Inner product size.
  - Dot product length per PE.
- NPE (Number of PEs)
  - Outer product size.
  - Number of PEs in the array.
- Reuse
  - Output reuse factor.
  - Locally accumulated FP32 partial results.
- Peak Throughput
  - $VS \times NPE$  MACs/cycle.

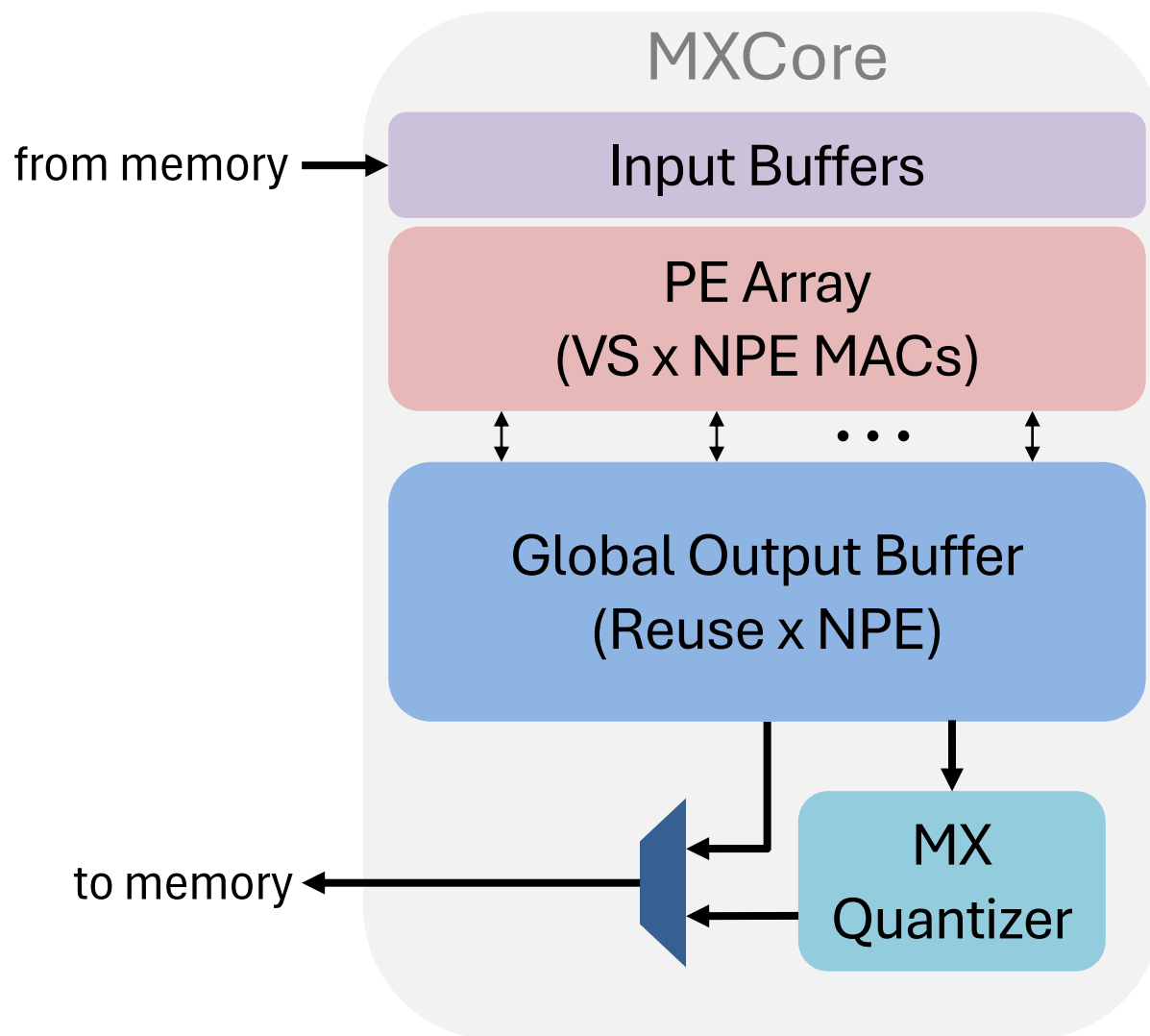


# Workload Mapping and Dataflow



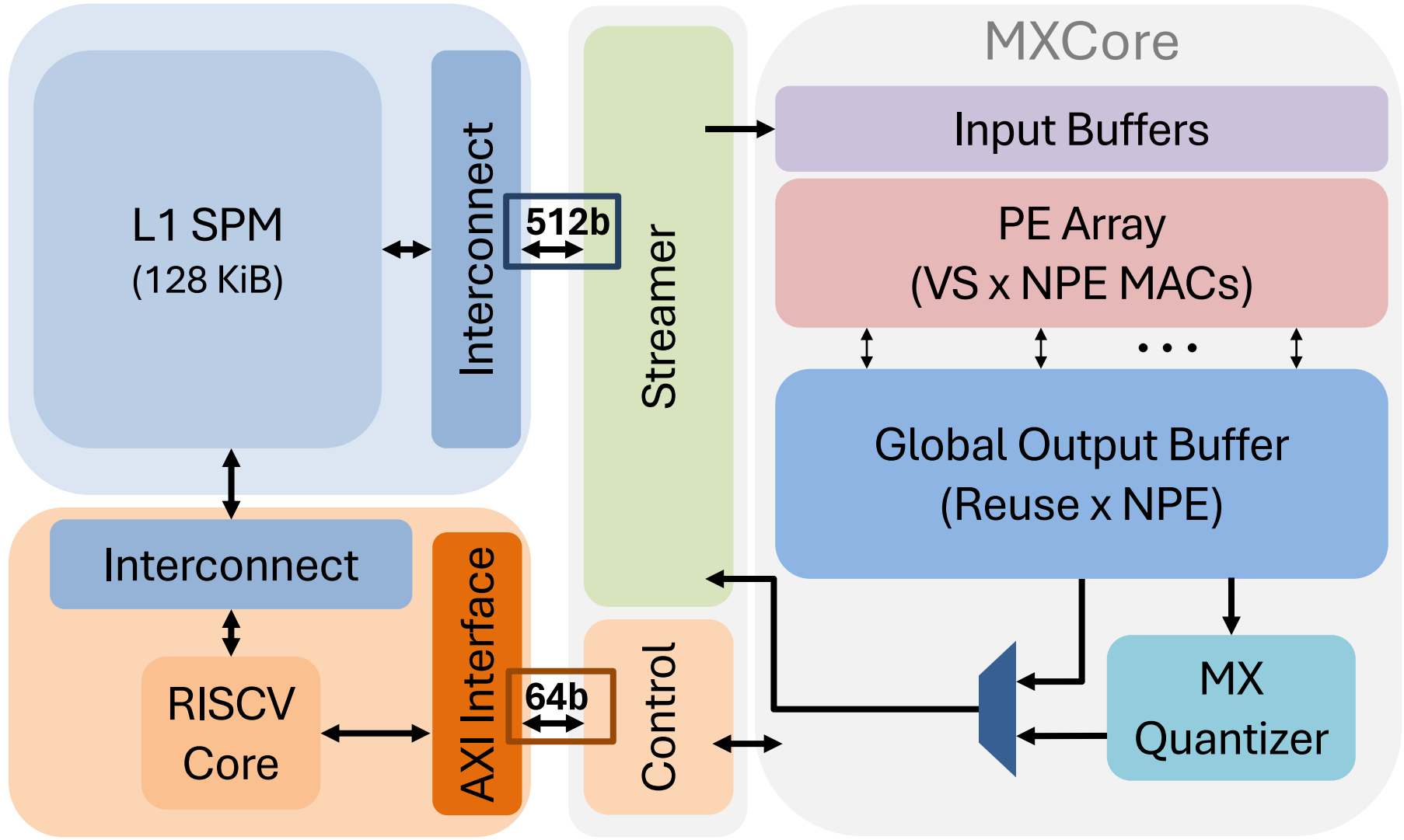
MXFP8 Matrix-Matrix Multiplication  
(Scale matrices omitted for brevity)

# Architecture of MXCore



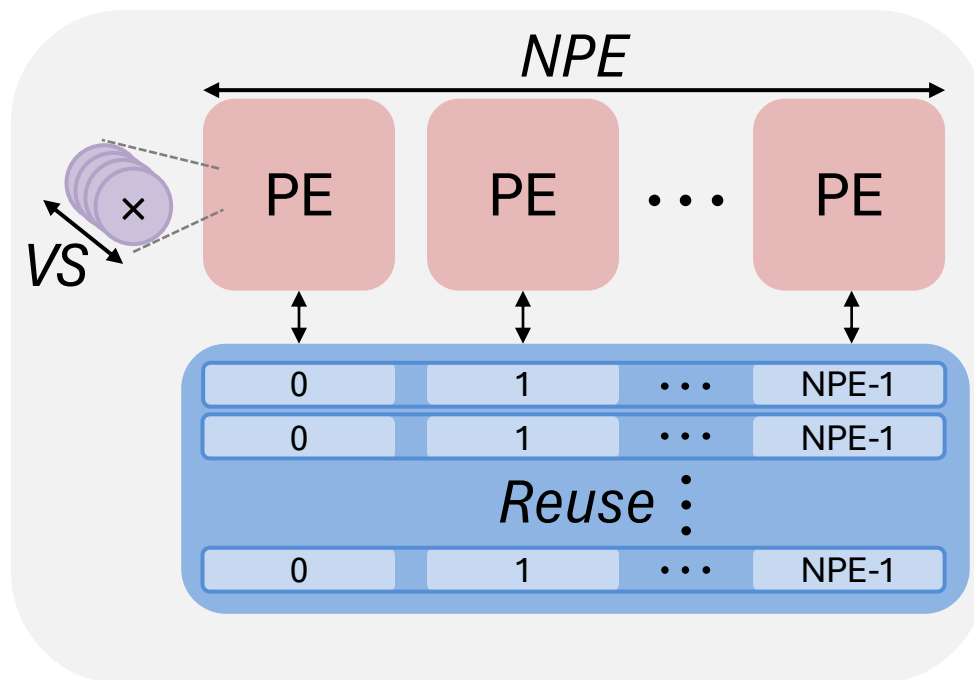


# System-level Integration of MXCore

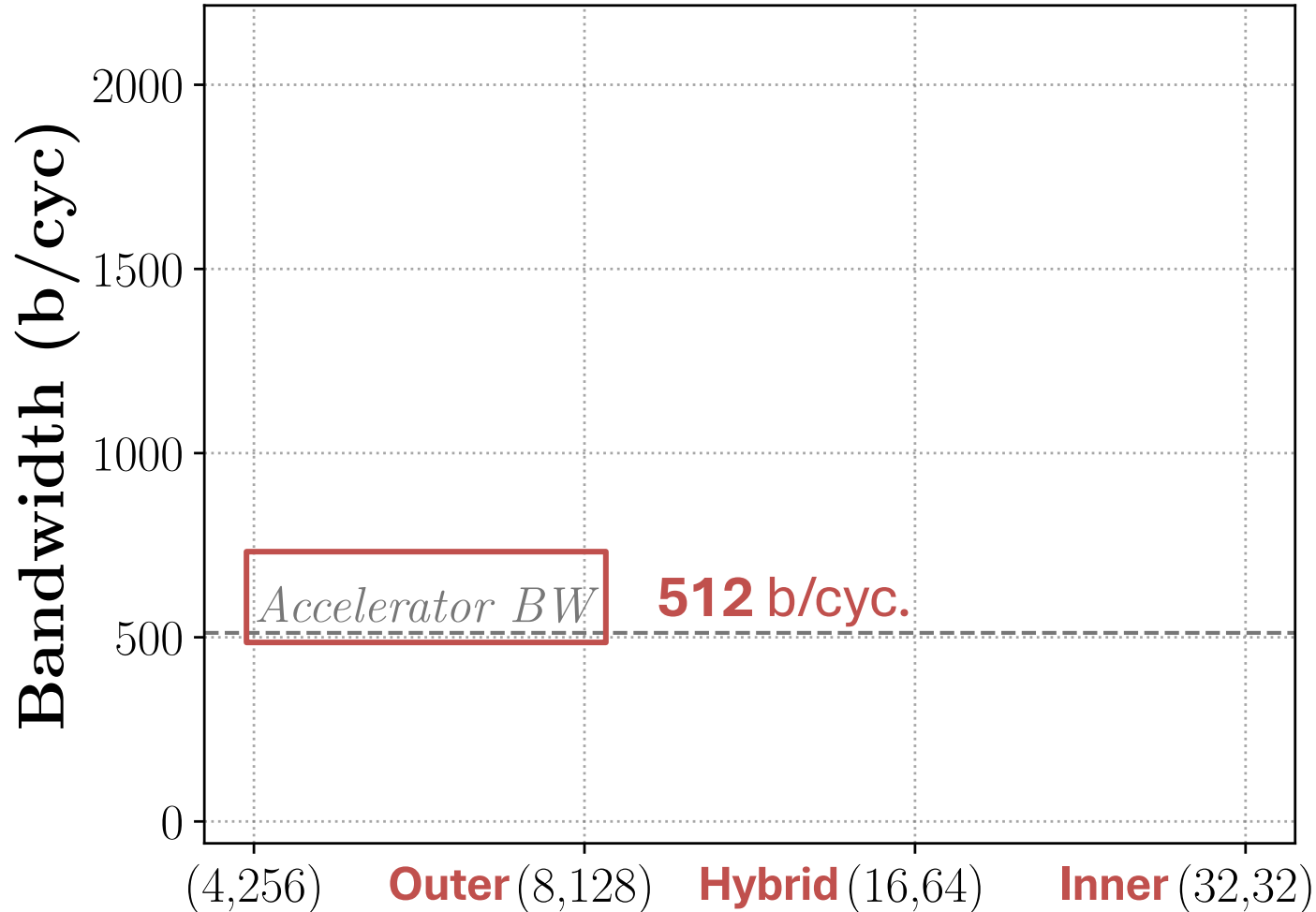


# Target Configurations and Parameters

- Iso-Peak-Throughput
  - Fixed #MACs = **1024**.
    - $VS \times NPE = \text{1024 MACs}$ .
    - Representative of accelerators targeting edge-to mid-range inference platforms.
- Equi-boundary-bandwidth
  - Accelerator BW = **512 b/cyc**.
- (VS, NPE)
  - (**32**, **32**): Inner
  - (**16**, **64**): Hybrid
  - (**8**, **128**): Outer
- $Reuse \in \{32, 64, 128\}$

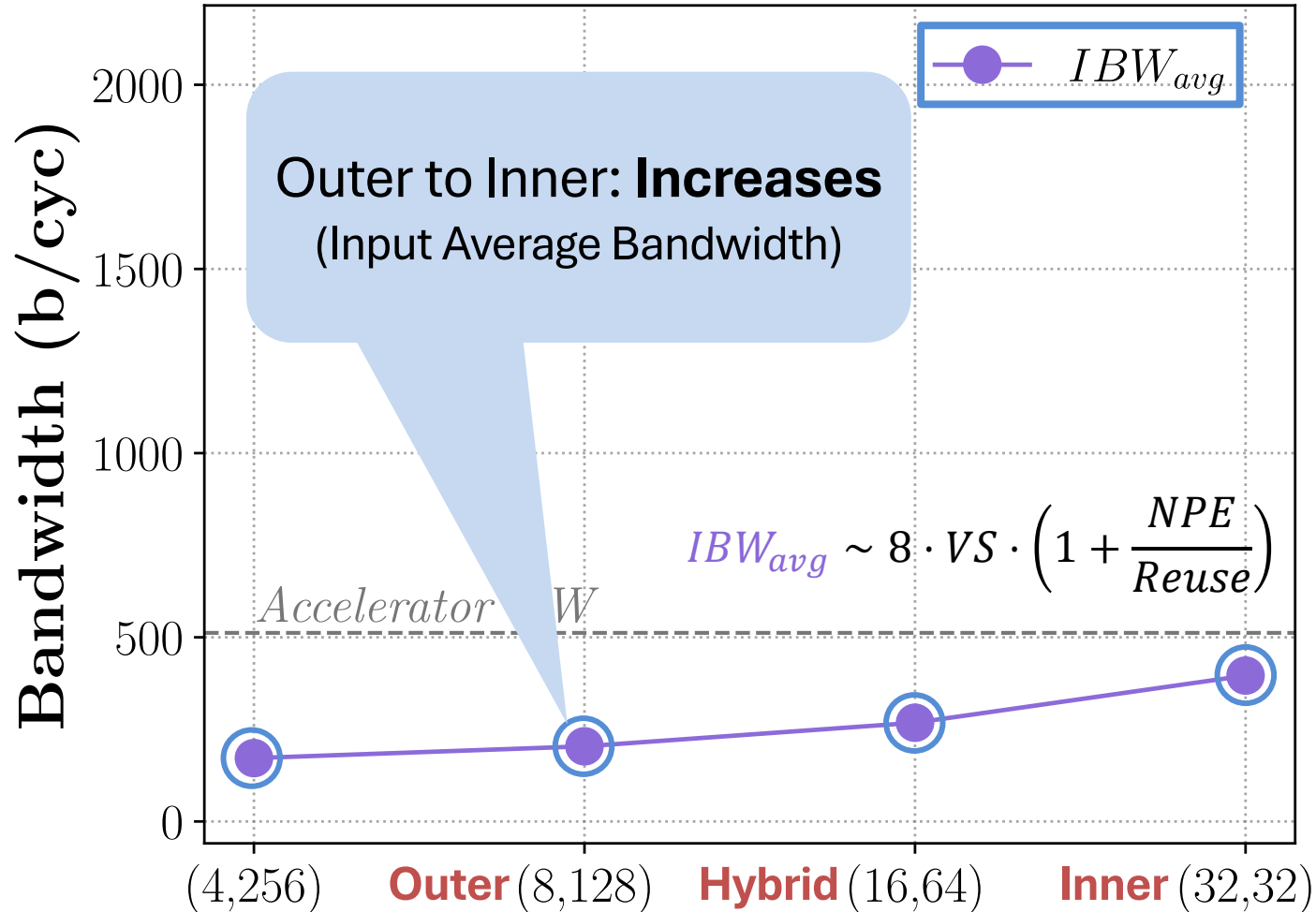


# Bandwidth: Inner vs Outer



(Inner, Outer) sizes of MXCore: Inner x Outer = 1024 MACs

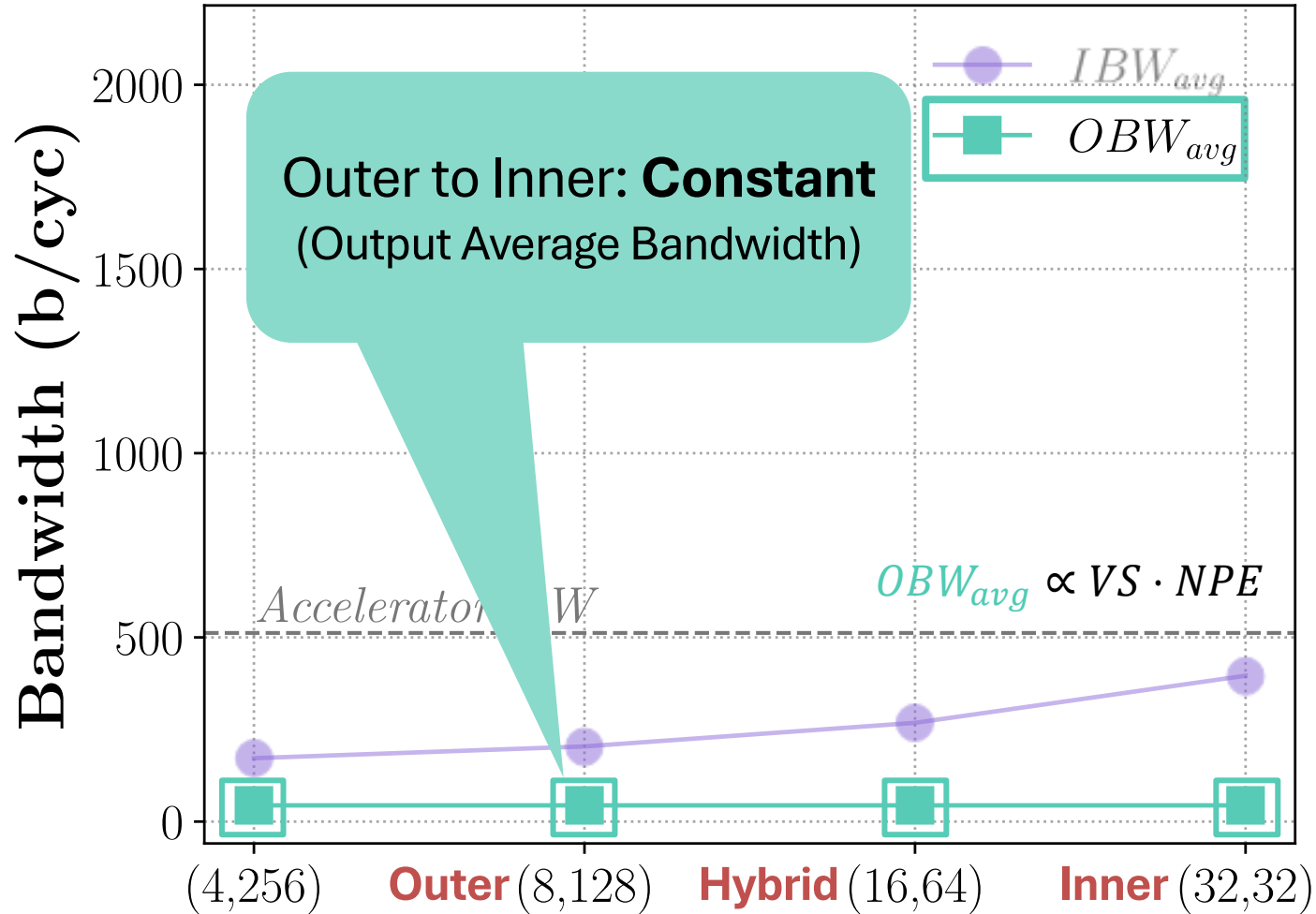
# Bandwidth: Inner vs Outer



(Inner, Outer) sizes of MXCore: Inner x Outer = 1024 MACs

**Reuse fixed at 64**

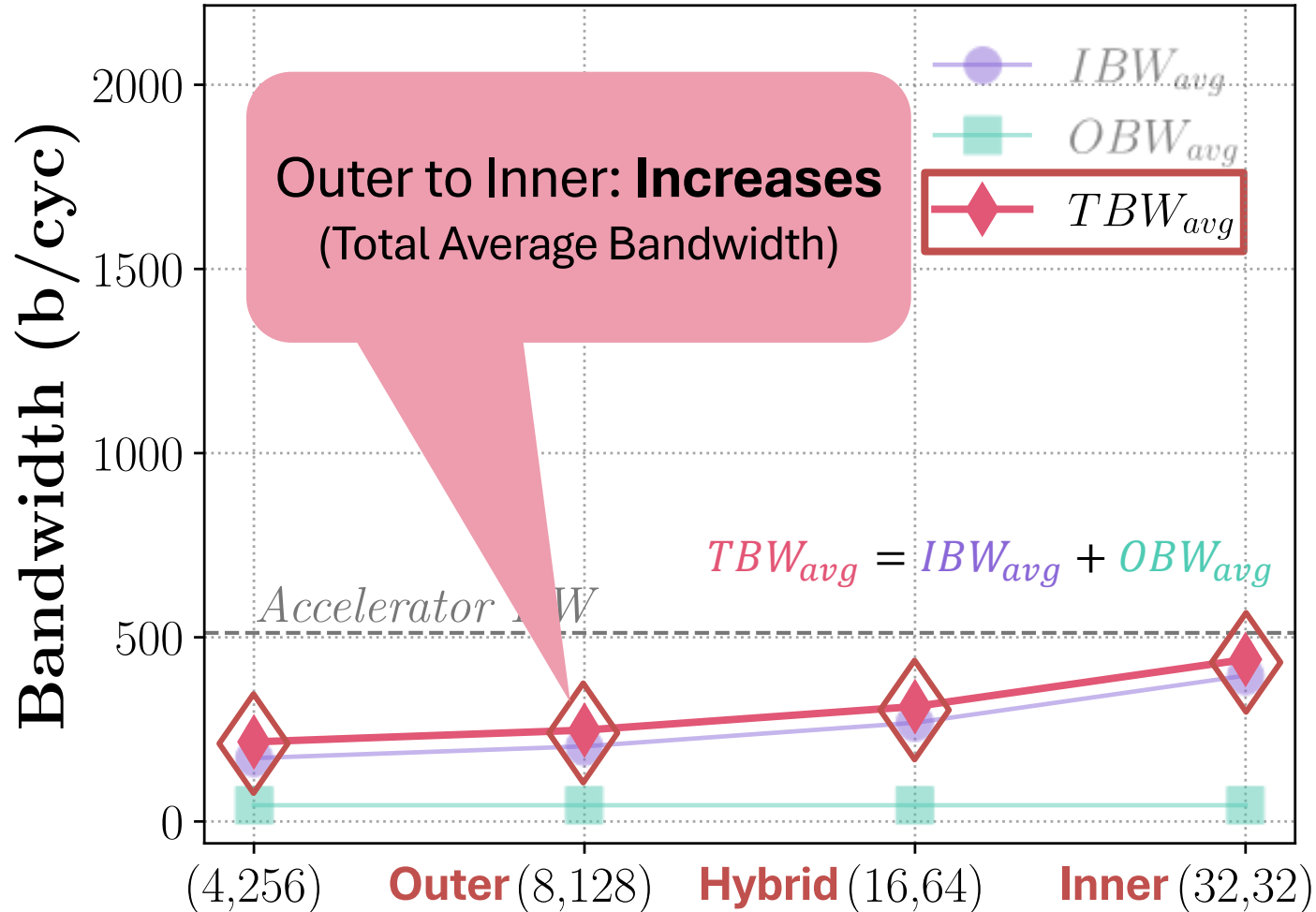
# Bandwidth: Inner vs Outer



(Inner, Outer) sizes of MXCore: Inner x Outer = 1024 MACs

Reuse fixed at 64

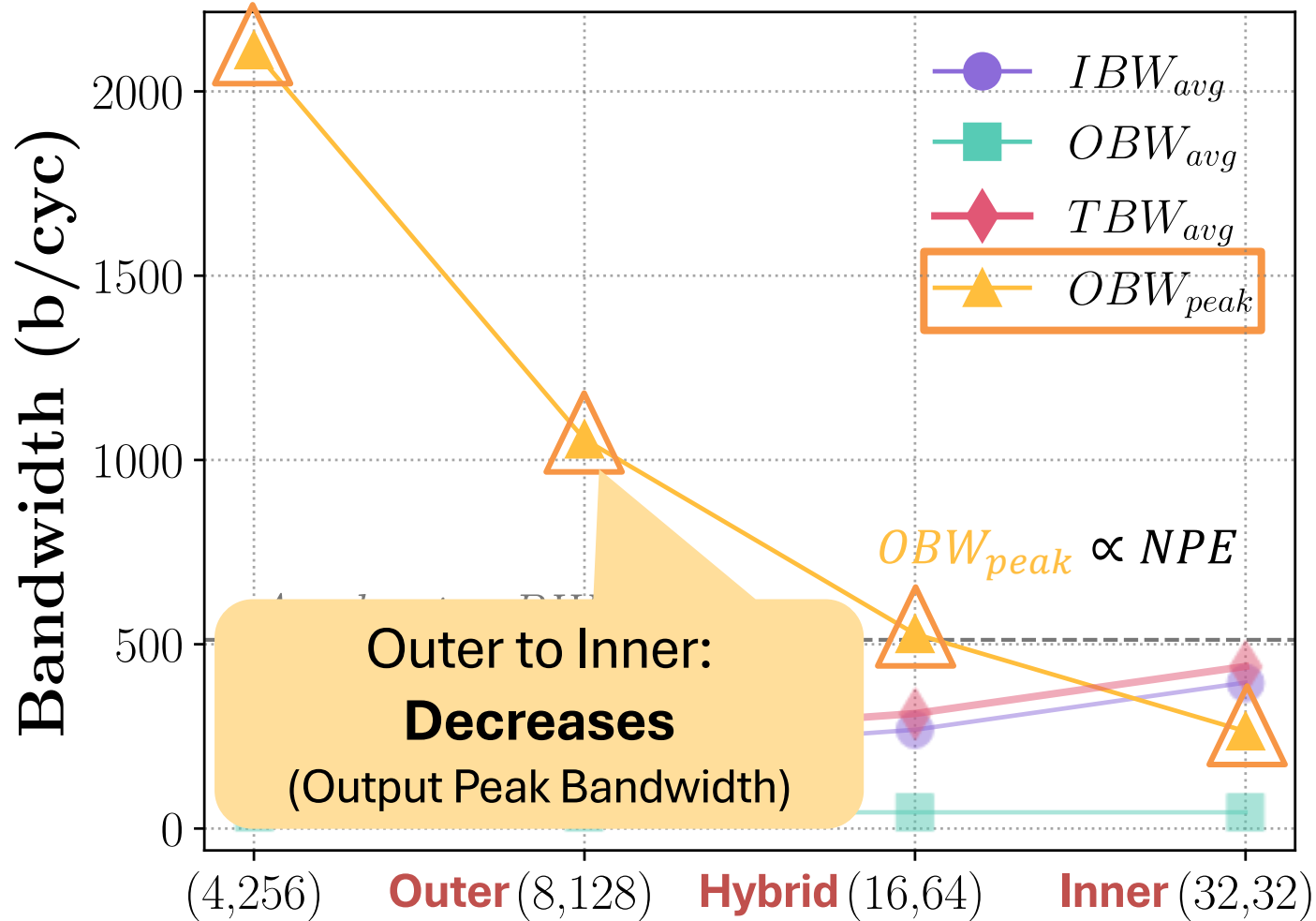
# Bandwidth: Inner vs Outer



(Inner, Outer) sizes of MXCore: Inner x Outer = 1024 MACs

Reuse fixed at 64

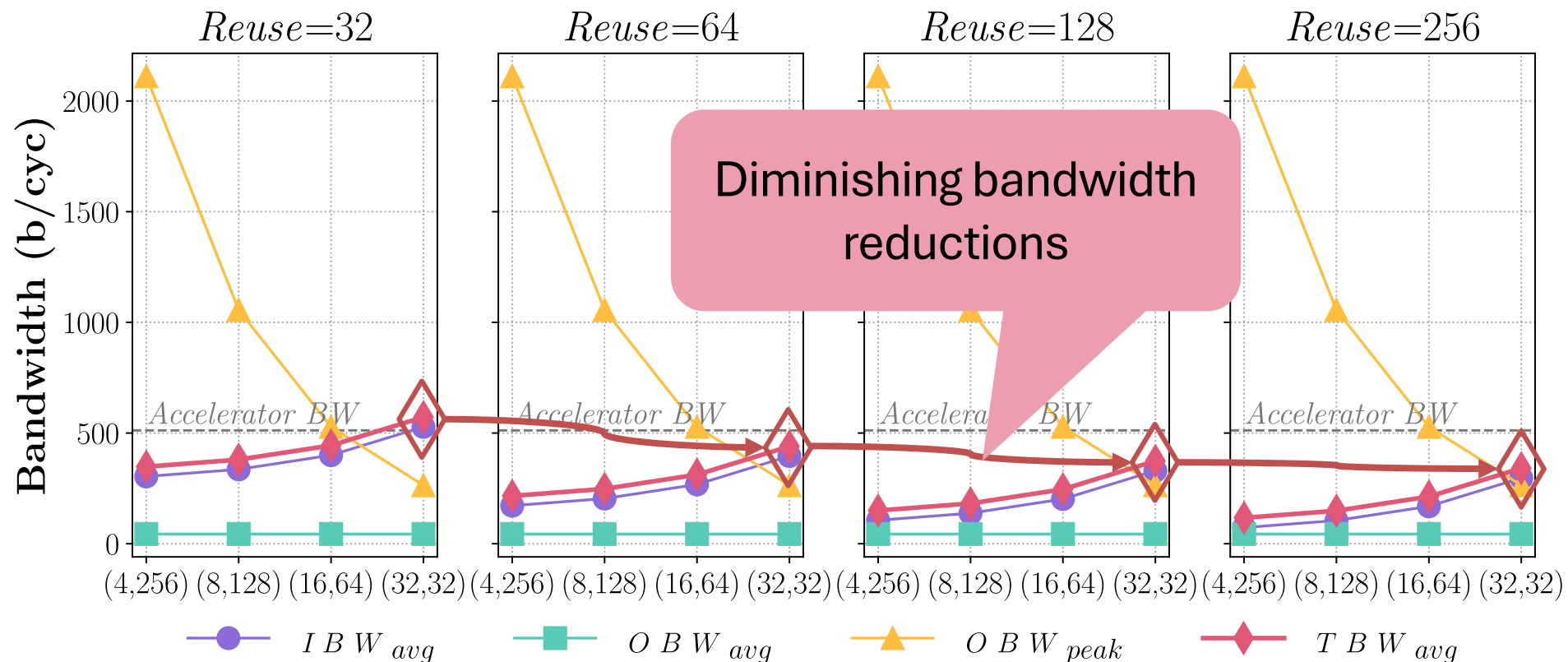
# Bandwidth: Inner vs Outer



(Inner, Outer) sizes of MXCore: Inner x Outer = 1024 MACs

Reuse fixed at 64

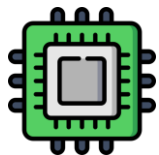
# Bandwidth: Across Reuse



**(Inner, Outer) sizes of MXCore: Inner x Outer = 1024 MACs**  
**Across different Reuse**



# Physical Implementation of MXCore



GlobalFoundries' **12nm** FinFET technology



Target Frequency **0.8 GHz** @ (SS/0.72V/125°C)



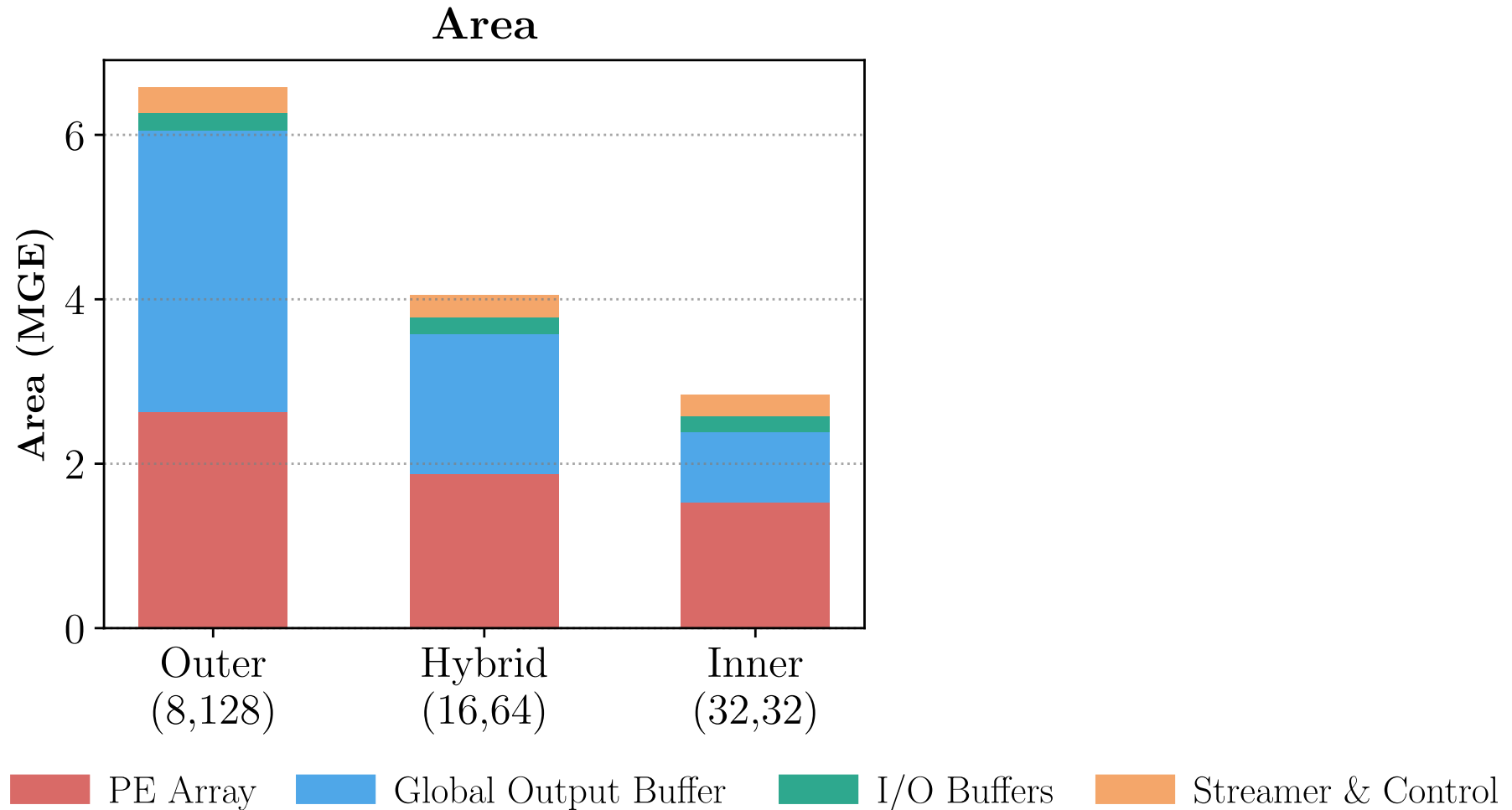
Outer to Inner: **0.98 – 1.07 GHz** @ (TT/0.8V/25°C)



Post-layout simulations at **0.95 GHz** on *five* MXFP8-quantized samples extracted from *DeiT-Tiny*

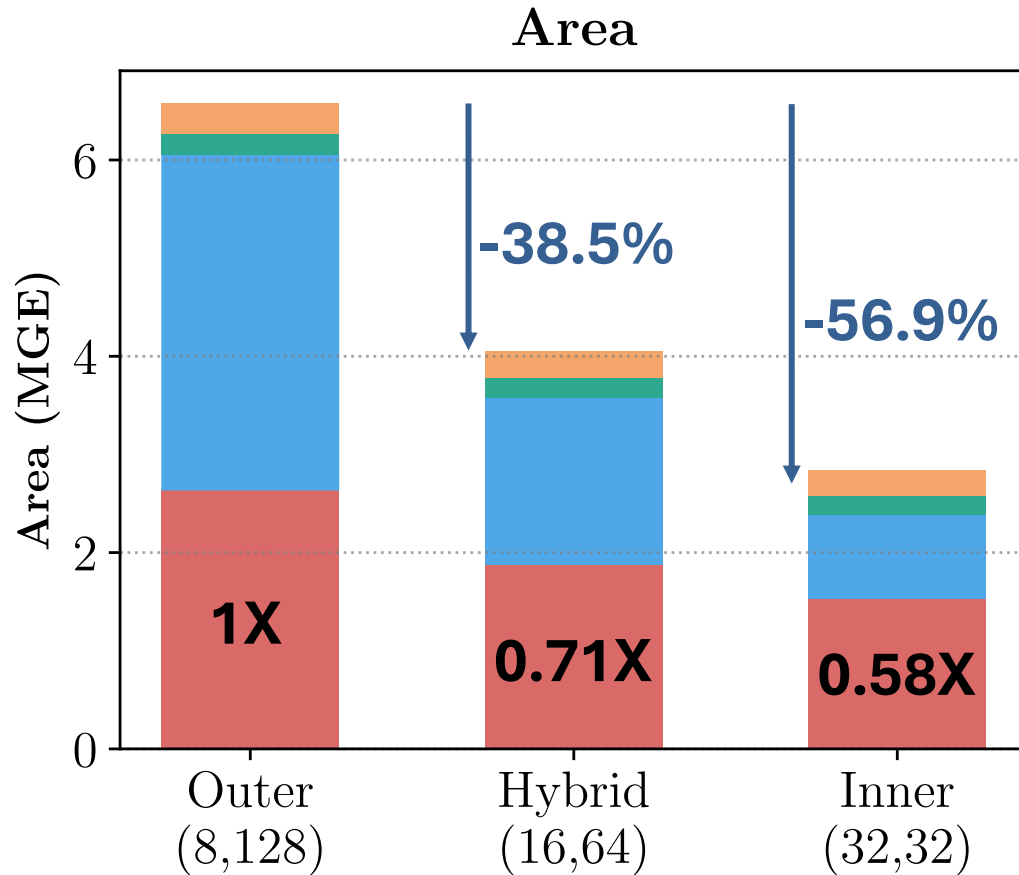


# Area Breakdown: Inner to Outer



**Area Breakdown from Inner to Outer at a fixed *Reuse* of 64**

# Area Breakdown: Inner to Outer



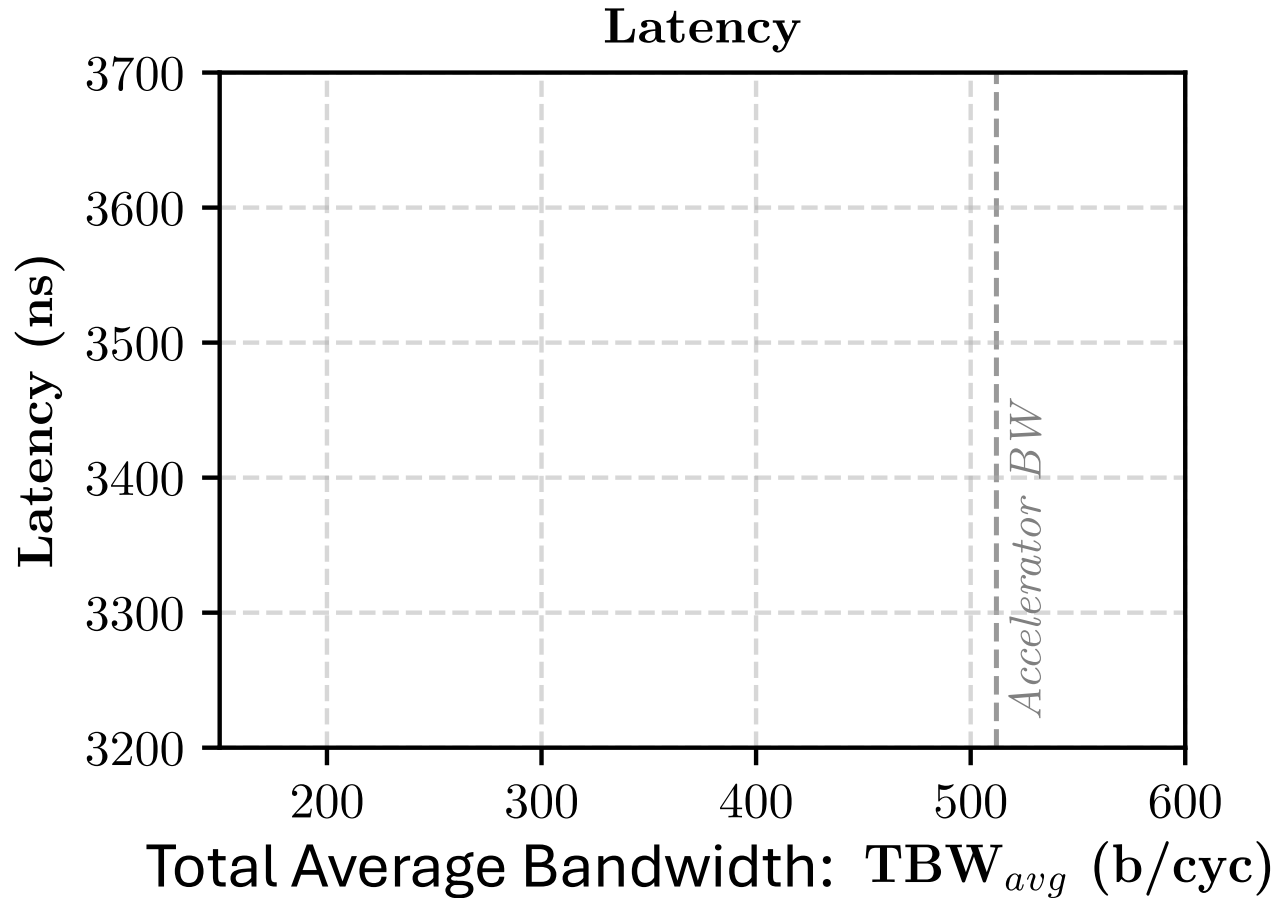
## Observations

- *Inner* arrays favour area.
  - Increasing *VS* only adds narrow multipliers.
- *Outer* arrays are area-inefficient.
  - Increasing *NPE* replicates logic.

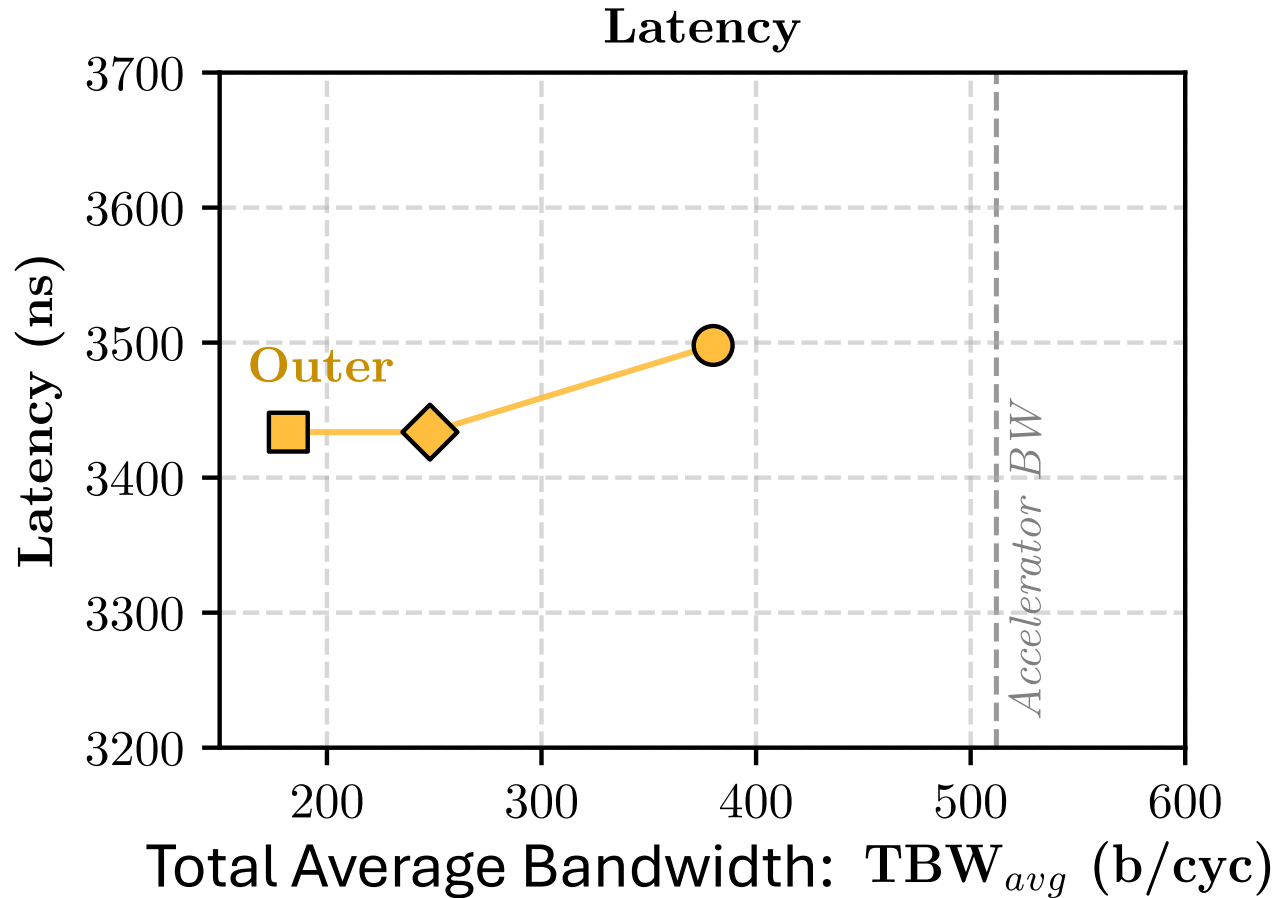
PE Array Global Output Buffer I/O Buffers Streamer & Control

Area Breakdown from Inner to Outer at a fixed *Reuse* of 64

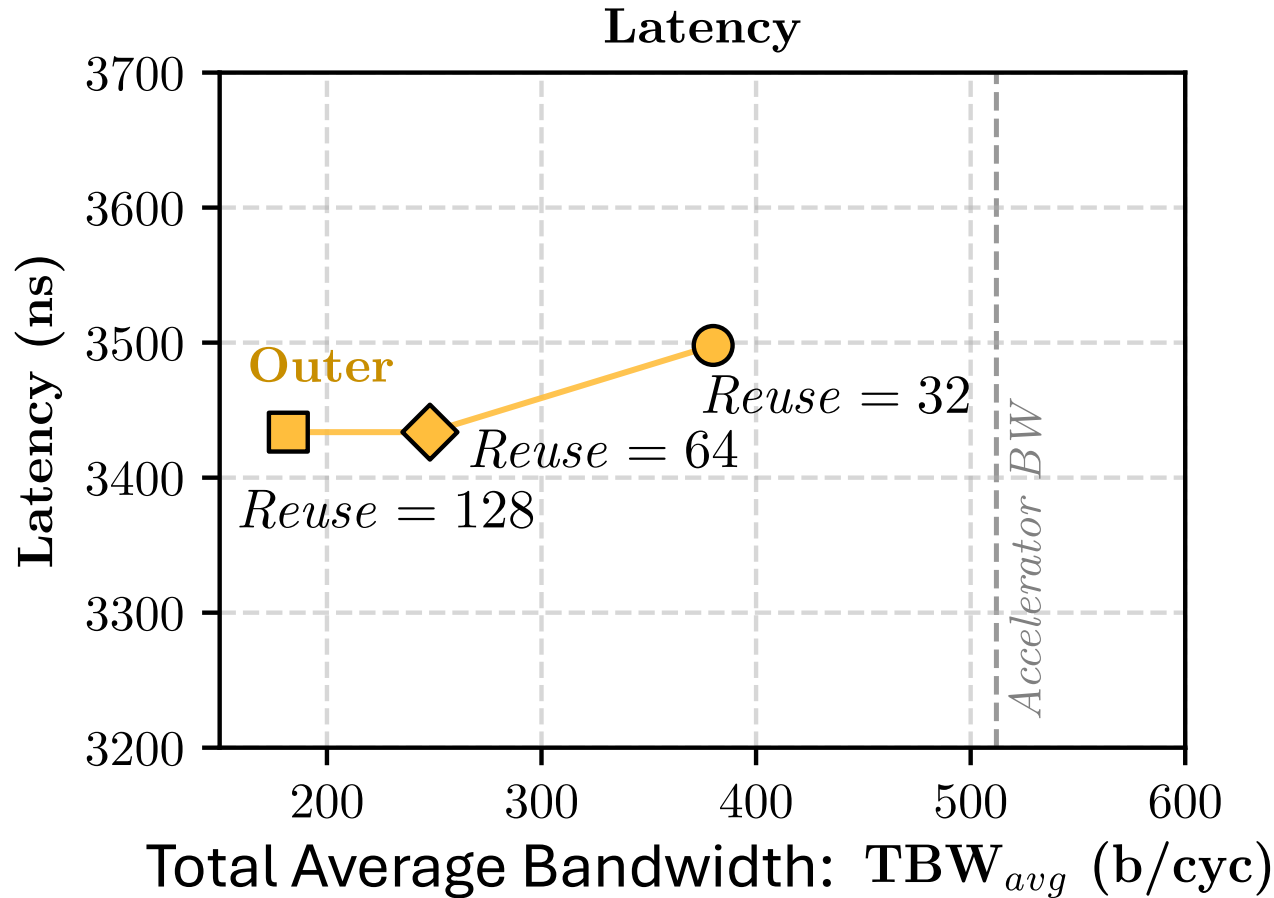
# Pareto-optimal Designs: Latency



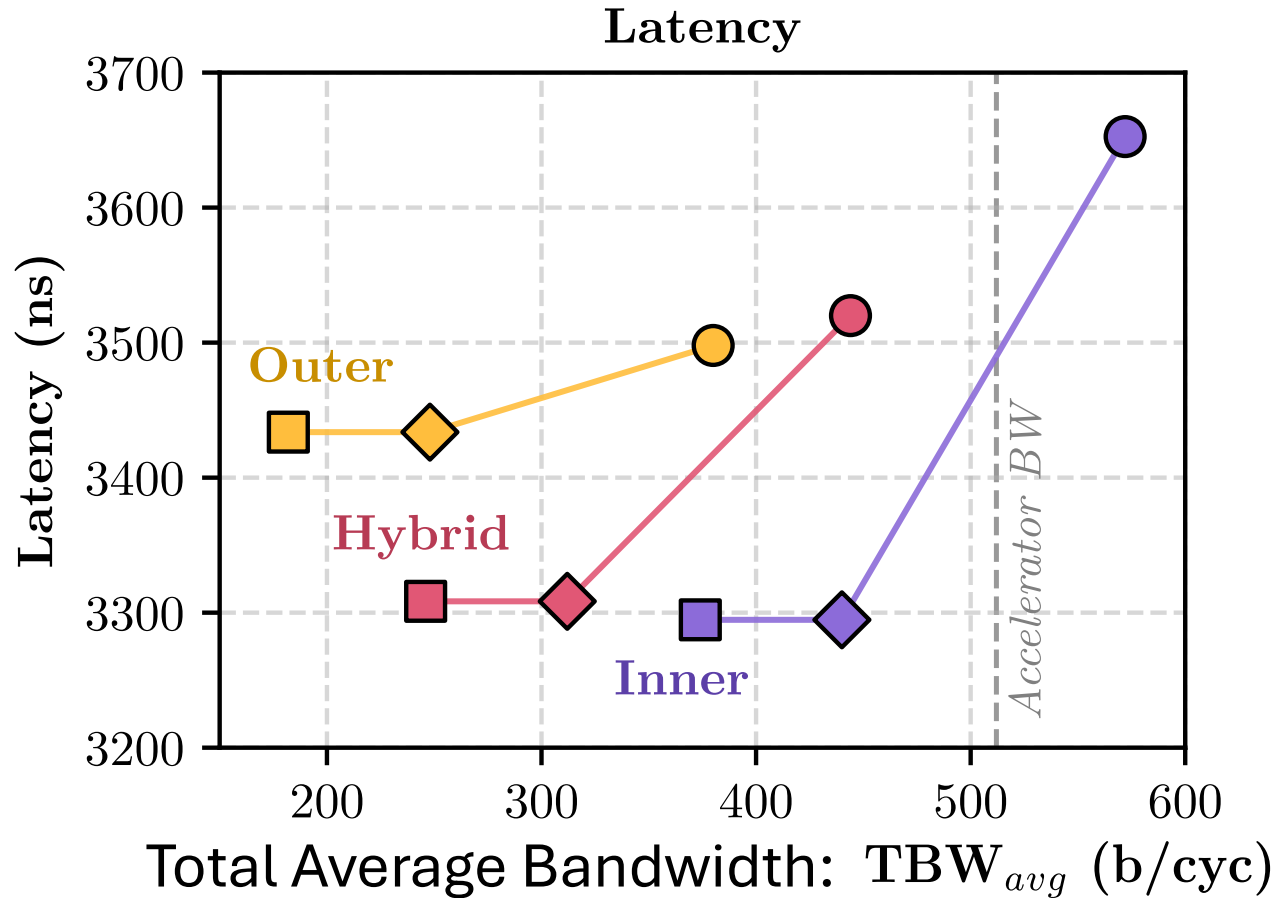
# Pareto-optimal Designs: Latency



# Pareto-optimal Designs: Latency



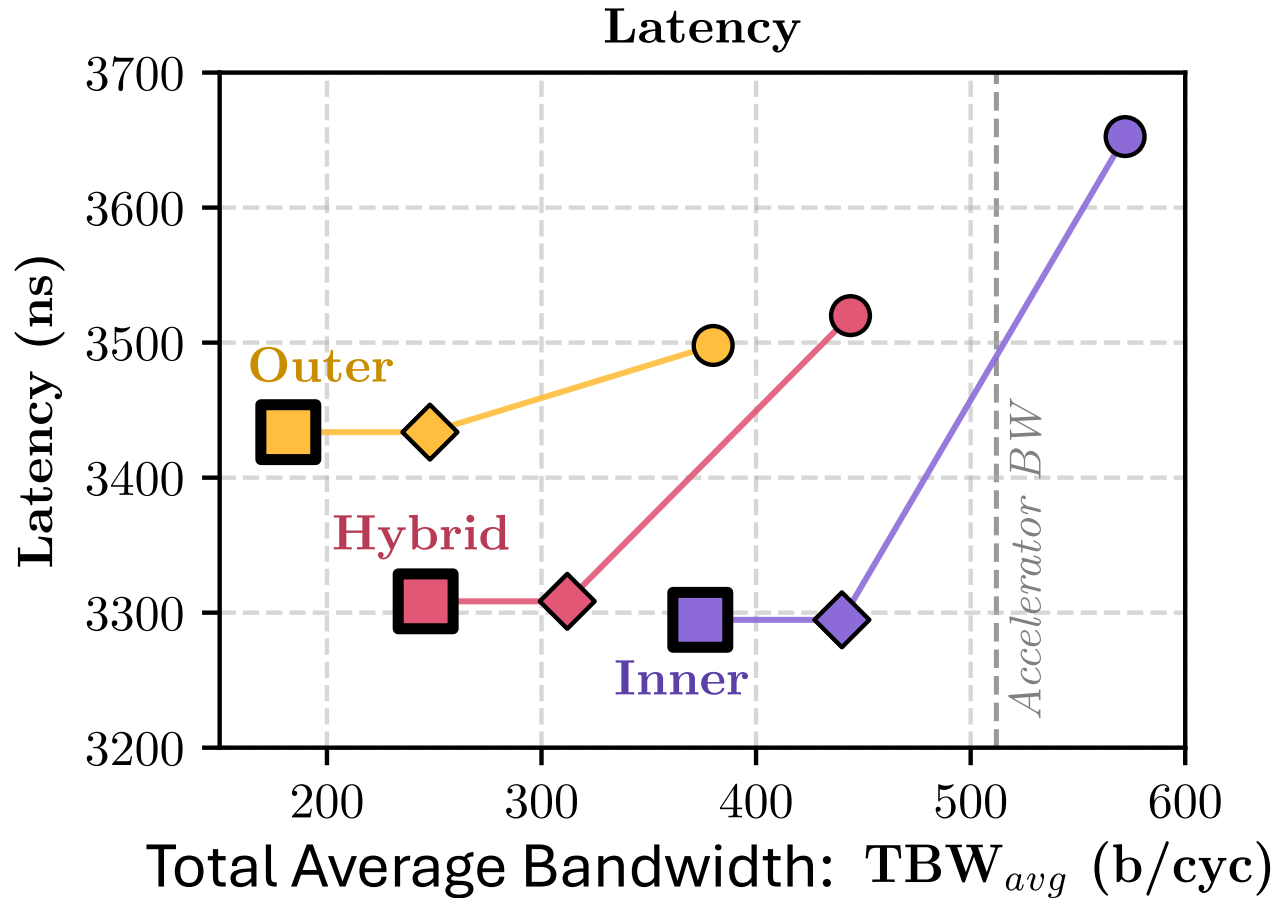
# Pareto-optimal Designs: Latency



—●— ( $VS = 8, NPE = 128$ )    —●— ( $VS = 16, NPE = 64$ )    —●— ( $VS = 32, NPE = 32$ )

●  $Reuse = 32$     ◆  $Reuse = 64$     ■  $Reuse = 128$

# Pareto-optimal Designs: Latency

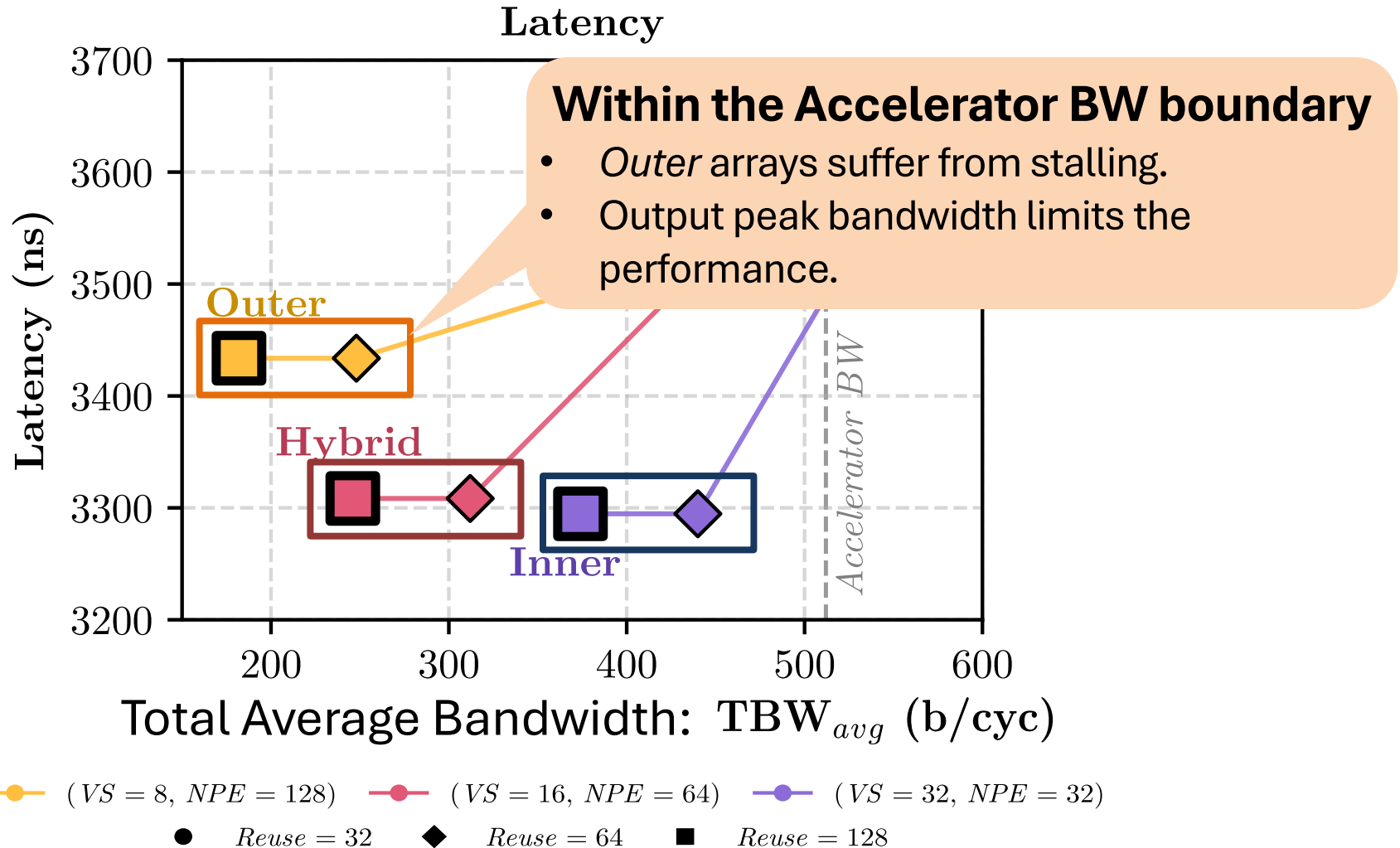


—●— ( $VS = 8, NPE = 128$ )    —●— ( $VS = 16, NPE = 64$ )    —●— ( $VS = 32, NPE = 32$ )

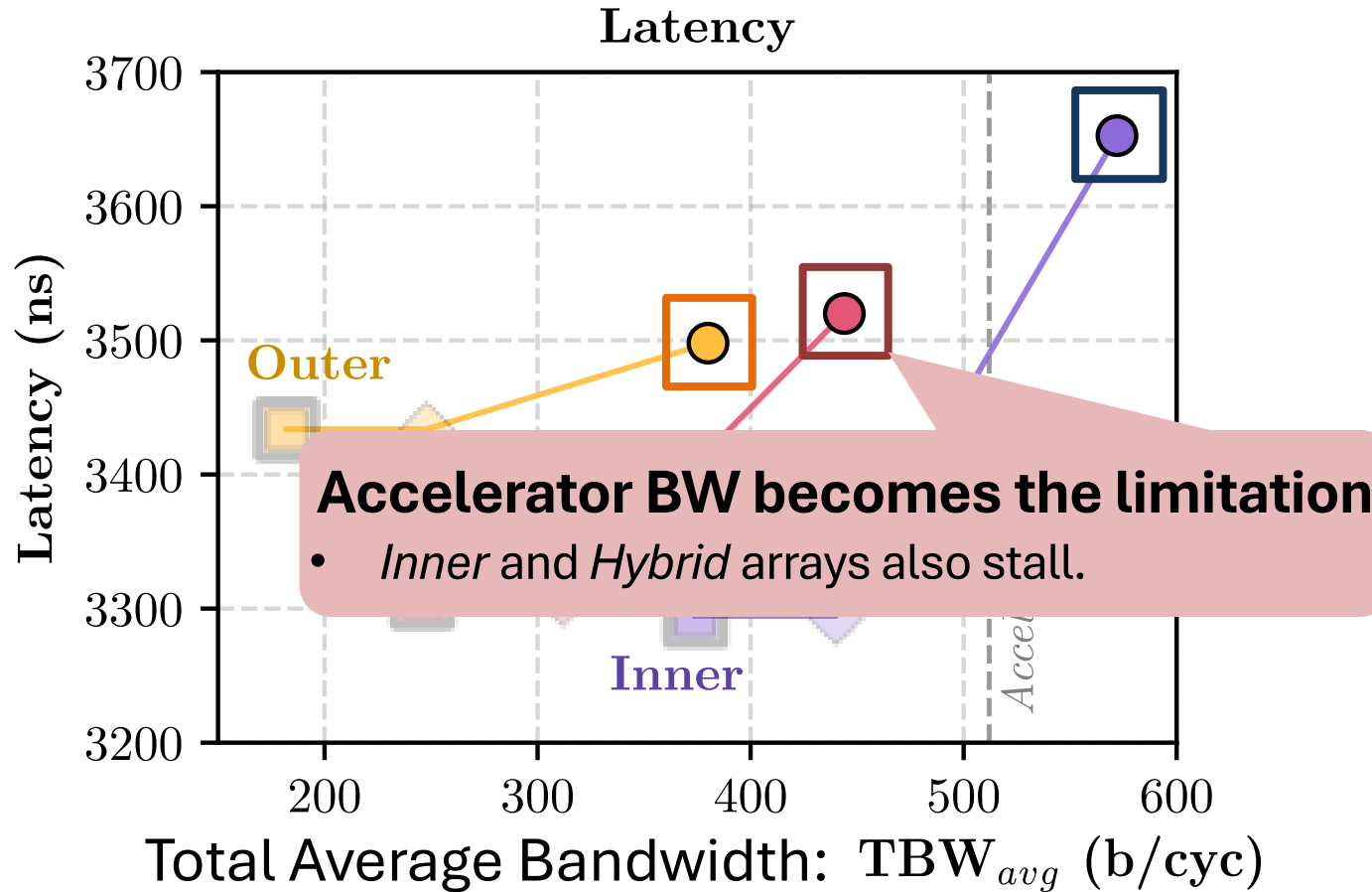
●  $Reuse = 32$     ◆  $Reuse = 64$     ■  $Reuse = 128$



# Pareto-optimal Designs: Latency



# Pareto-optimal Designs: Latency

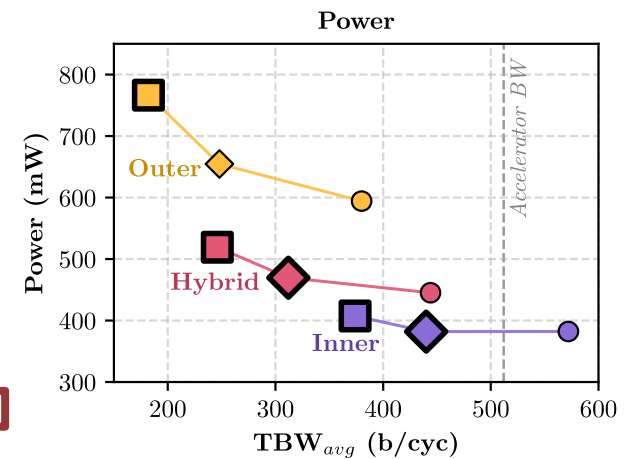
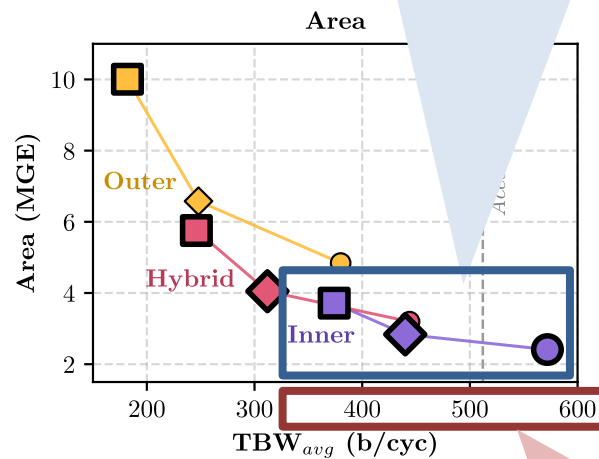
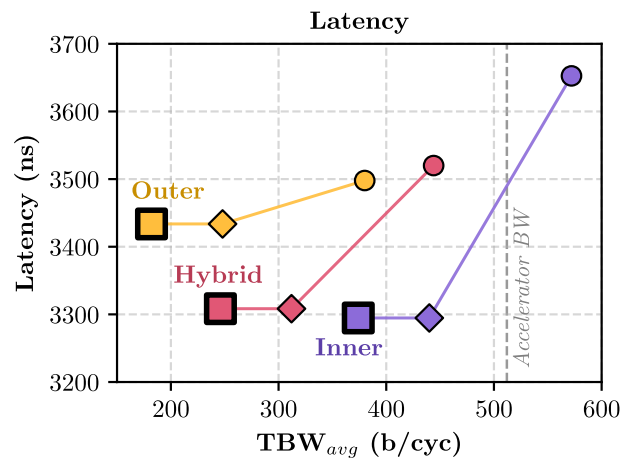


—●— ( $VS = 8, NPE = 128$ )    —●— ( $VS = 16, NPE = 64$ )    —●— ( $VS = 32, NPE = 32$ )

●  $Reuse = 32$     ◆  $Reuse = 64$     ■  $Reuse = 128$

# Pareto-optimal Designs: Trade-offs

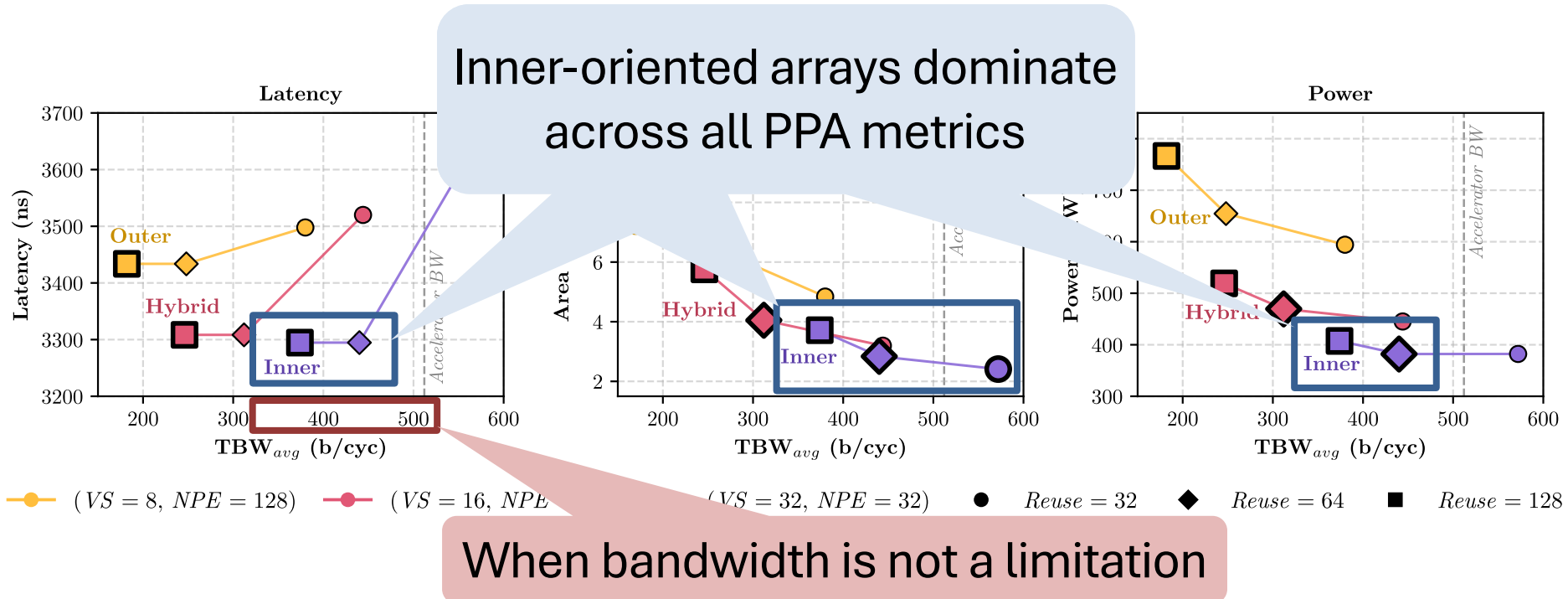
Inner-oriented arrays offer superior area efficiency



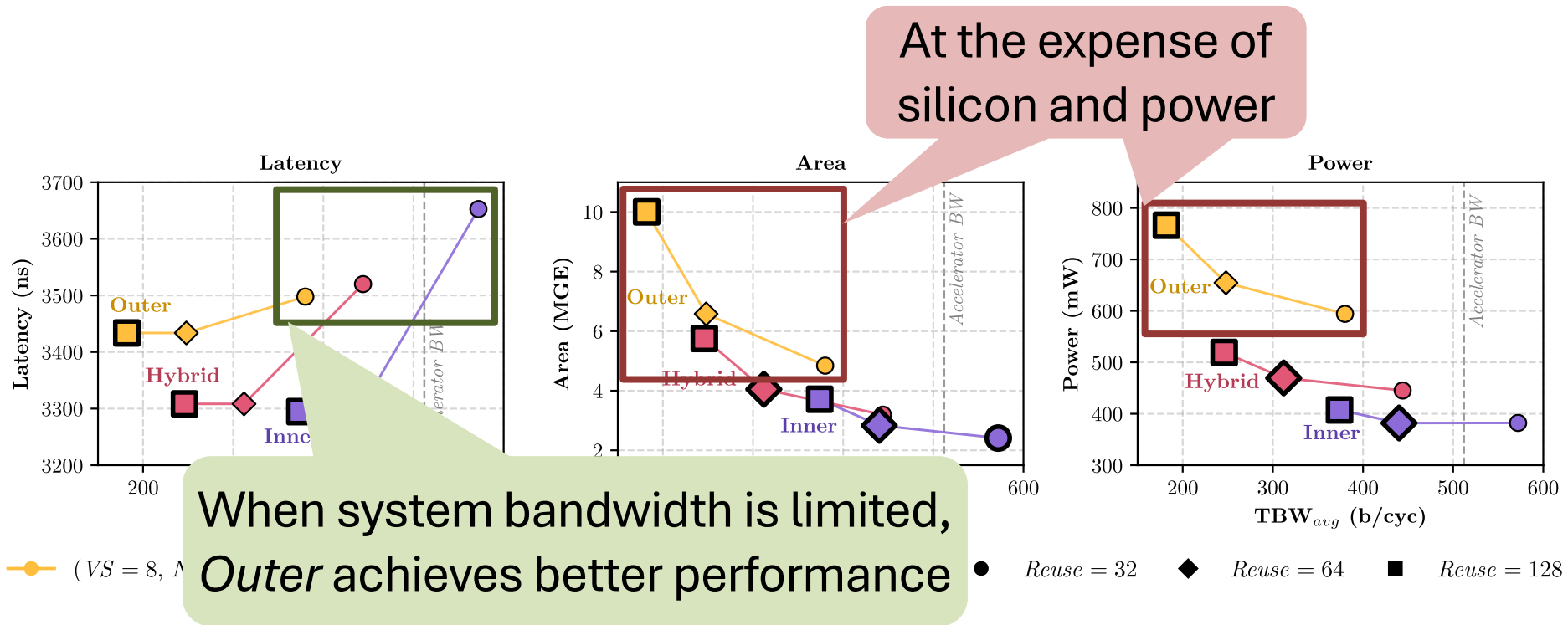
—●— ( $VS = 8, NPE = 128$ )   
 —●— ( $VS = 16, NPE = 64$ )   
 —●— ( $VS = 32, NPE = 32$ )   
 ● ( $VS = 32, NPE = 32$ )   
 ◆  $Reuse = 64$    
 ■  $Reuse = 128$

At the cost of bandwidth

# Pareto-optimal Designs: Trade-offs



# Pareto-optimal Designs: Trade-offs



# Results: State-of-the-Art Comparison

| Work                                 | Tech.<br>(nm) | Freq.<br>(MHz) | Precision      | Area Eff.<br>GFLOPS/mm <sup>2</sup> | Energy Eff.<br>GFLOPS/W  |
|--------------------------------------|---------------|----------------|----------------|-------------------------------------|--------------------------|
| Gemmini                              | 22            | 210            | INT8           | 59.4                                | 106.1                    |
| OpenGEMM                             | 16            | 200            | INT2/4/8       | 329                                 | 4680                     |
| Cuyckens <i>et. al.</i>              | 16            | 400            | All MX Formats | 1469                                | 943-1042                 |
| VEGETA                               | 65            | 180            | MXFP8          | 182/140                             | 6460/5680                |
| <b>MXCore Inner<br/>(32, 32, 64)</b> | 12            | 950            | MXFP8          | 3190                                | 5093/4438<br>(8202/6742) |

Results shown for E5M2/E4M3; Values in parantheses correspond to PE array efficiency.

Cuyckens *et. al.* 2025. Efficient Precision-Scalable Hardware for Microscaling (MX) Processing in Robotics Learning. doi:10.48550/arXiv.2505.22404

Nine *et. al.* 2025. Optimizing Sparse/Dense VEGETA Accelerator Performance with Microscaling Quantization. ISCAS. doi:10.1109/ISCAS56072.2025.11043957



# Results: State-of-the-Art Comparison

| Work                         | Tech.<br>(nm)                                | Freq.<br>(MHz) | Precision | Area Eff.<br>GFLOPS/mm <sup>2</sup> | Energy Eff.<br>GFLOPS/W  |
|------------------------------|----------------------------------------------|----------------|-----------|-------------------------------------|--------------------------|
| Gemmini                      | 22                                           | 210            | INT8      | 48X                                 | 106.1                    |
| OpenGEMM                     | 16                                           | 200            | INT2/4/8  |                                     | 4680                     |
| Cuyckens <i>et. al.</i>      | No MX support<br>Widely used SOTA frameworks |                |           |                                     | 1469                     |
| VEGETA                       |                                              |                |           | 182/140                             | 182/140                  |
| MXCore Inner<br>(32, 32, 64) | 12                                           | 950            | MXFP8     | 3190                                | 5093/4438<br>(8202/6742) |

Results shown for E5M2/E4M3; Values in parantheses correspond to PE array efficiency.

Cuyckens *et. al.* 2025. Efficient Precision-Scalable Hardware for Microscaling (MX) Processing in Robotics Learning. doi:10.48550/arXiv.2505.22404

Nine *et. al.* 2025. Optimizing Sparse/Dense VEGETA Accelerator Performance with Microscaling Quantization. ISCAS. doi:10.1109/ISCAS56072.2025.11043957

# Results: State-of-the-Art Comparison

| Work                                 | Tech.<br>(nm) | Freq.<br>(MHz) | Precision      | Area Eff.<br>GFLOPS/mm <sup>2</sup> | Energy Eff.<br>GFLOPS/W  |
|--------------------------------------|---------------|----------------|----------------|-------------------------------------|--------------------------|
| Cuyckens <i>et. al.</i>              | 16            | 400            | All MX Formats | 1469                                | 943-1042                 |
| VEGETA                               | 65            | 180            | M 2.2X         | 182/140                             | 6460/5680                |
| <b>MXCore Inner<br/>(32, 32, 64)</b> | 12            | 950            | MXFP8          | 3190                                | 5093/4438<br>(8202/6742) |

Results shown for E5M2/E4M3; Values in parantheses correct efficiency.

7.2 – 7.9X

Cuyckens *et. al.* 2025. Efficient Precision-Scalable Hardware for Microscaling (MX) Processing in Robotics Learning. doi:10.48550/arXiv.2505.22404

Nine *et. al.* 2025. Optimizing Sparse/Dense VEGETA Accelerator Performance with Microscaling Quantization. ISCAS. doi:10.1109/ISCAS56072.2025.11043957



# Results: State-of-the-Art Comparison

| Work                                 | Tech.<br>(nm) | Freq.<br>(MHz) | Precision      | Area Eff.<br>GFLOPS/mm <sup>2</sup> | Energy Eff.<br>GFLOPS/W  |
|--------------------------------------|---------------|----------------|----------------|-------------------------------------|--------------------------|
| Cuyckens <i>et. al.</i>              | 16            | 400            | All MX Formats | 1469                                | 943-1042                 |
| VEGETA                               | 65            | 180            | MXFP8          | 182/140                             | 6460/5680                |
| <b>MXCore Inner<br/>(32, 32, 64)</b> | 12            | 950            | MXFP8          | 3190                                | 5093/4438<br>(8202/6742) |

**1.7 – 2.2X**

(After technology scaling)

Results shown for E5M2 correspond to PE array efficiency.

Cuyckens *et. al.* 2025. Efficient Precision-Scalable Hardware for Microscaling (MX) Processing in Robotics Learning. doi:10.48550/arXiv.2505.22404

Nine *et. al.* 2025. Optimizing Sparse/Dense VEGETA Accelerator Performance with Microscaling Quantization. ISCAS. doi:10.1109/ISCAS56072.2025.11043957

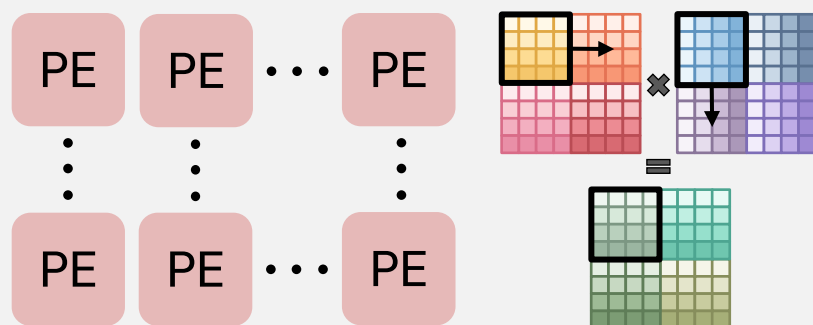


# Conclusion

- **MXCore: MX Matrix Multiplication Accelerator**
  - Fully parametric and configurable.
  - System-level integration capabilities.
  - FP32-to-MX quantizer for output bandwidth reduction.
  - Pareto-optimal designs and PPA comparisons.
- **Inner vs Outer Product Trade-offs**
  - Outer to Inner: **58.6%** higher bandwidth demand.
  - Outer to Inner: **2.1X** higher energy efficiency.
  - Outer to Inner: **3.7X** higher area efficiency.
- **State-of-the-Art Comparison**
  - **2.2X** higher area efficiency.
  - **1.2 – 7.9X** higher energy efficiency.



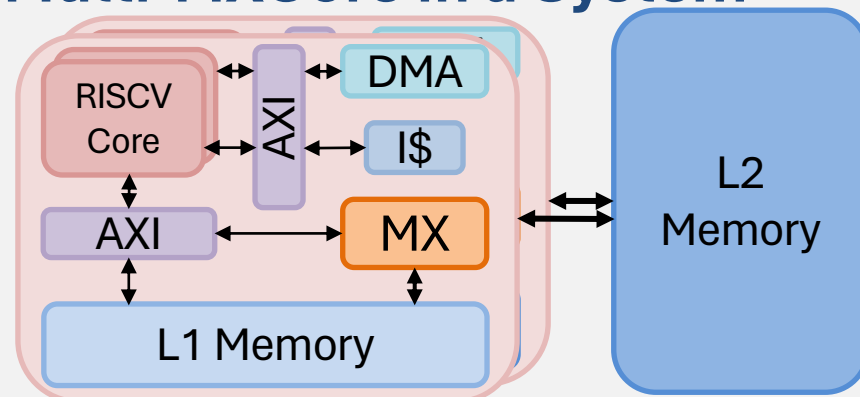
## Architectural Explorations



PE Organization

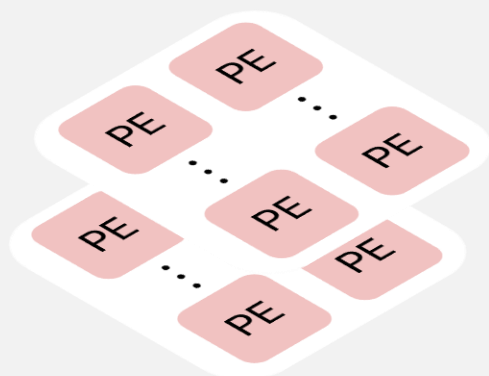
Dataflow

## Multi-MXCore in a System

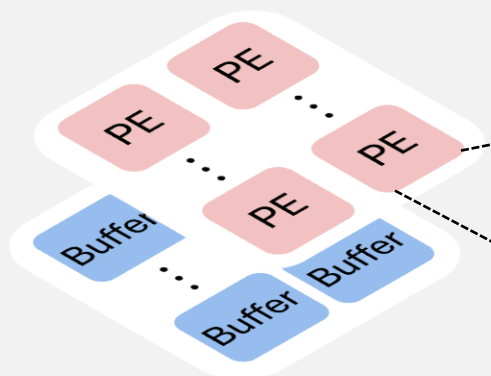


System-level Performance & Scalability

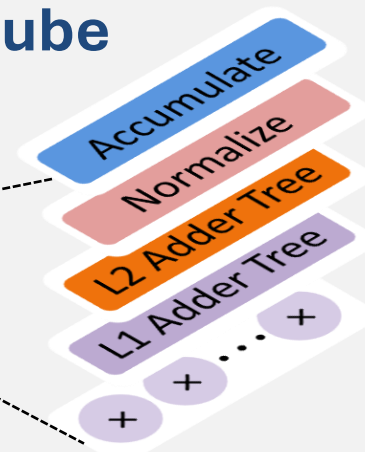
## Novel Design Paradigms: Going 3D with MXCube



Compute Scalability



Shorter Interconnects



PE Optimization

# THANK YOU

