
MX+: Pushing the Limits of Microscaling Formats for Efficient Large Language Model Serving

Jungi Lee

Junyong Park

Soohyun Cha

Jaehoon Cho

Jaewoong Sim

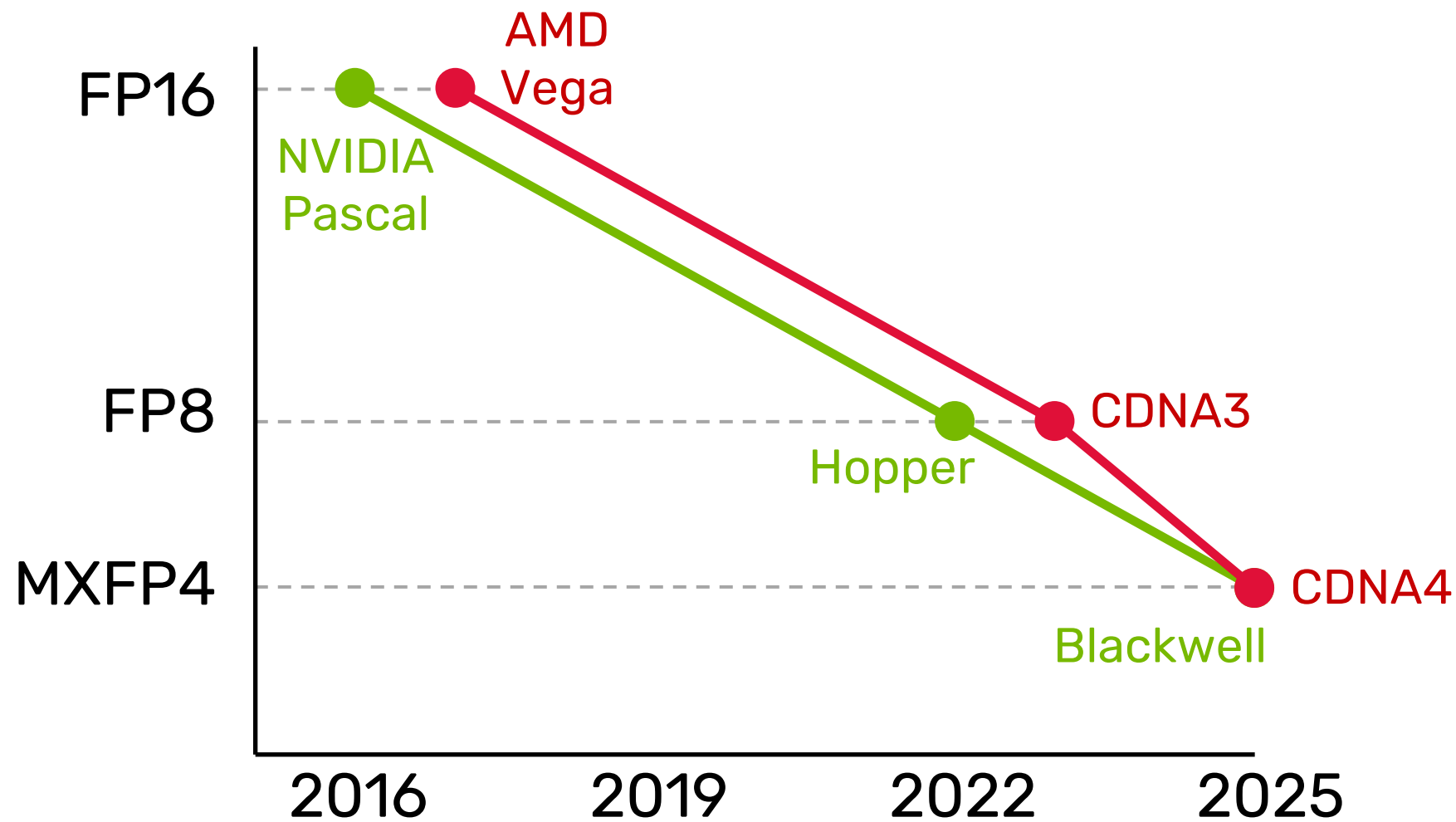
Seoul National University



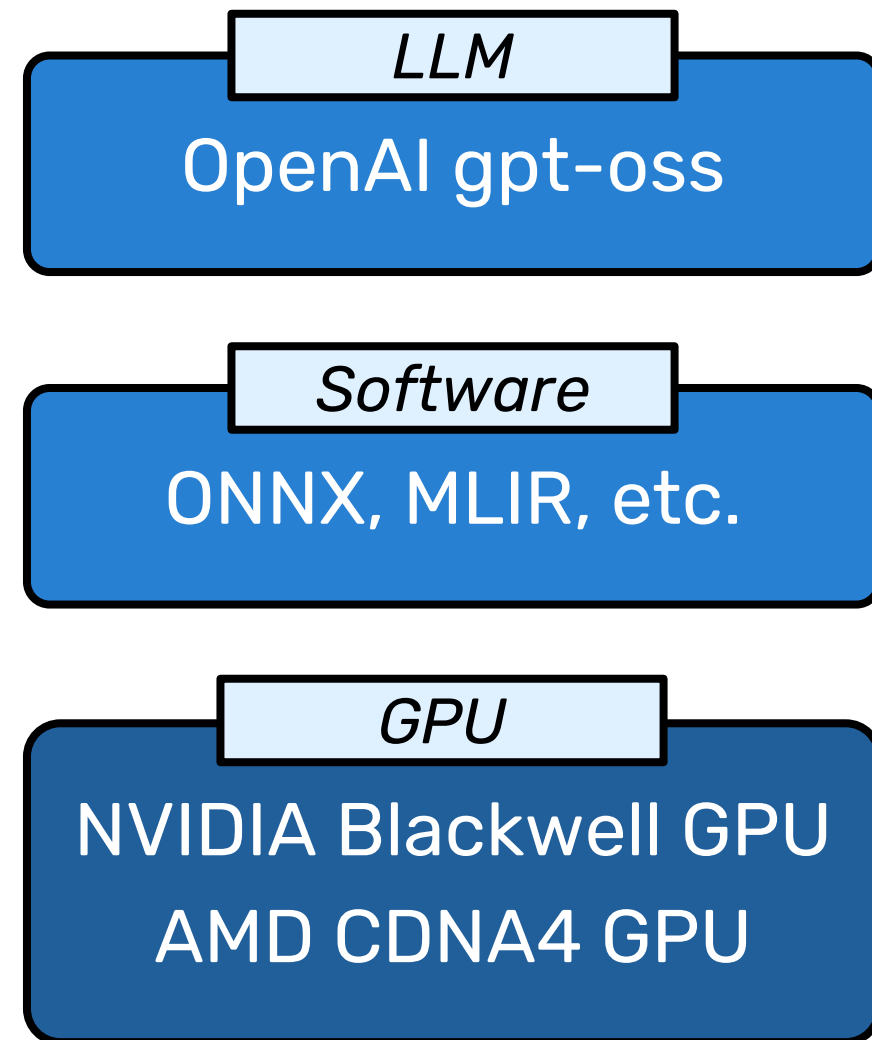
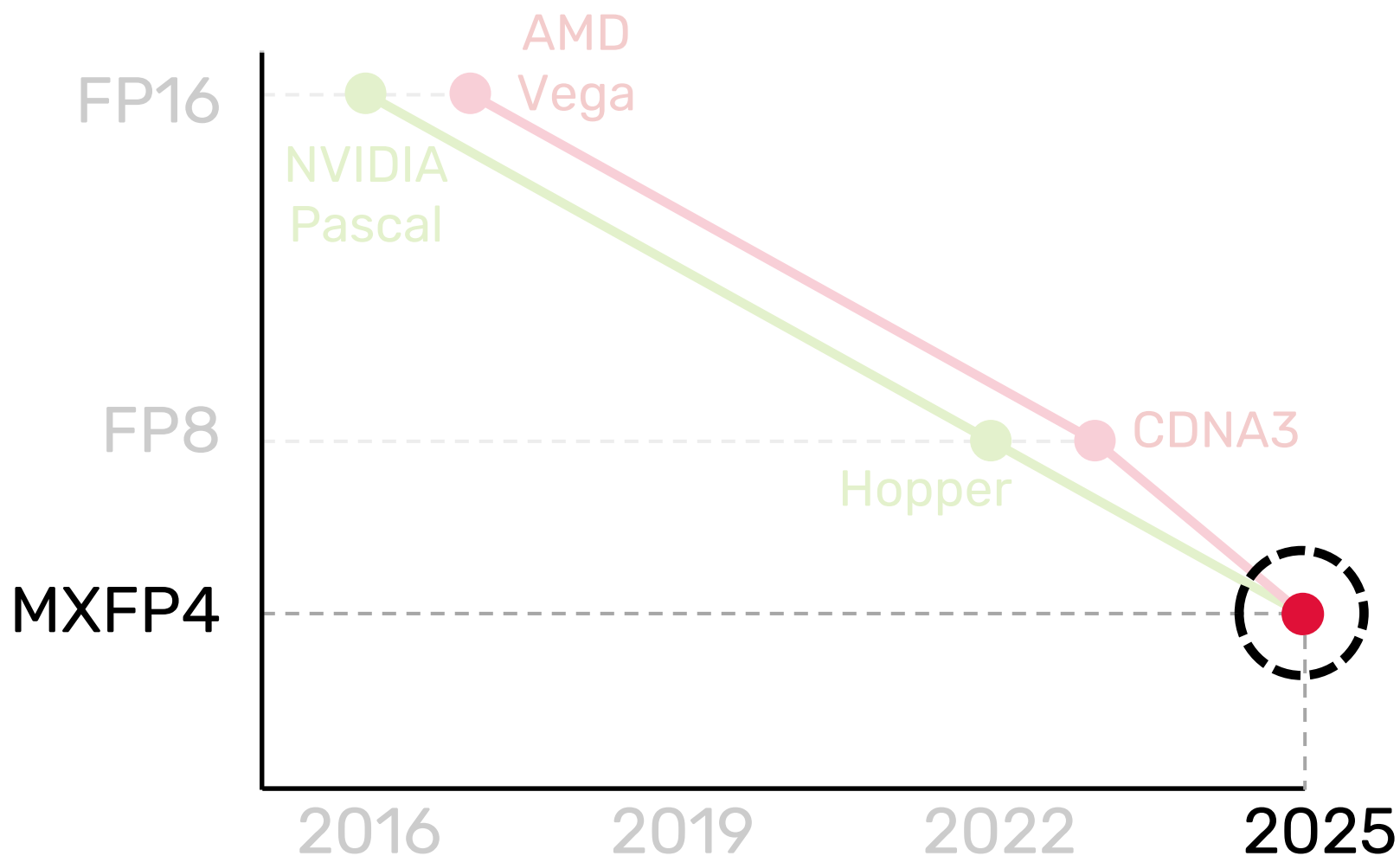
Outline

- **Introduction to MX Formats**
- **Motivation**
 - Implications for Low-bit LLM Serving
 - Analysis on Low-bit MX Formats
- **MX+: Cost-Effective and Non-intrusive Extension to MX Formats**
 - Key Insight
 - Software & Hardware Integration on GPUs
- **Evaluation**
- **Conclusion**

Increasing Support for MXFP4



Increasing Support for MXFP4



Microscaling (MX) Formats

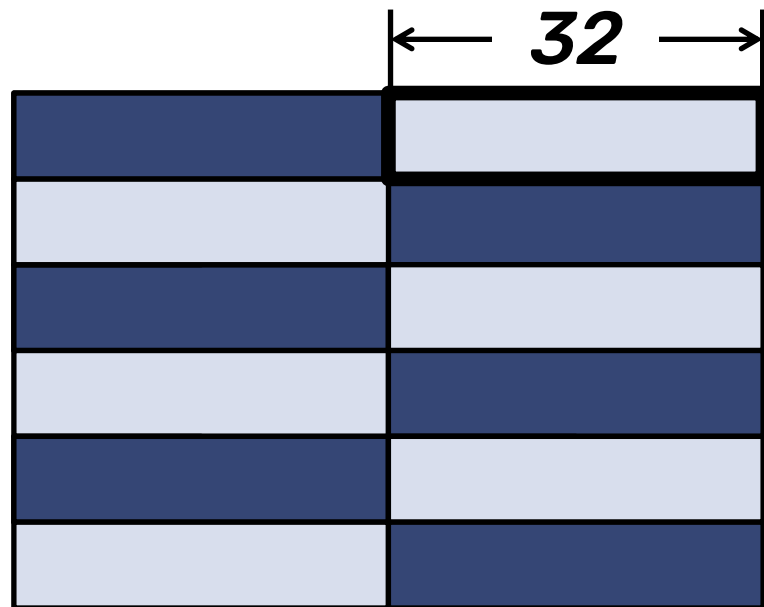
Open and Interoperable Data Formats

Standardized by the Industry *

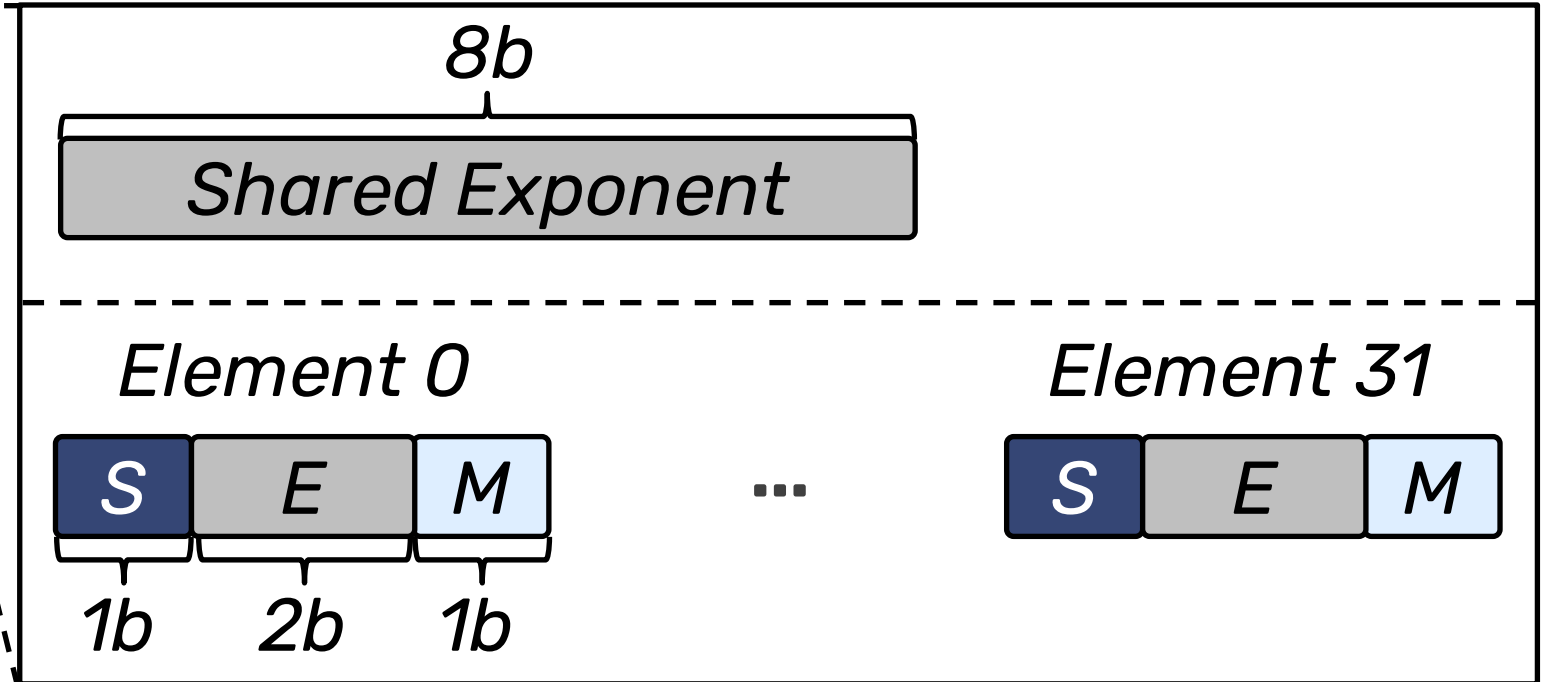
** AMD, Arm, Intel, Meta, Microsoft, NVIDIA, and Qualcomm*

Microscaling (MX) Formats

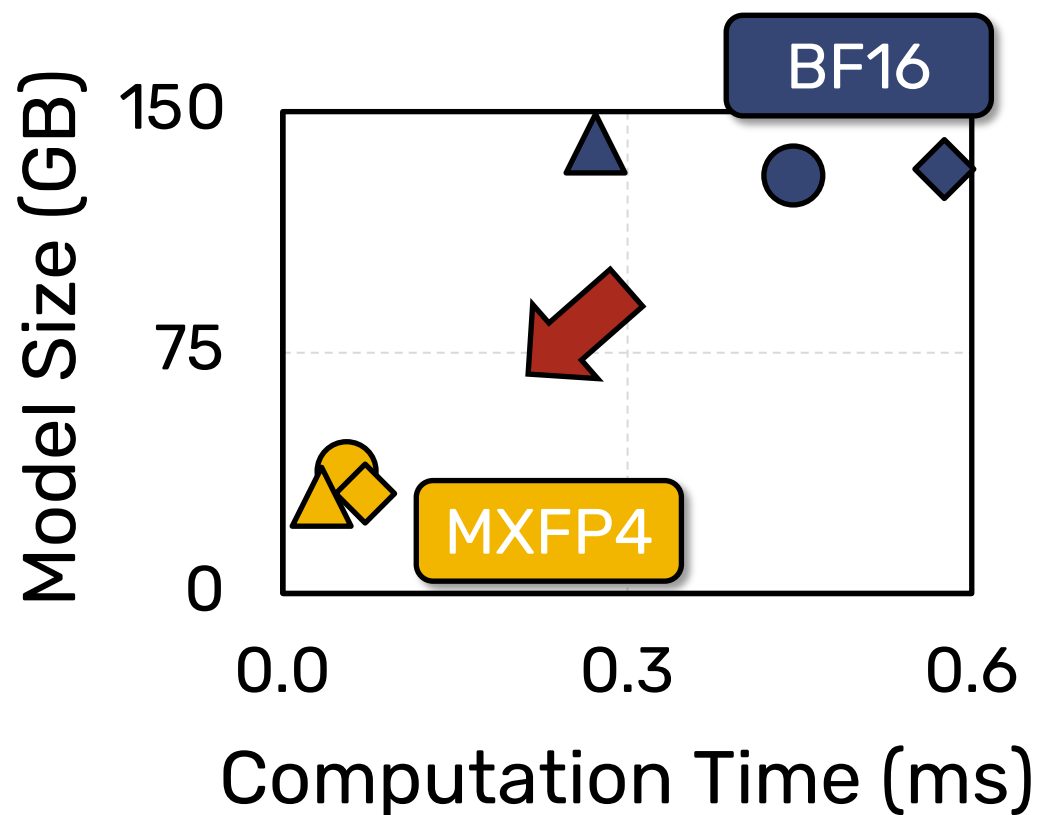
Tensor



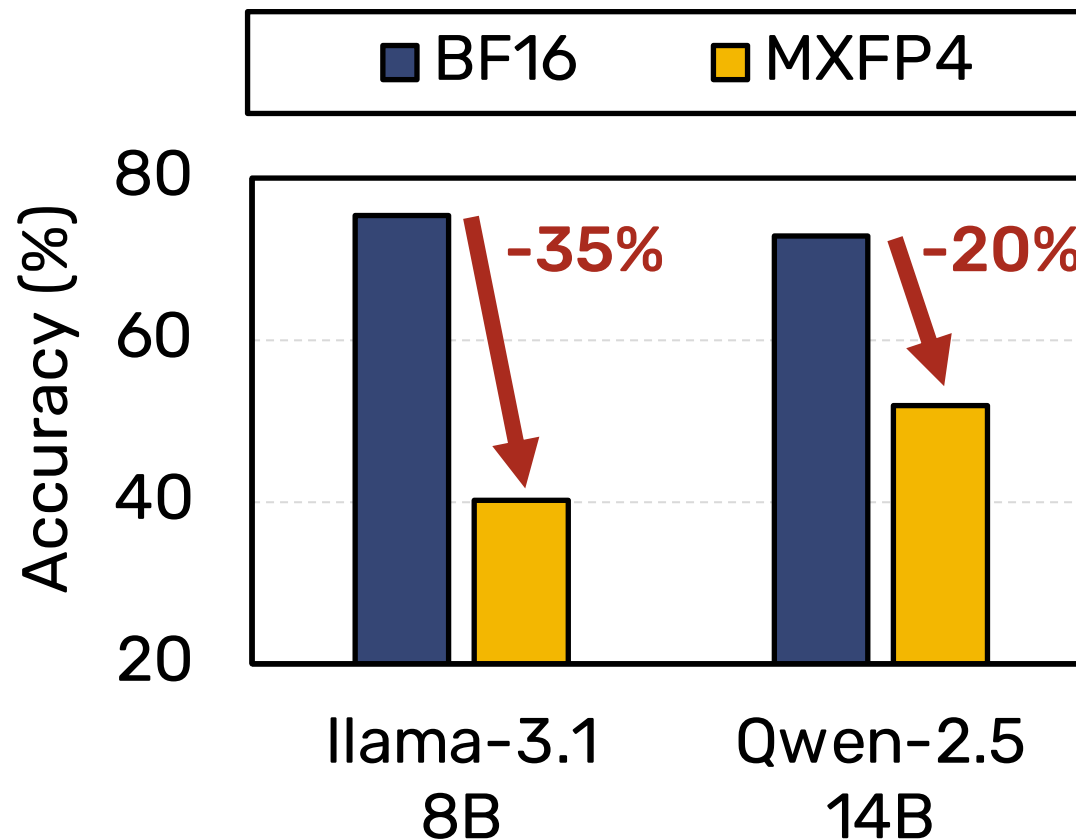
Data Layout



Implications on LLM Serving



Reduce Compute & Memory Overheads 😊



Low Accuracy 😞

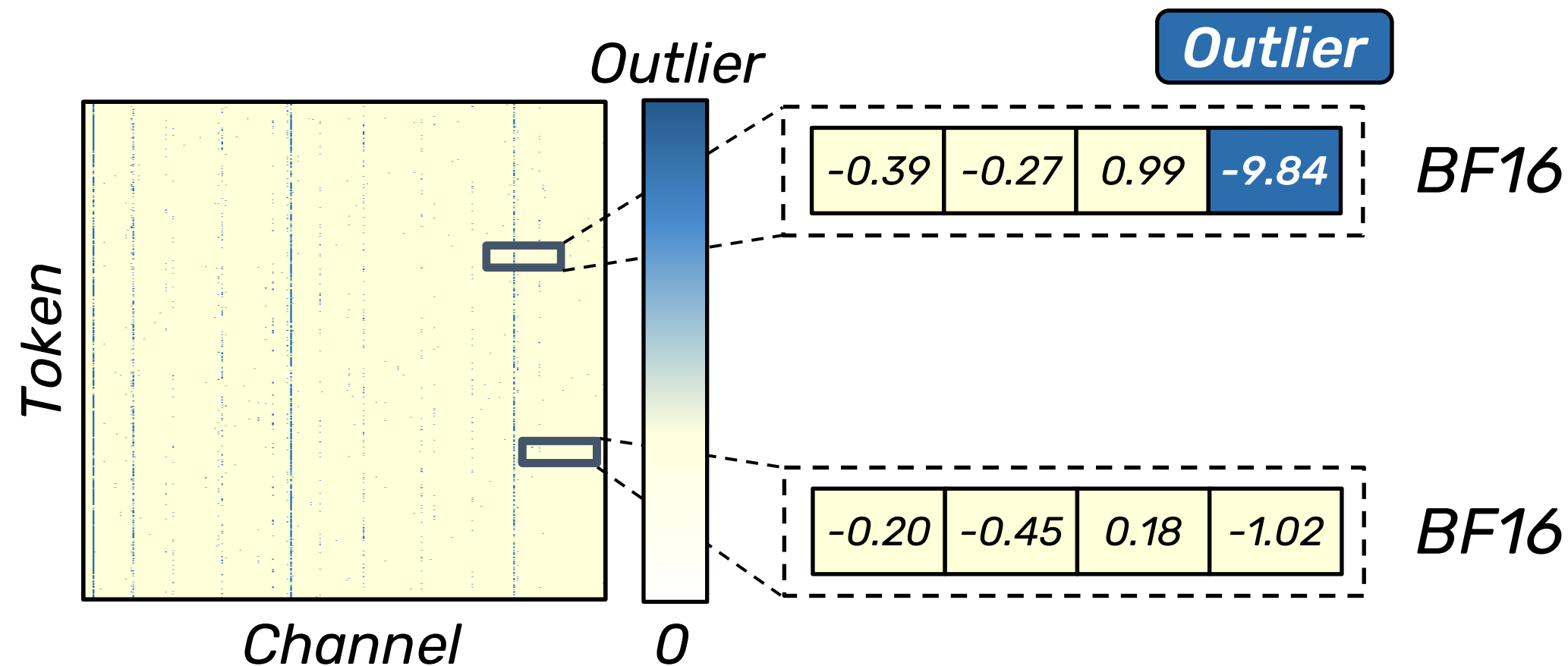
Implications on LLM Serving

How can we extend the MX formats to enable practical use in LLM serving?

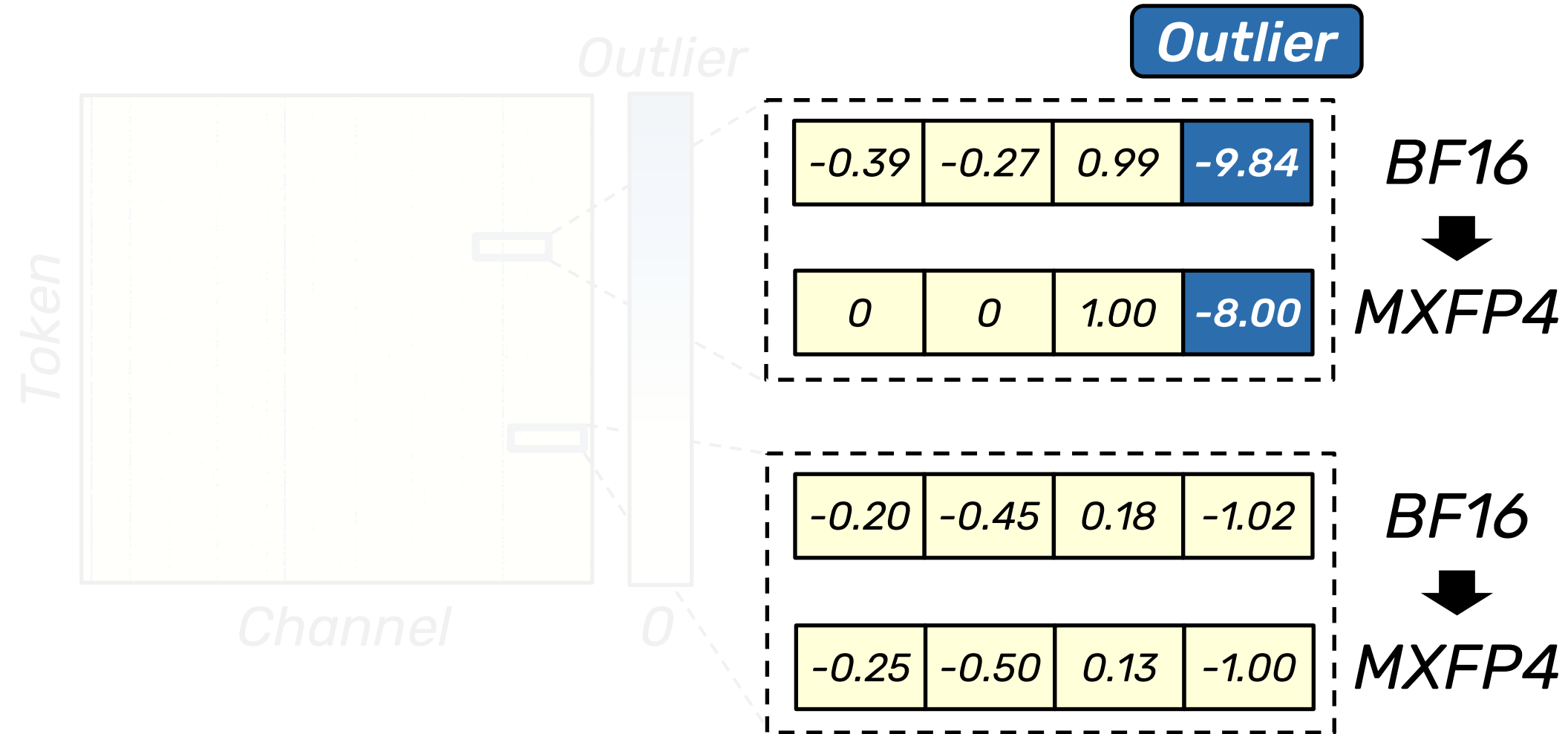
Reduce Compute & Memory Overheads 😊

Low Accuracy 😞

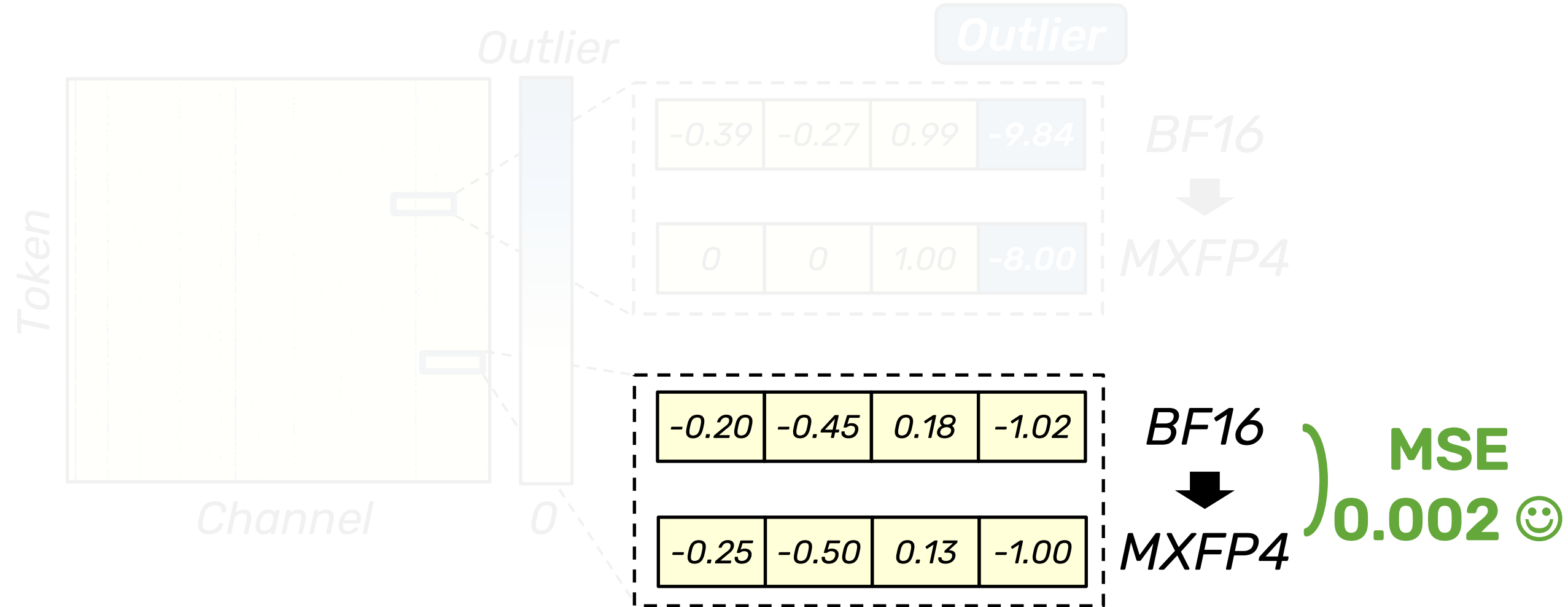
Analysis on Low-Bit MX Format



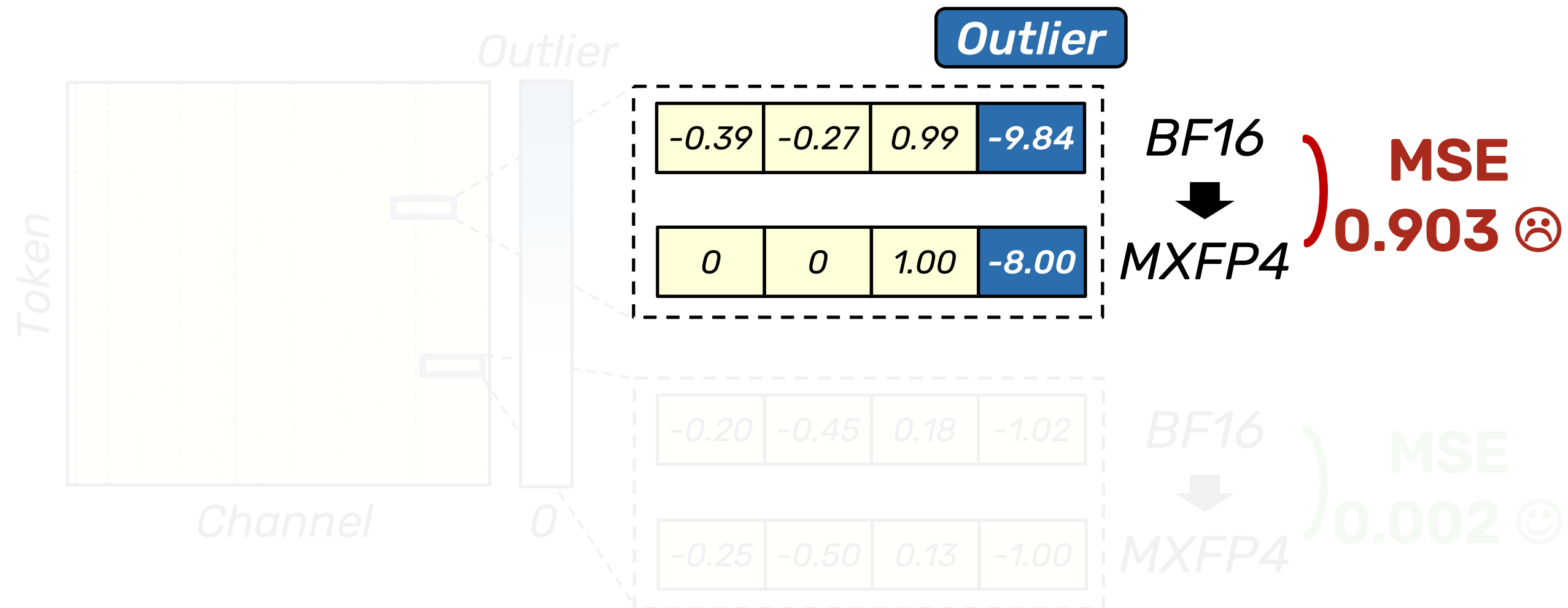
Analysis on Low-Bit MX Format



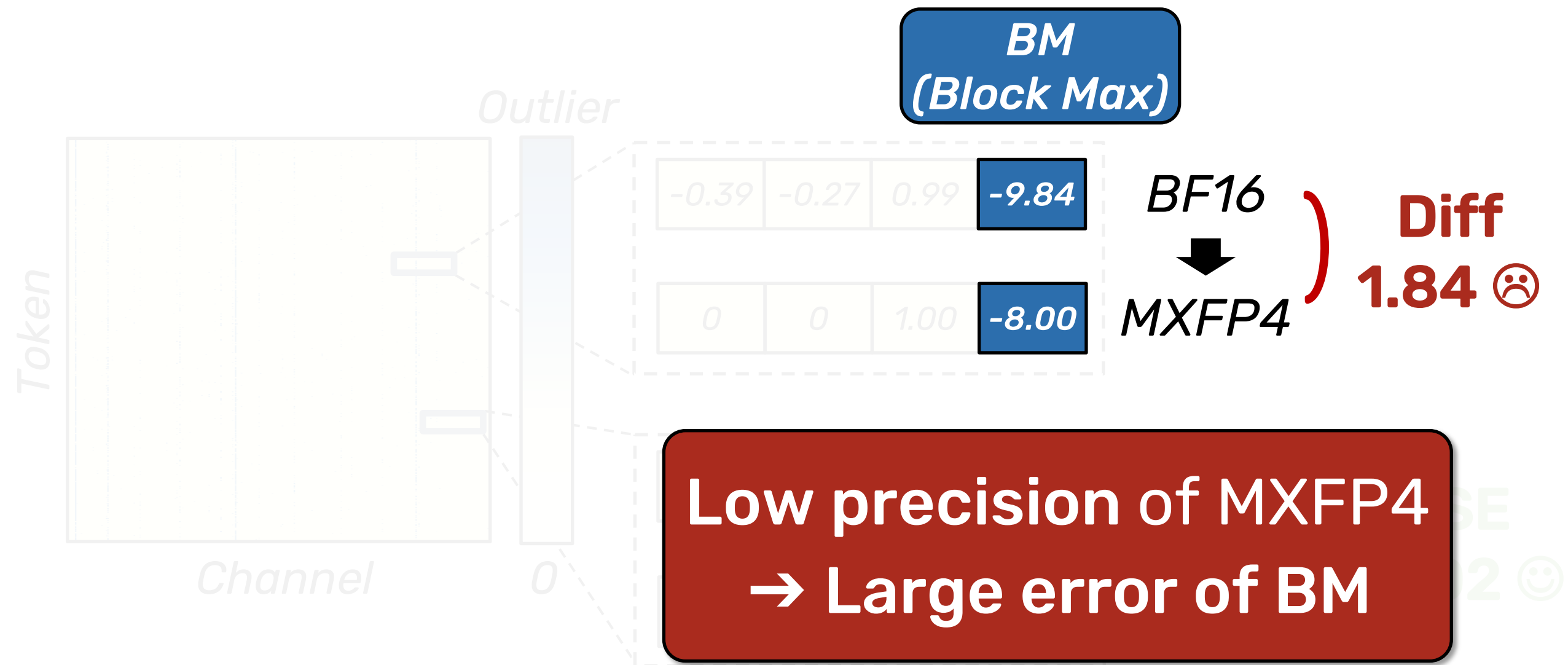
Analysis on Low-Bit MX Format



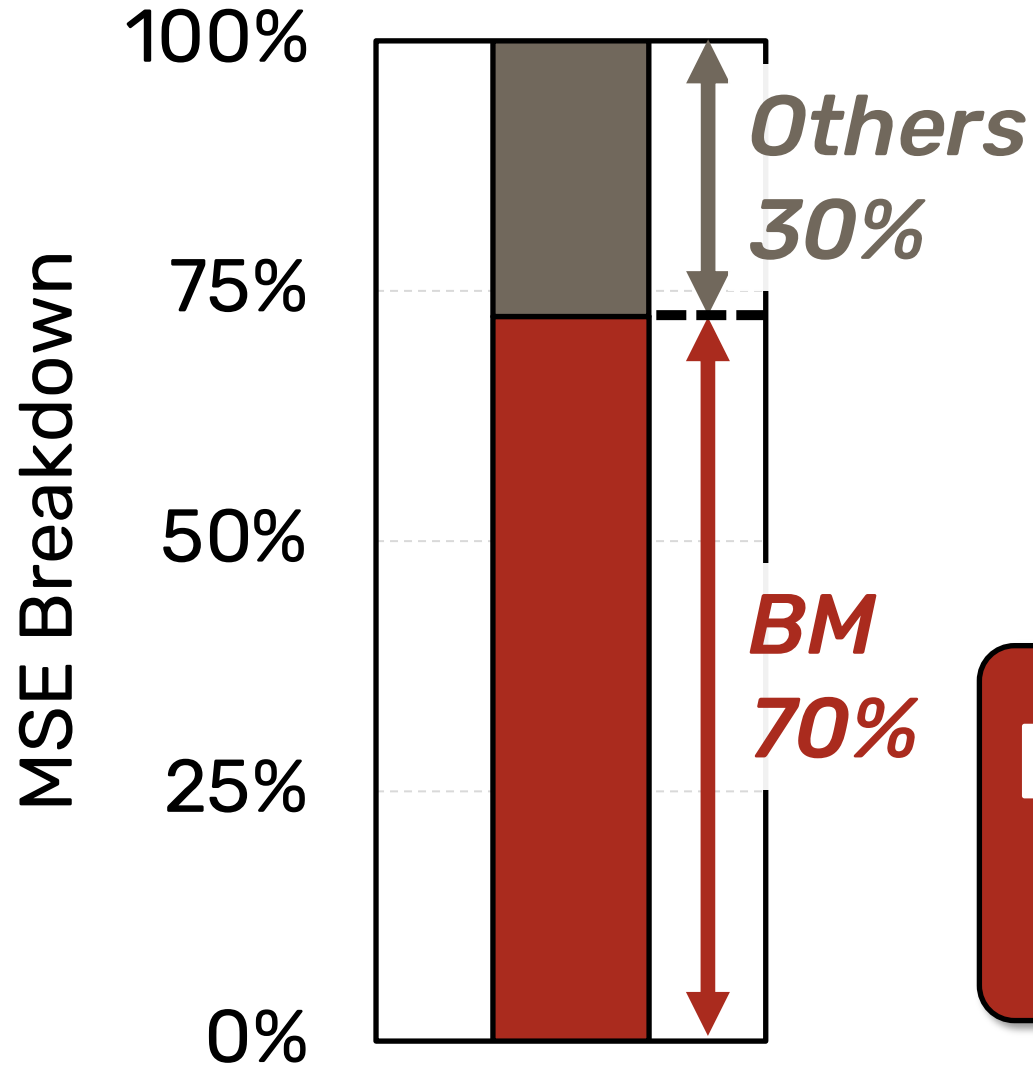
Analysis on Low-Bit MX Format



Analysis on Low-Bit MX Format

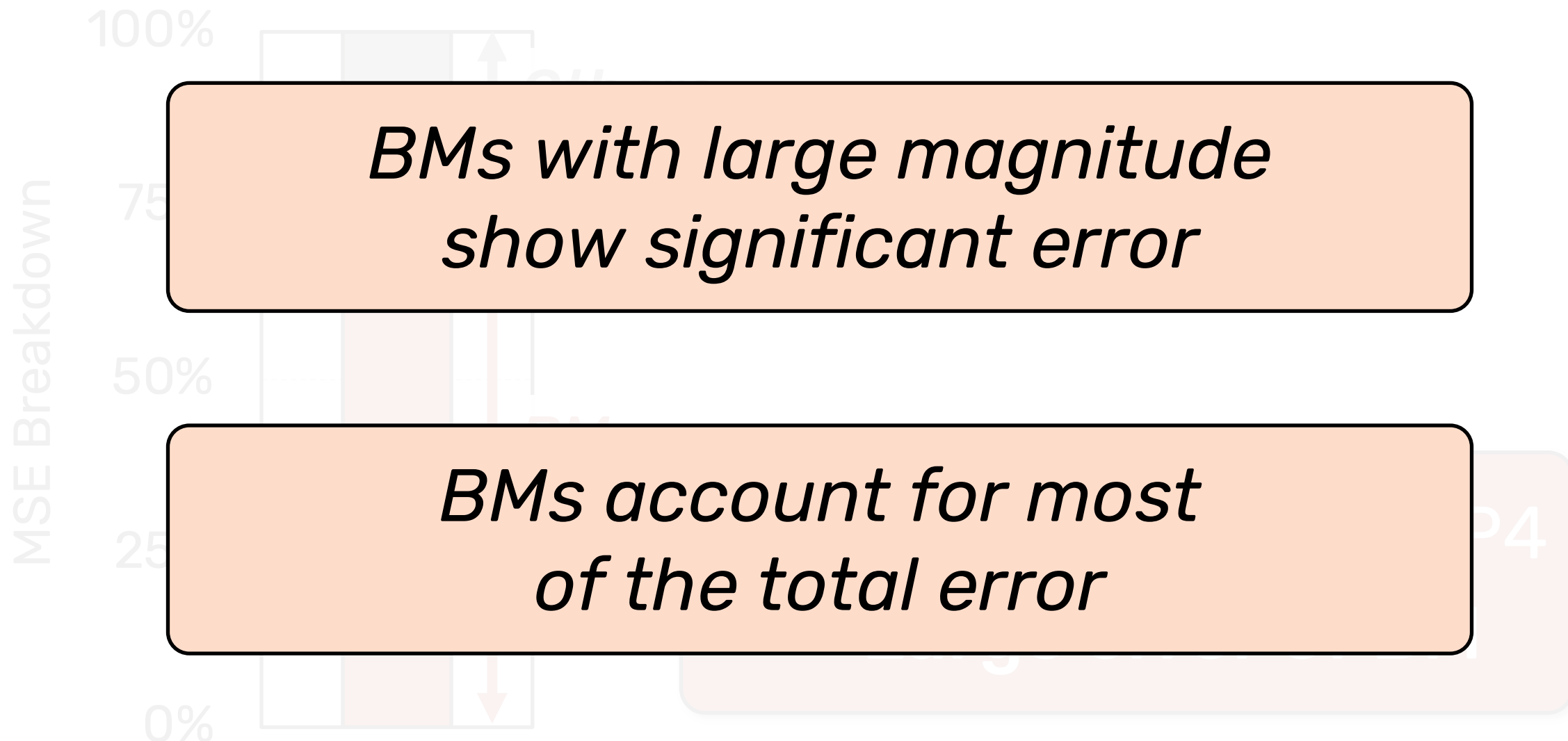


Analysis on Low-Bit MX Format



**Low precision of MXFP4
→ Large error of BM**

Analysis on Low-Bit MX Format



*BMs with large magnitude
show significant error*

*BMs account for most
of the total error*

Analysis on Low-Bit MX Format

BMs with large magnitude

**BM elements need more
precise representation**

*BMs account for most
of the total error*

Outline

- Introduction to MX Formats
- Motivation
 - Implications for Low-bit LLM Serving
 - Analysis on Low-bit MX Formats
- ***MX+*: Cost-Effective and Non-intrusive Extension to MX Formats**
 - Key Insight
 - Software & Hardware Integration on GPUs
- Evaluation
- Conclusion

MX+ Overview

MX



MX+



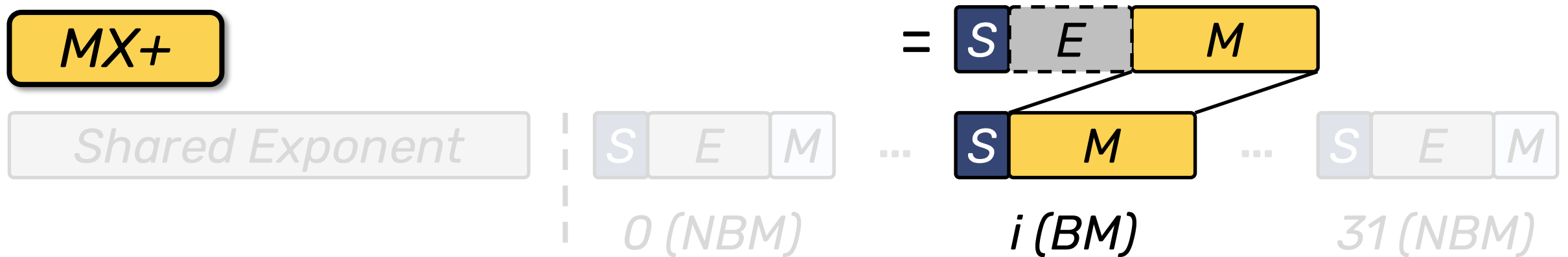
BM represented in high-precision 😊

MX+ Overview

MX



MX+



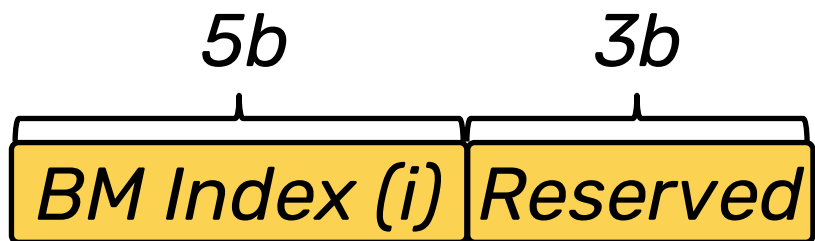
BM represented in high-precision 😊

MX+ Overview

MX

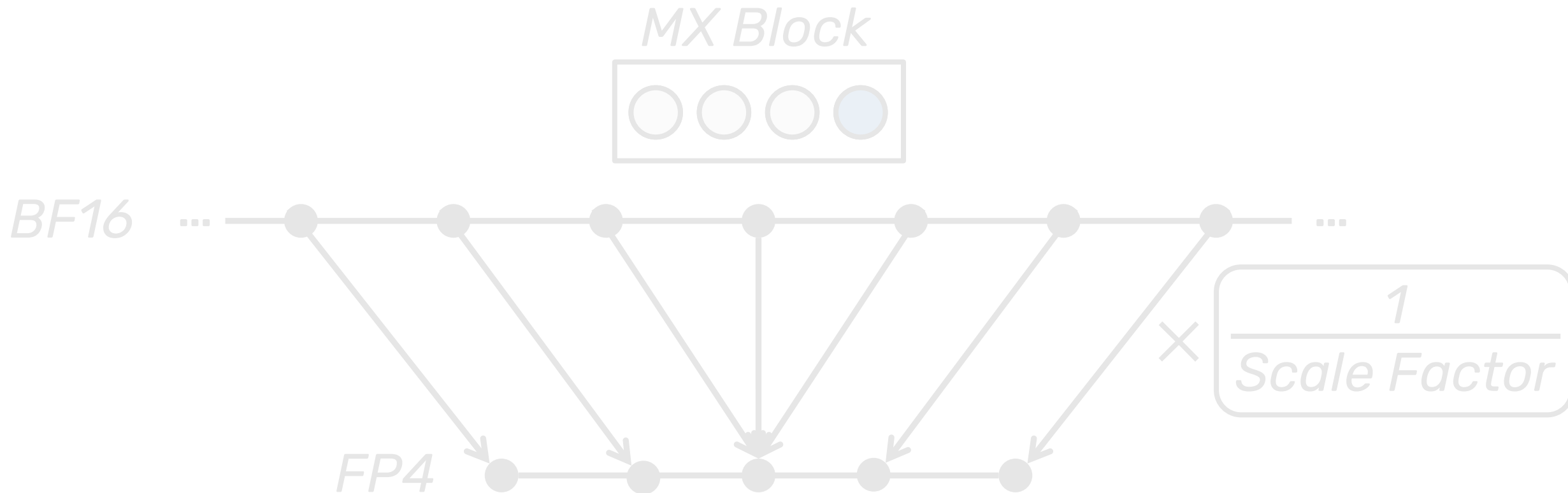


MX+



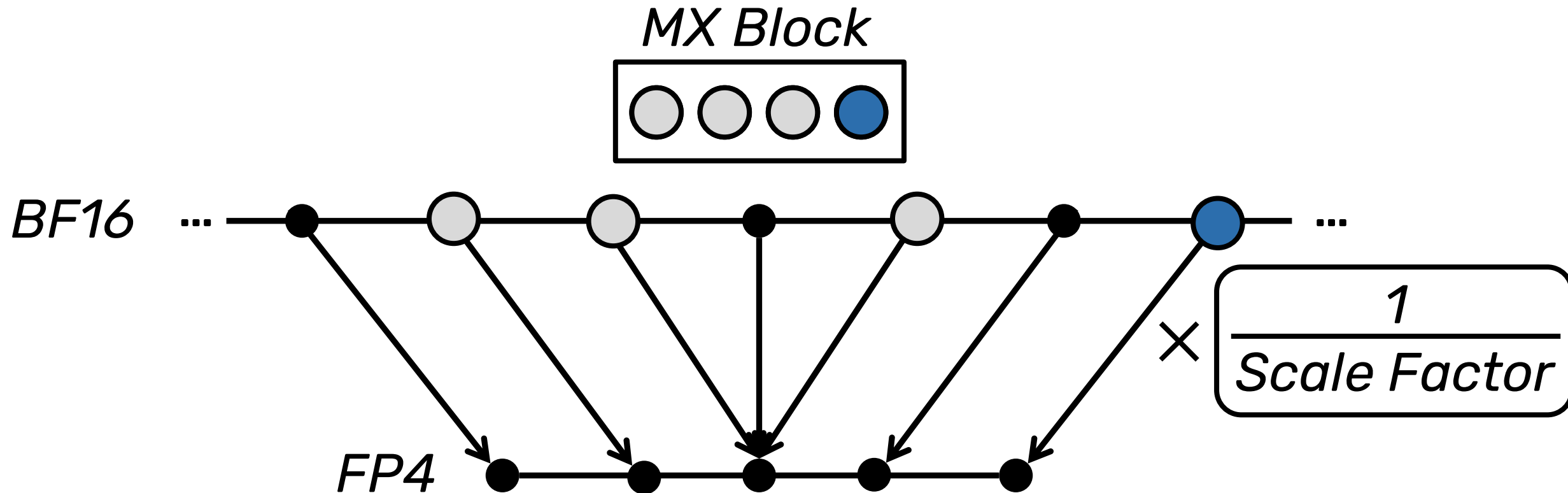
Key Insight

Exponent of BM is set to the max value



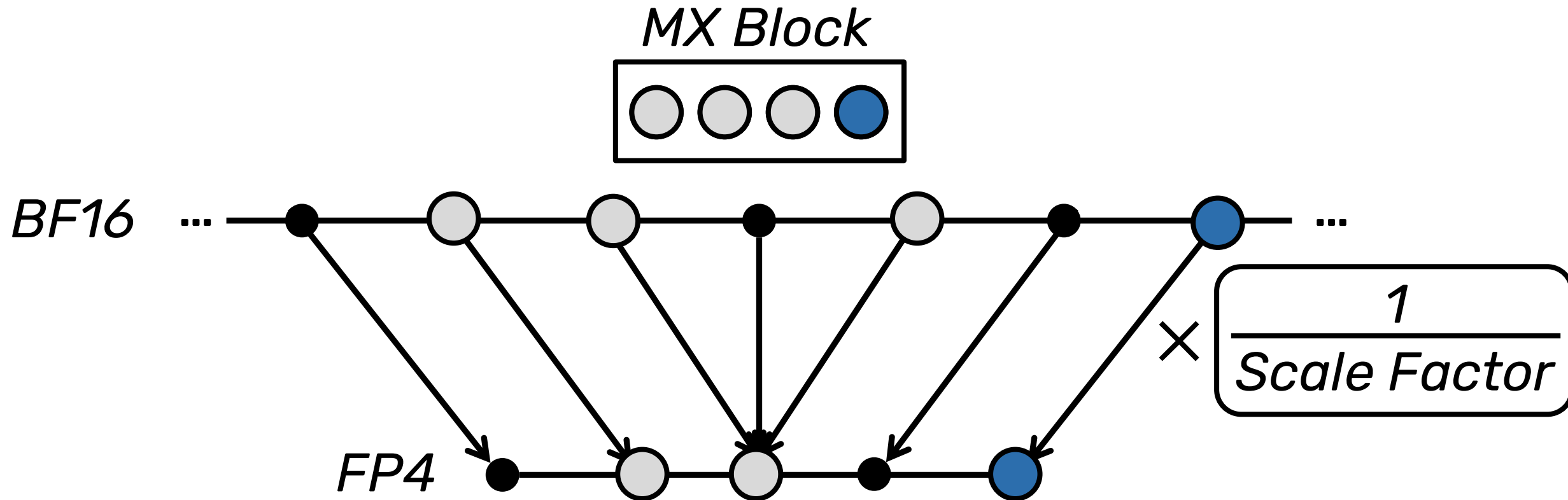
Key Insight

Exponent of BM is set to the max value



Key Insight

Exponent of BM is set to the max value

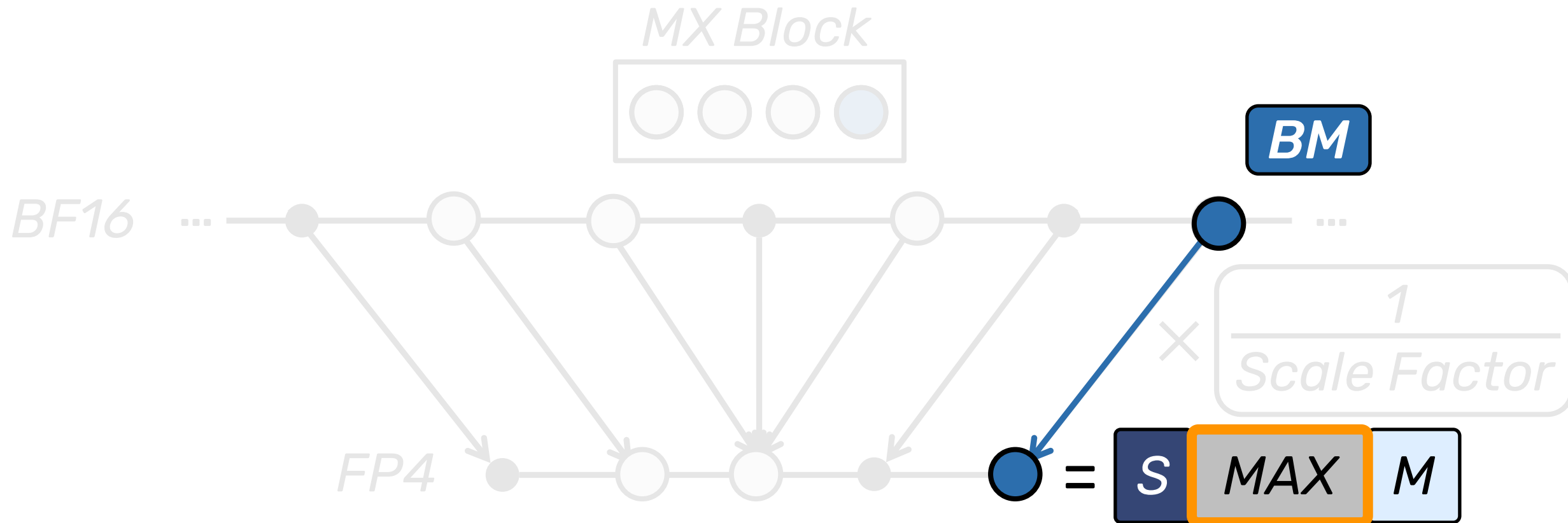


Exponent of BM is set to the max value

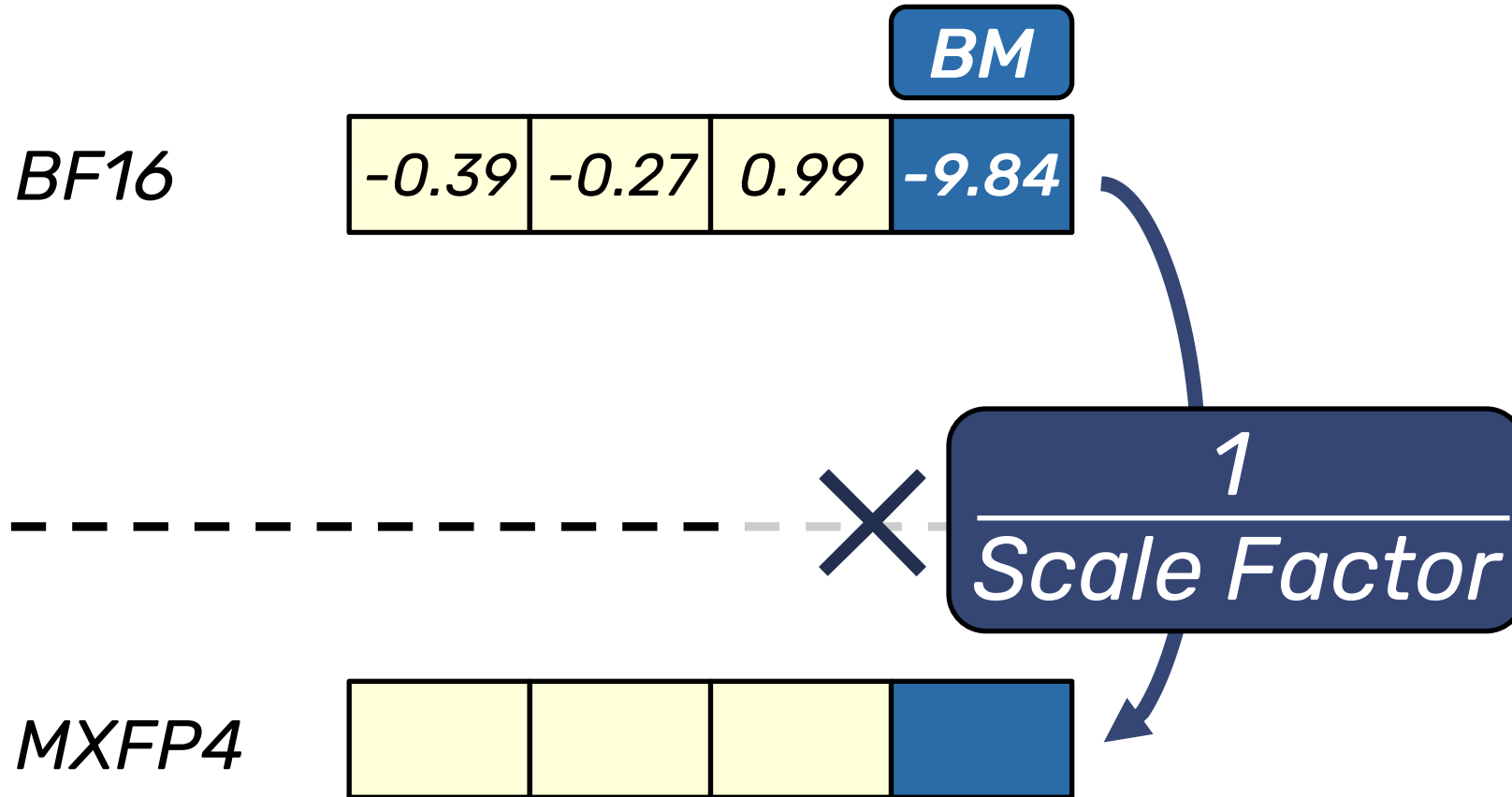


Key Insight

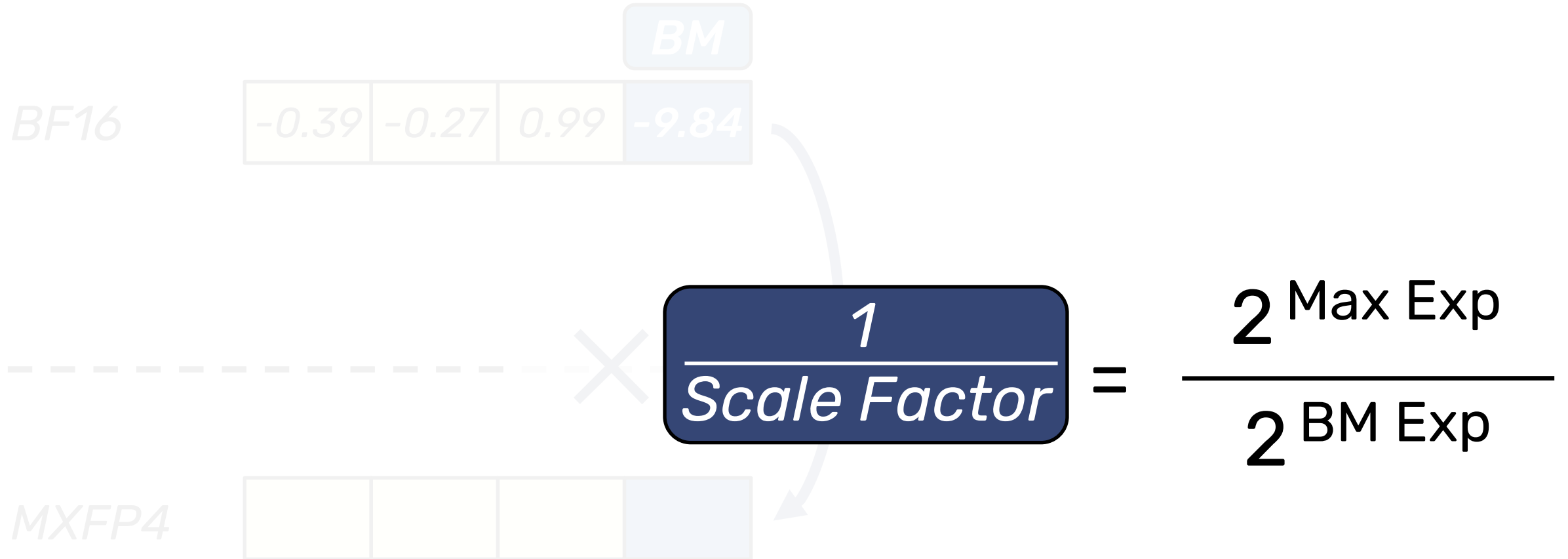
Exponent of BM is set to the max value



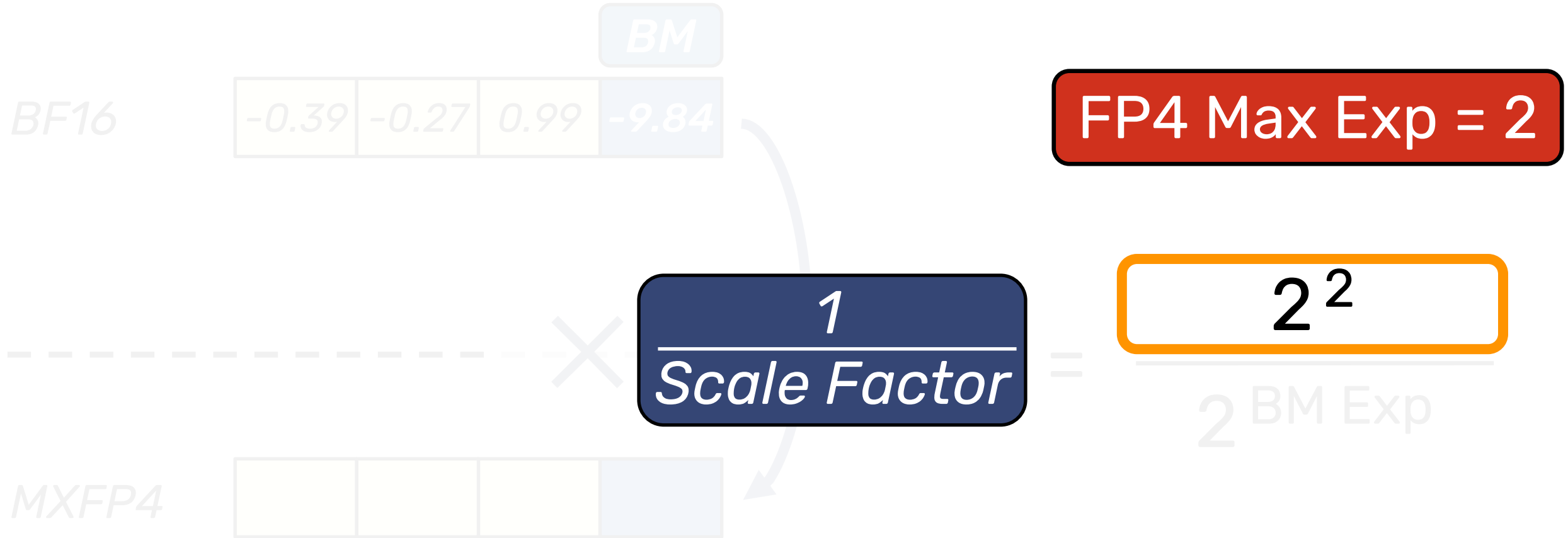
Conversion to MXFP4



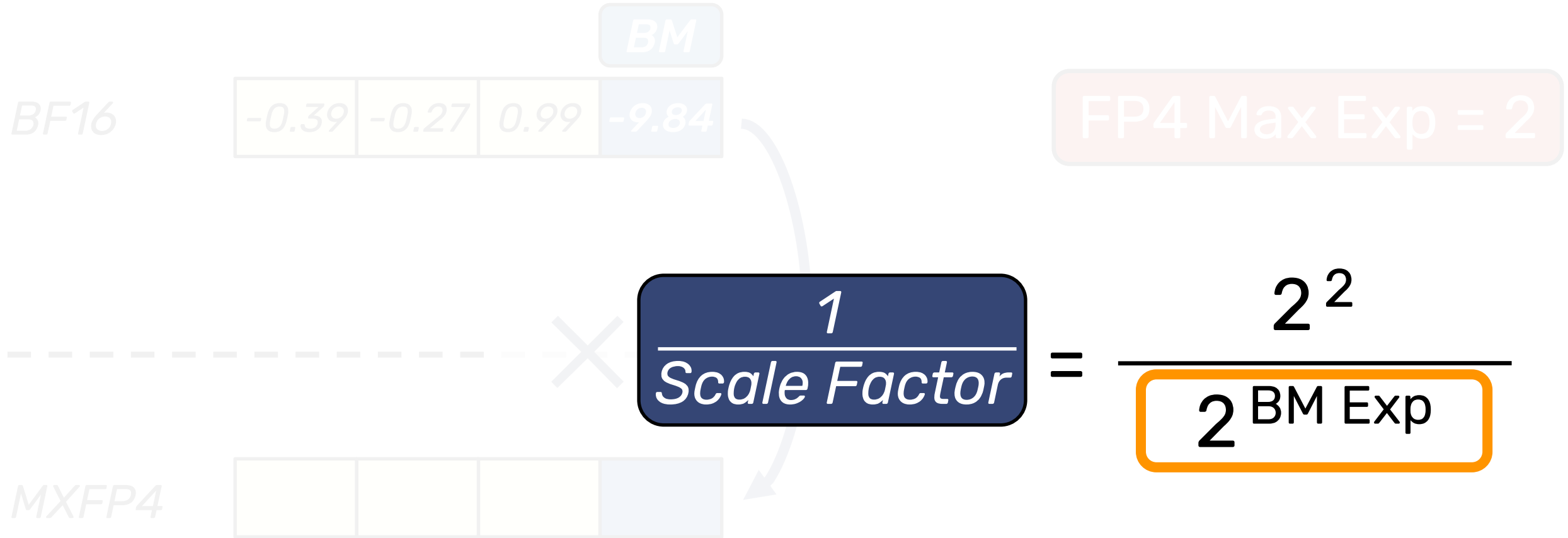
Conversion to MXFP4



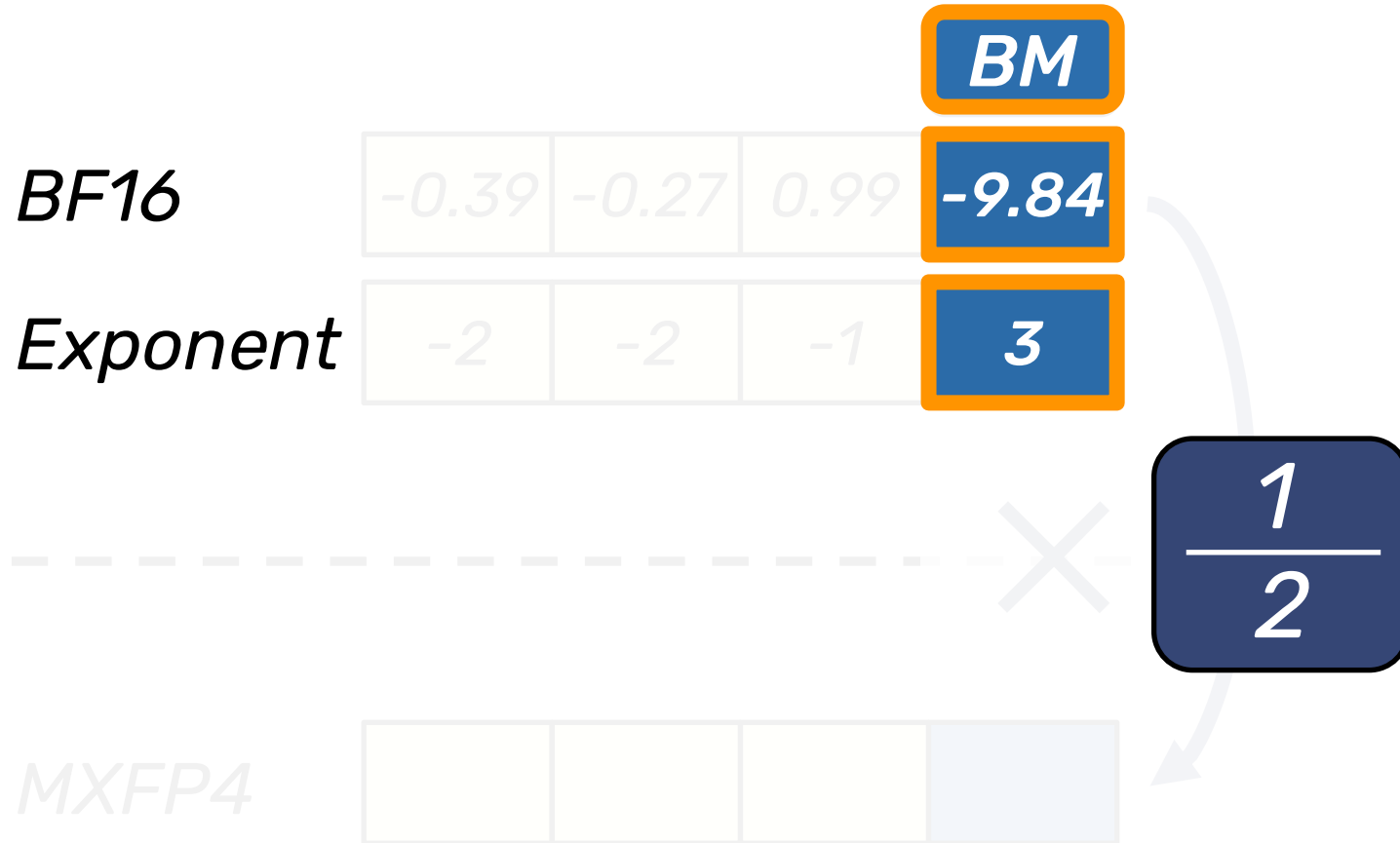
Conversion to MXFP4



Conversion to MXFP4



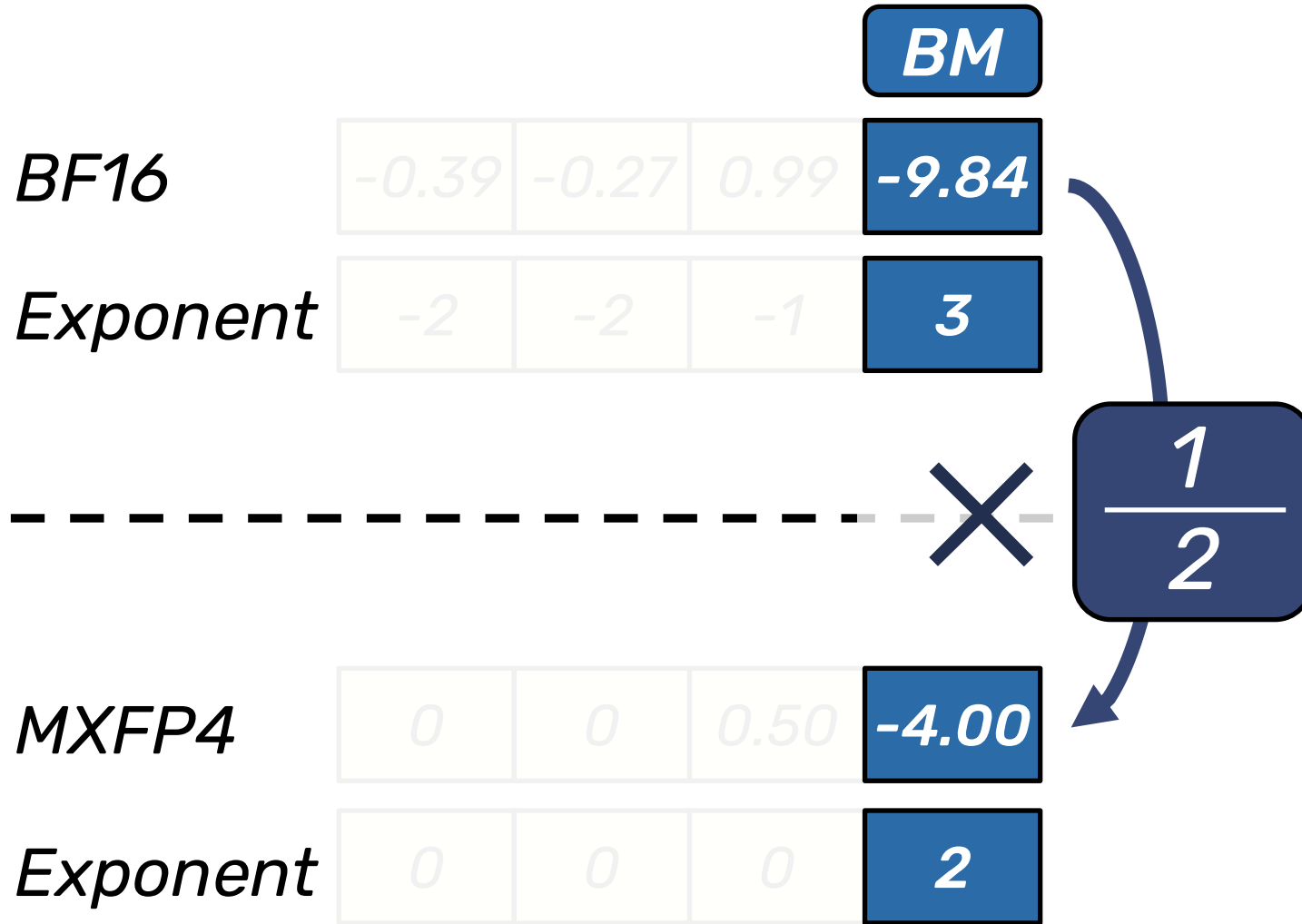
Conversion to MXFP4



FP4 Max Exp = 2

$$= \frac{2^2}{2^3}$$

Conversion to MXFP4



Conversion to MXFP4



FP4 Max Exp = 2

Conversion to MXFP4

BF16

			<i>BM</i>
-0.39	-0.27	0.99	-9.84
-2	-2	-1	3

Exponent

Store mantissa bits instead of exponent

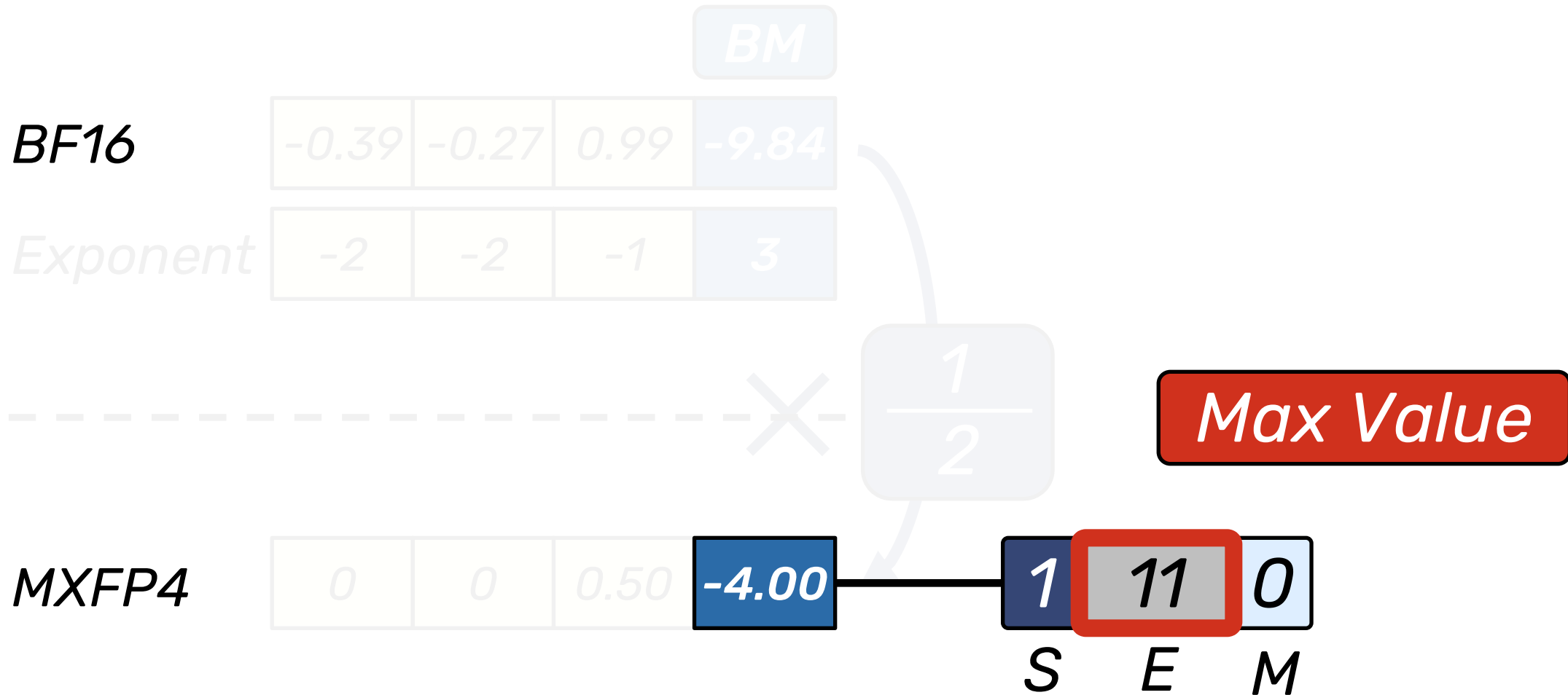
MXFP4

0	0	0.50	-4.00
0	0	0	2

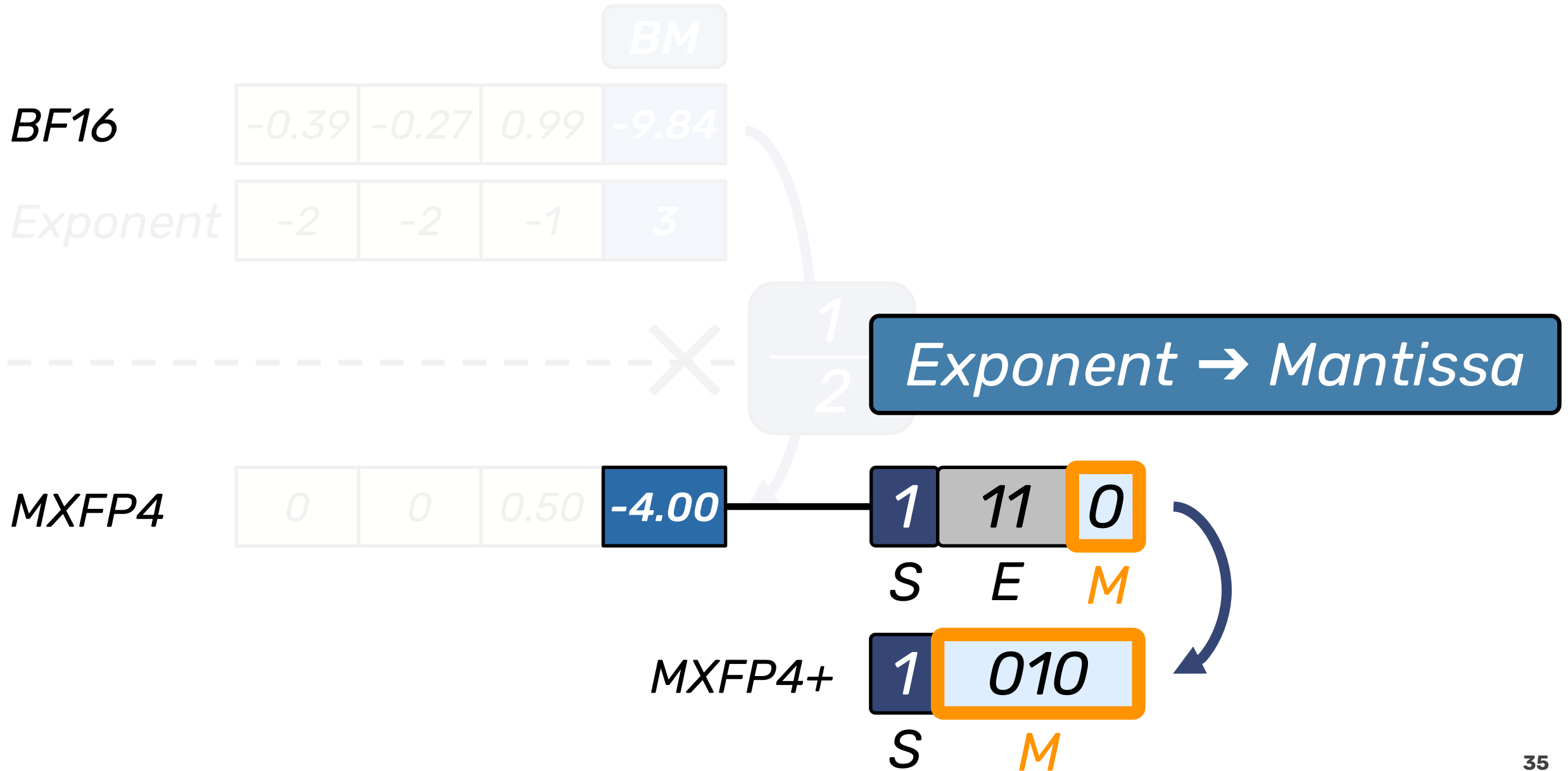
Exponent

FP4 Max Exp = 2

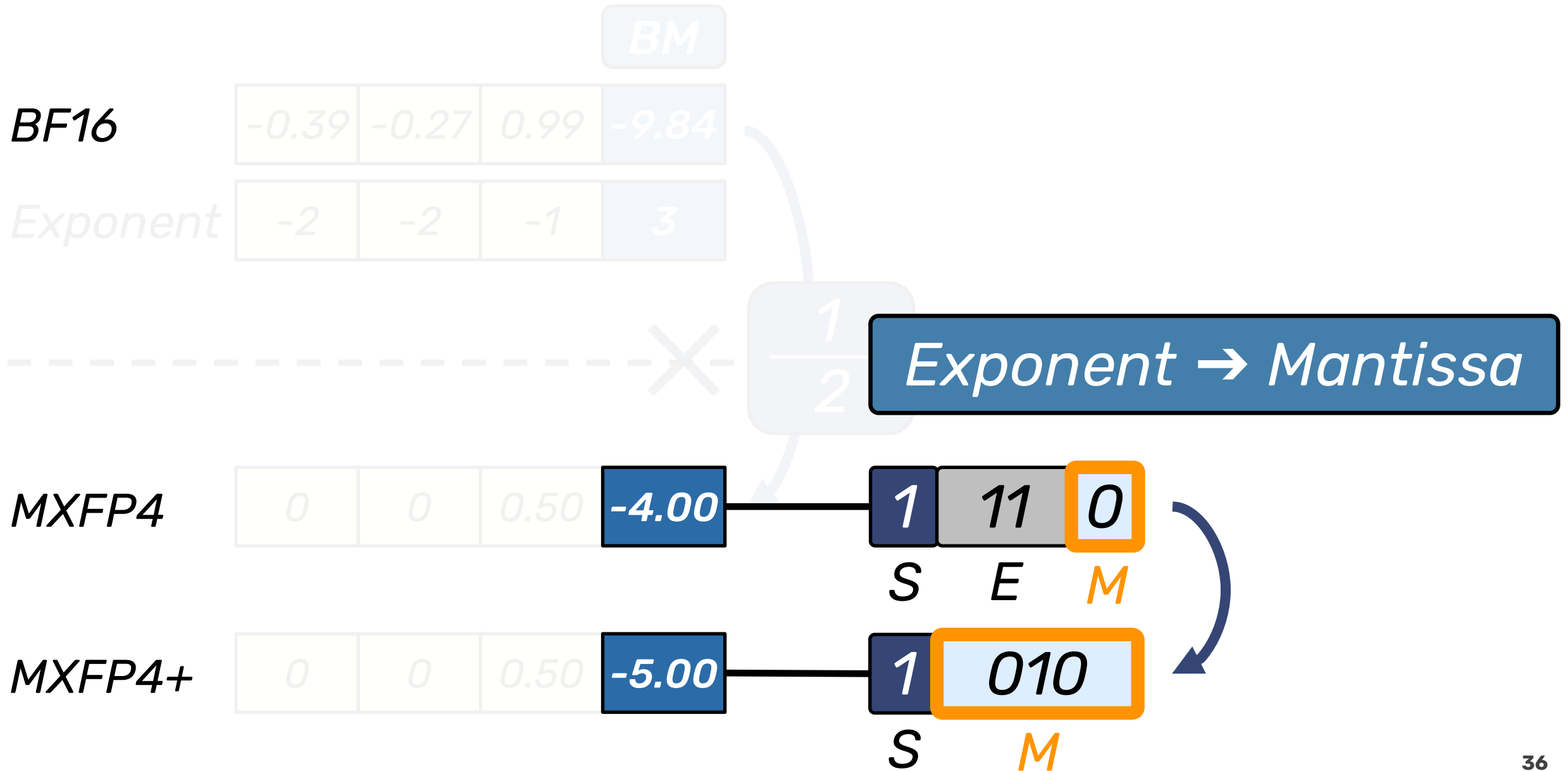
MX+ Extension



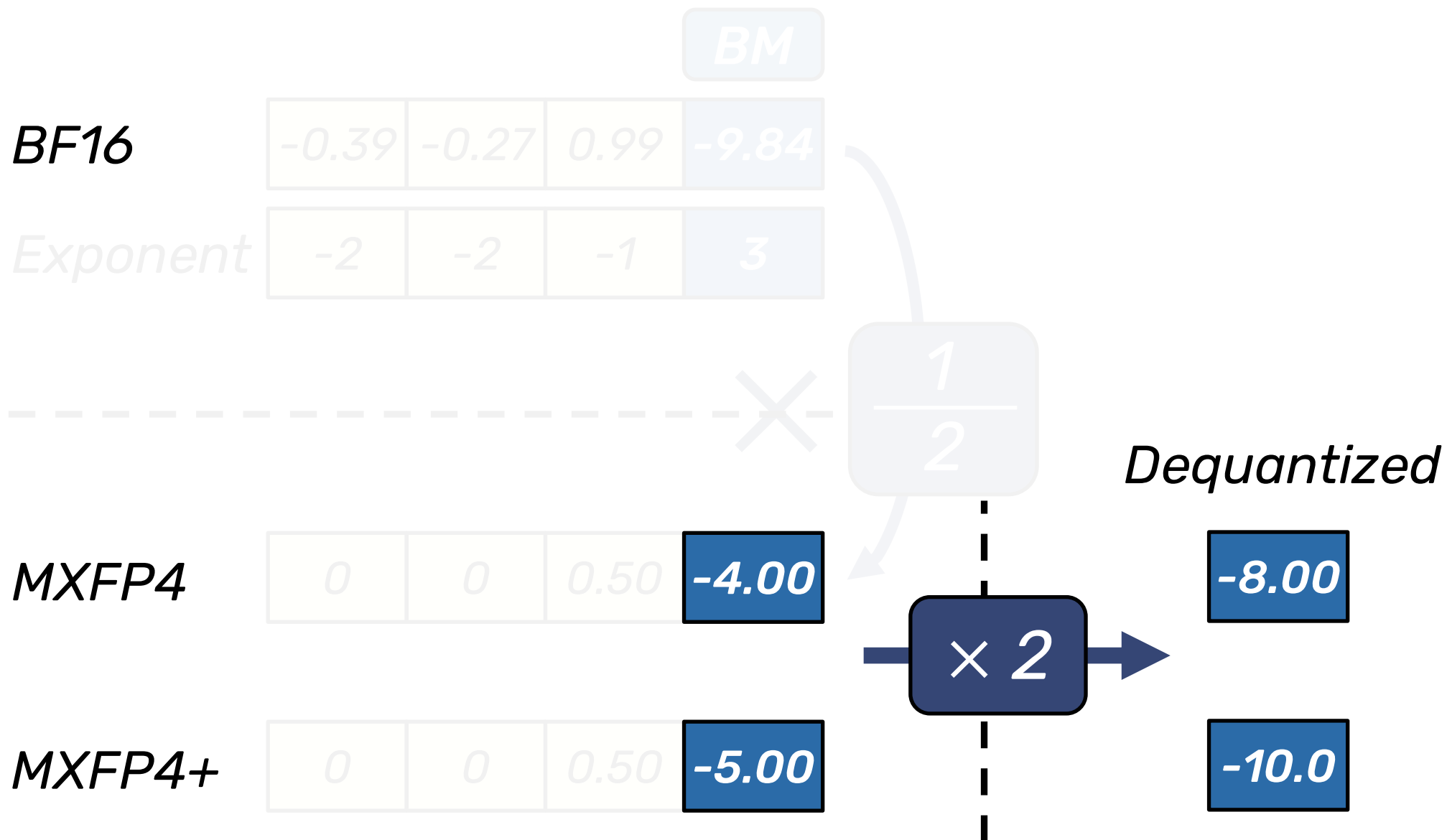
MX+ Extension



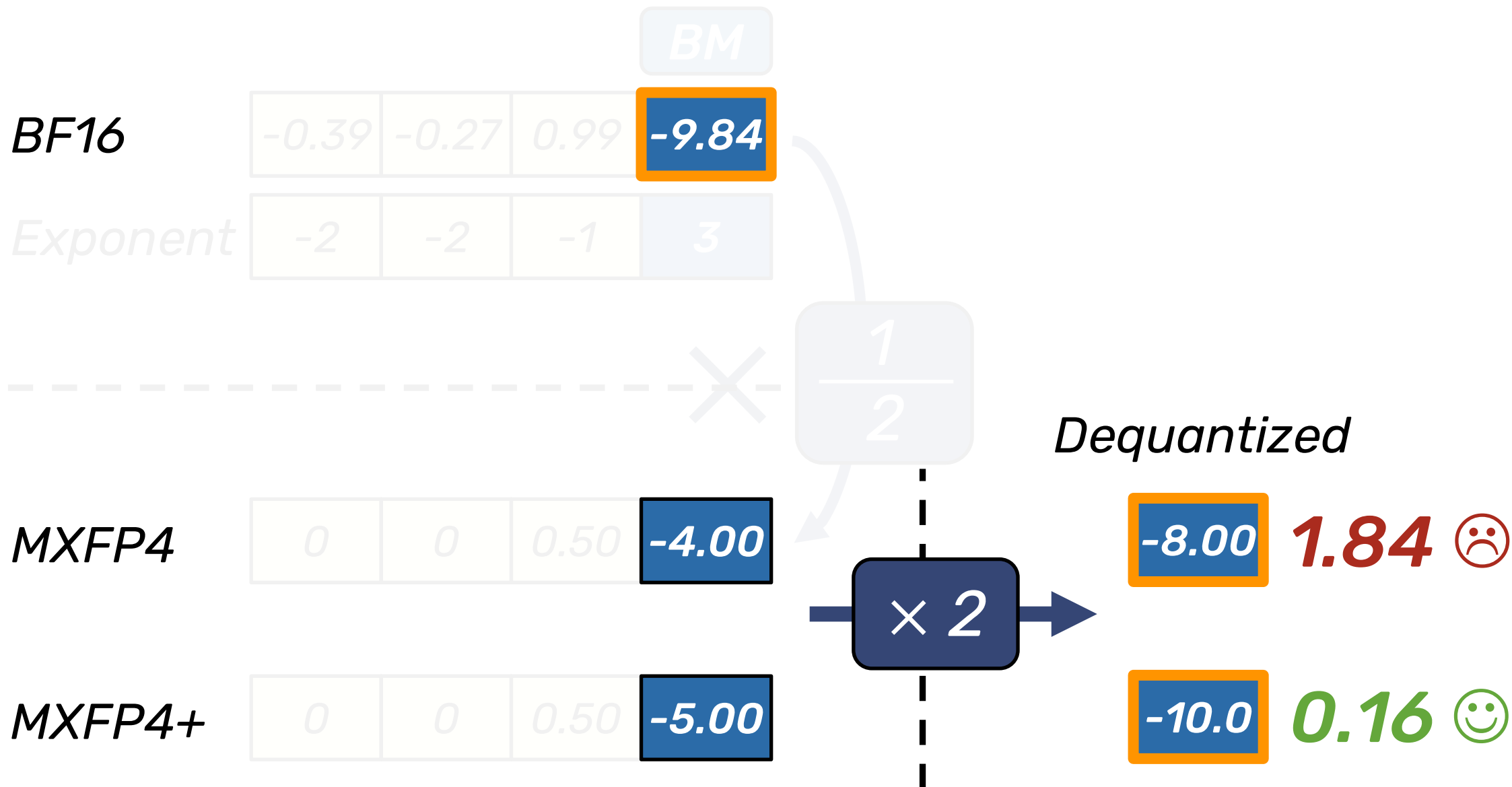
MX+ Extension



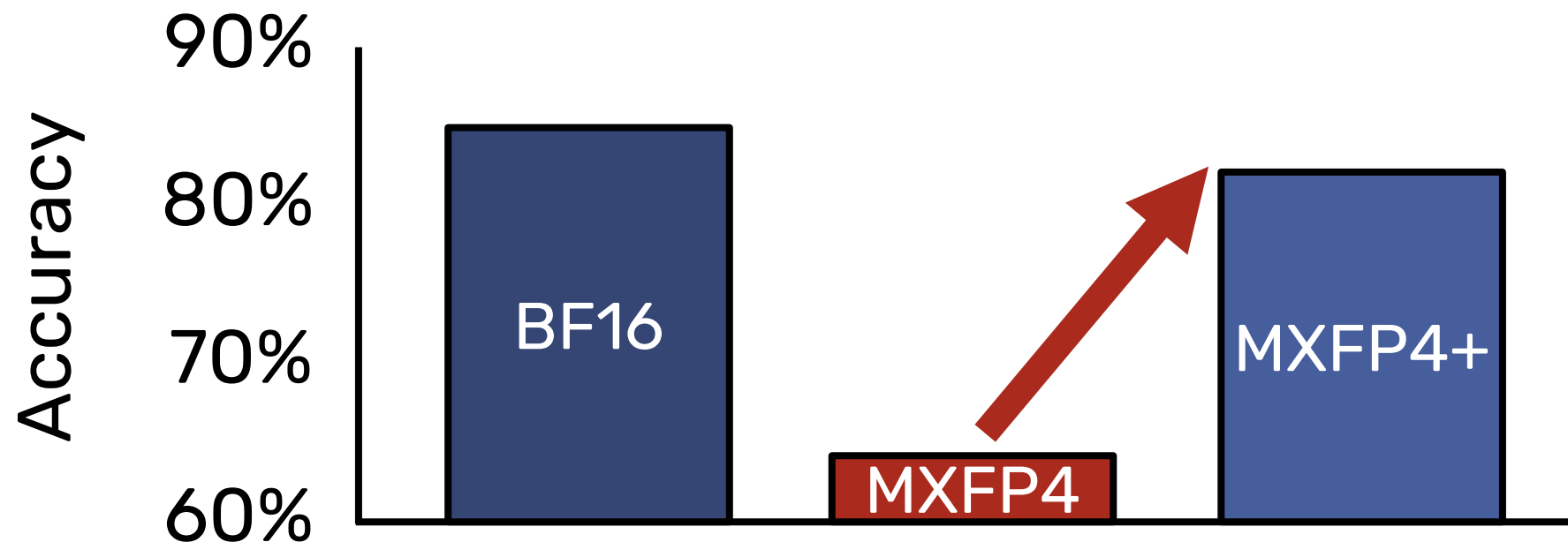
MX+ Extension



MX+ Extension



MX+ Extension



MXFP4

-4.00

MXFP4

-8.00

1.84 ☹️

MXFP4+

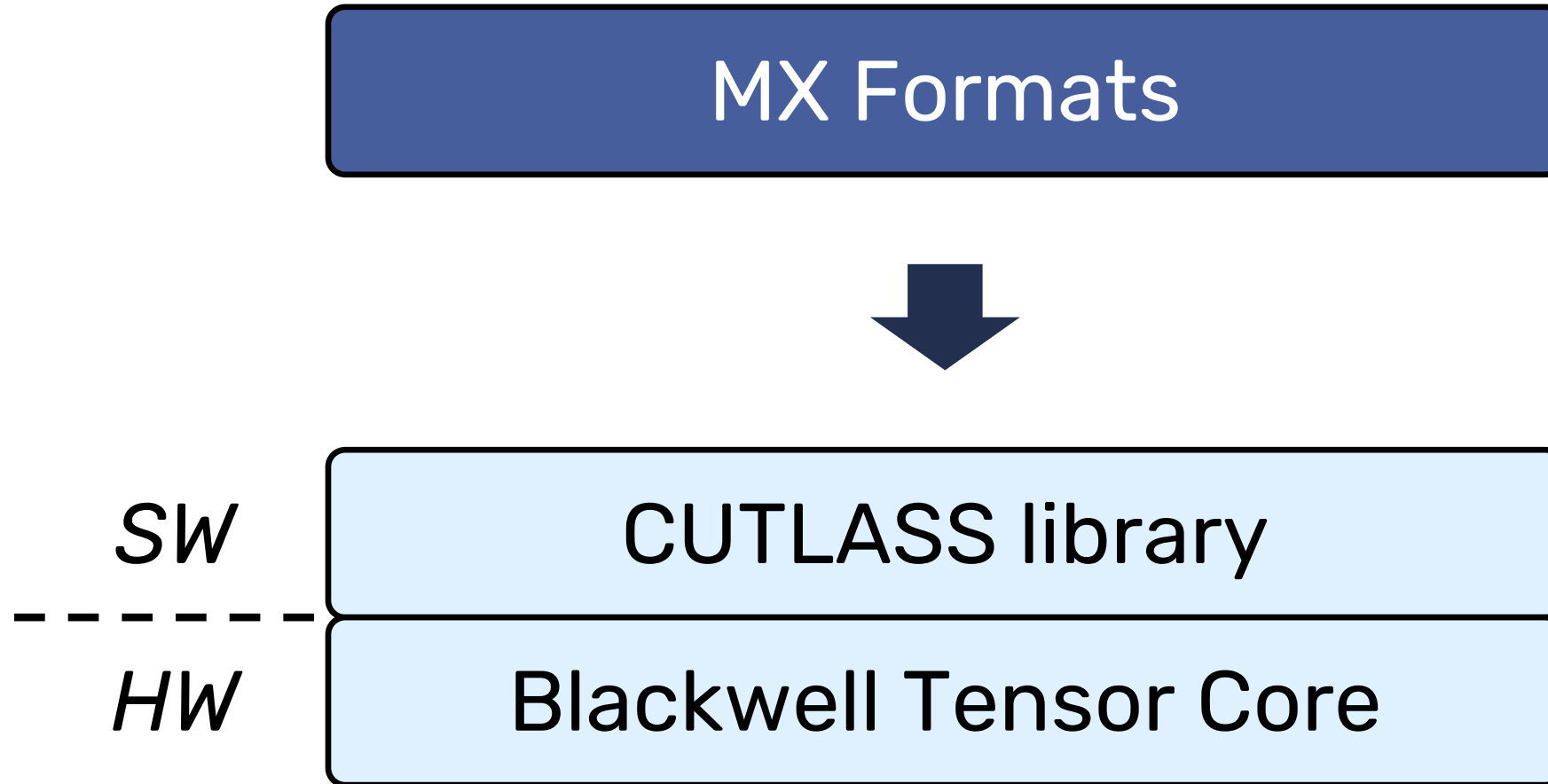
-5.00

MXFP4+

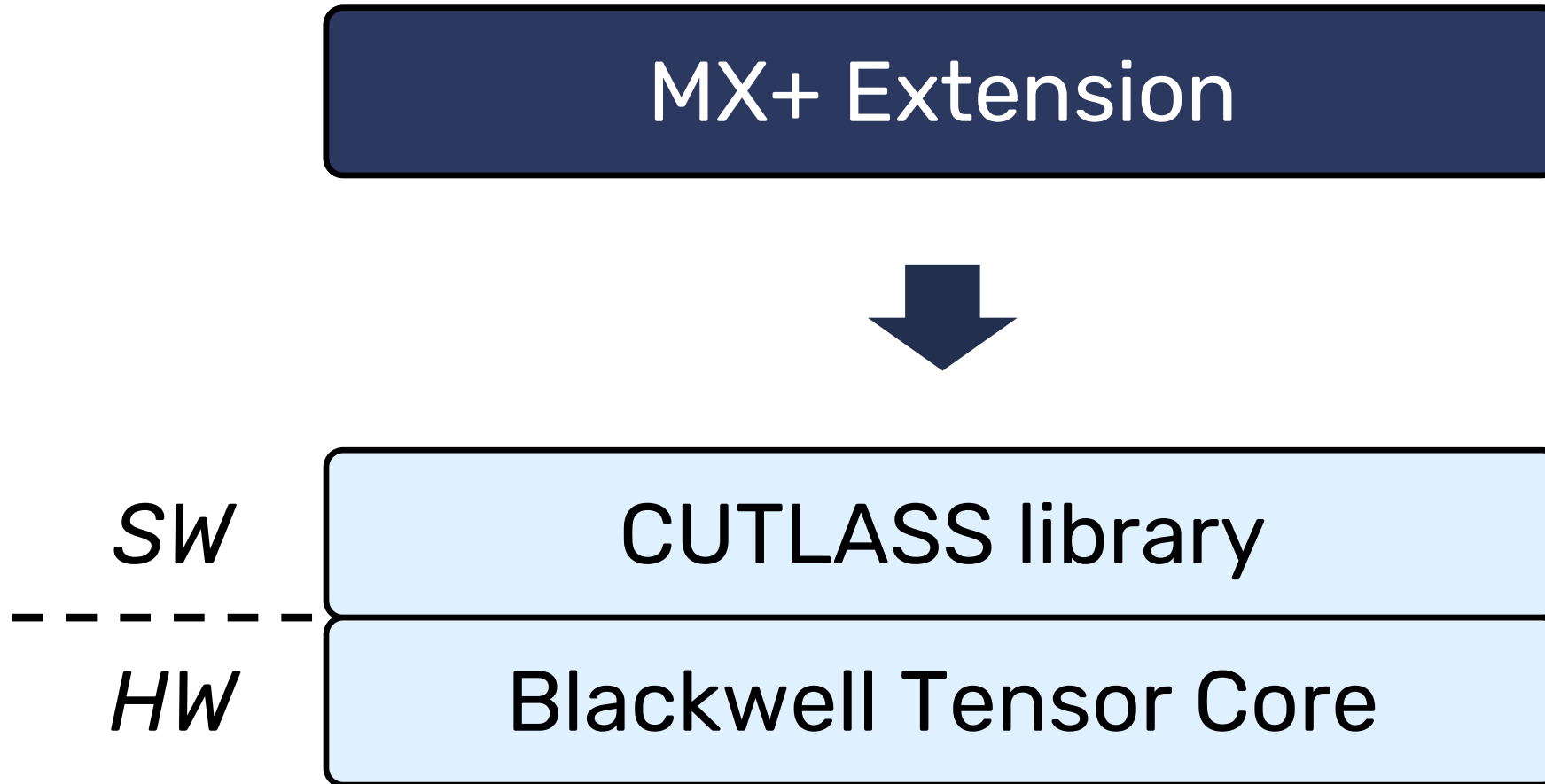
-10.0

0.16 😊

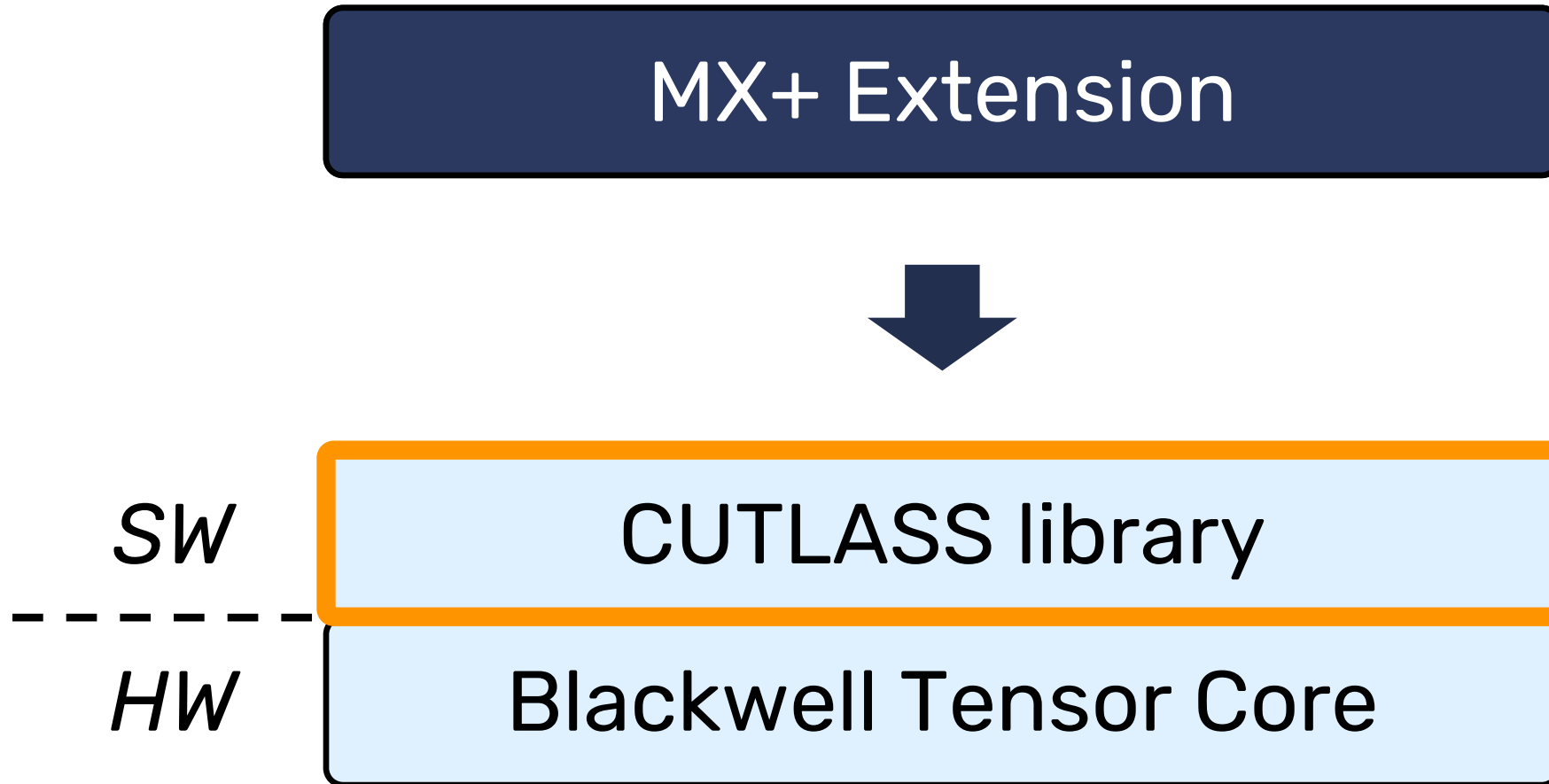
Integration of MX+ on GPUs



Integration of MX+ on GPUs



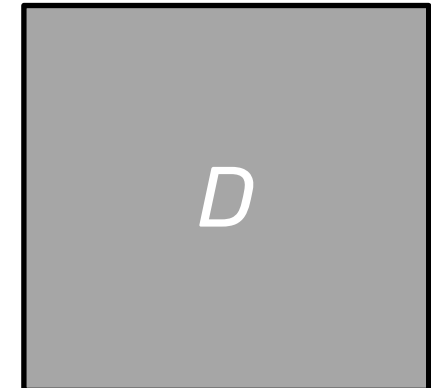
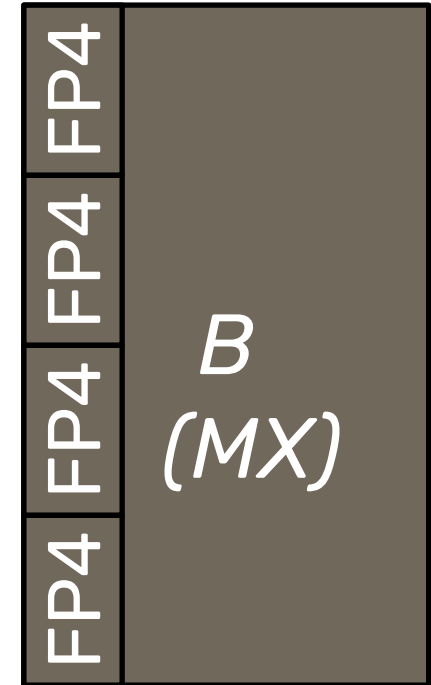
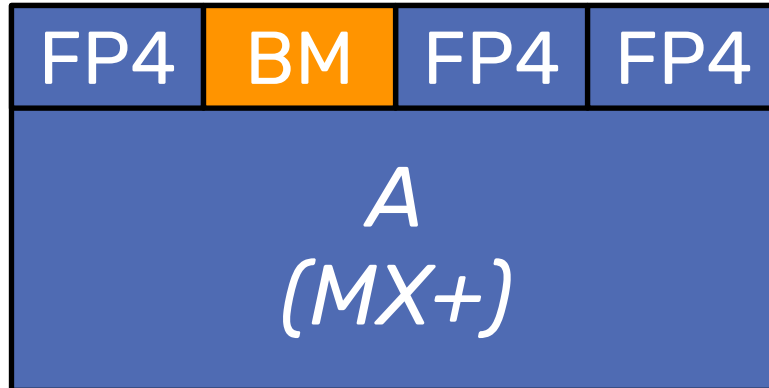
Integration of MX+ on GPUs



Software Integration of MX+ on GPUs

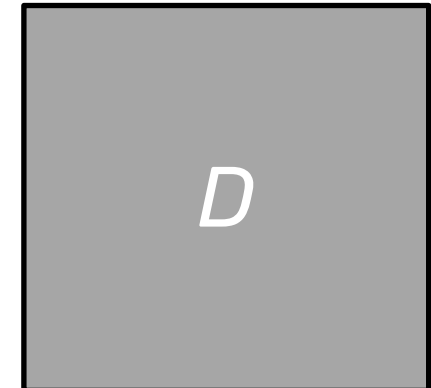
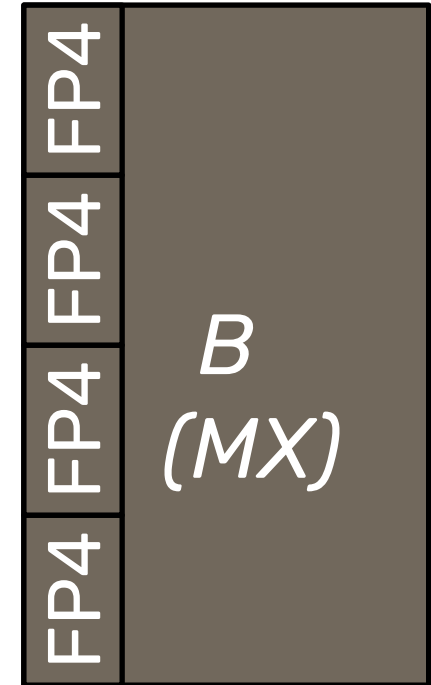
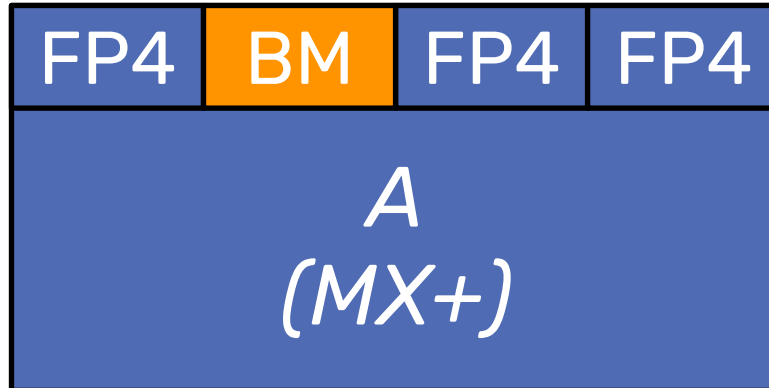
MMA (Matrix Multiply and Accumulate)

$$D += A \times B$$



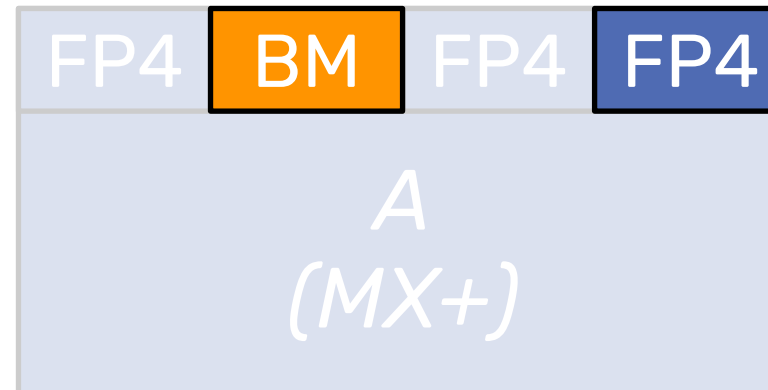
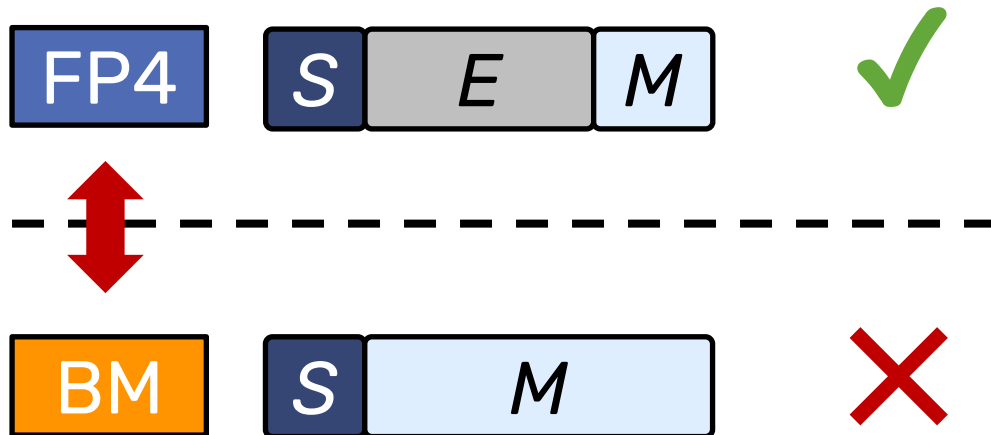
Software Integration of MX+ on GPUs

Tensor Core cannot perform
MMA operation with MX+ ☹️



Software Integration of MX+ on GPUs

Tensor Core cannot perform
MMA operation with MX+ 😞



Software Integration of MX+ on GPUs

Tensor Core cannot perform
MMA operation with MX+ ☹️

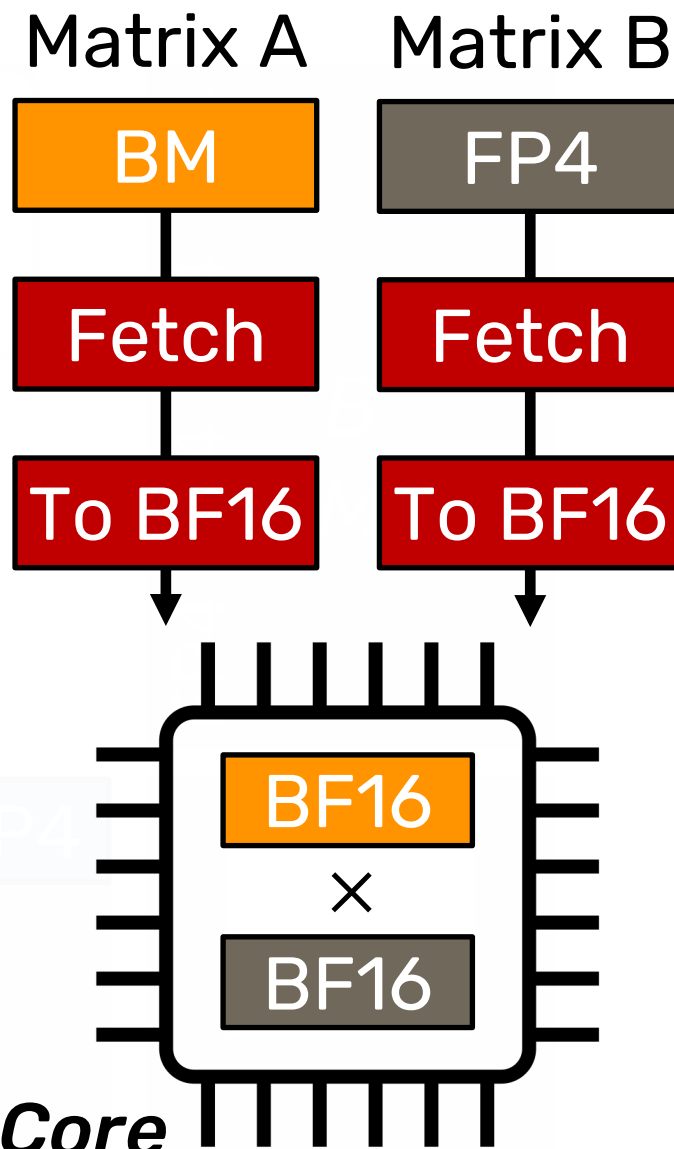
→ Compute BMs in CUDA Core?



Software Integration of MX+ on GPUs

Tensor Core cannot perform
MMA operation with MX+ ☹️

→ Compute BMs in CUDA Core? ❌



Software Integration of MX+ on GPUs

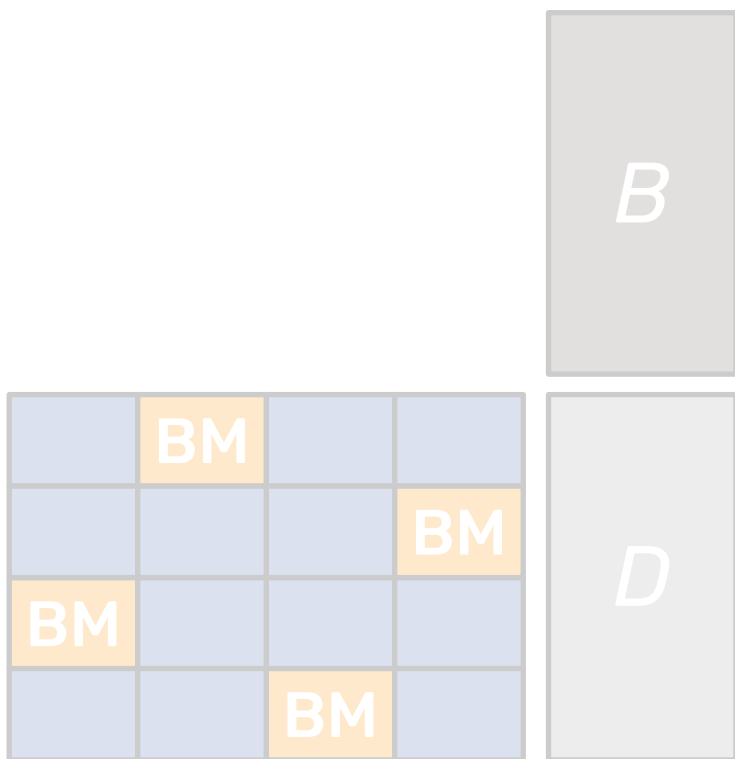
Tensor Core cannot perform
MMA operation with MX+ 😞

→ Compute BMs in Tensor Core ✓

Maintains Performance
in the Decode Stage 😊

Software Integration of MX+ on GPUs

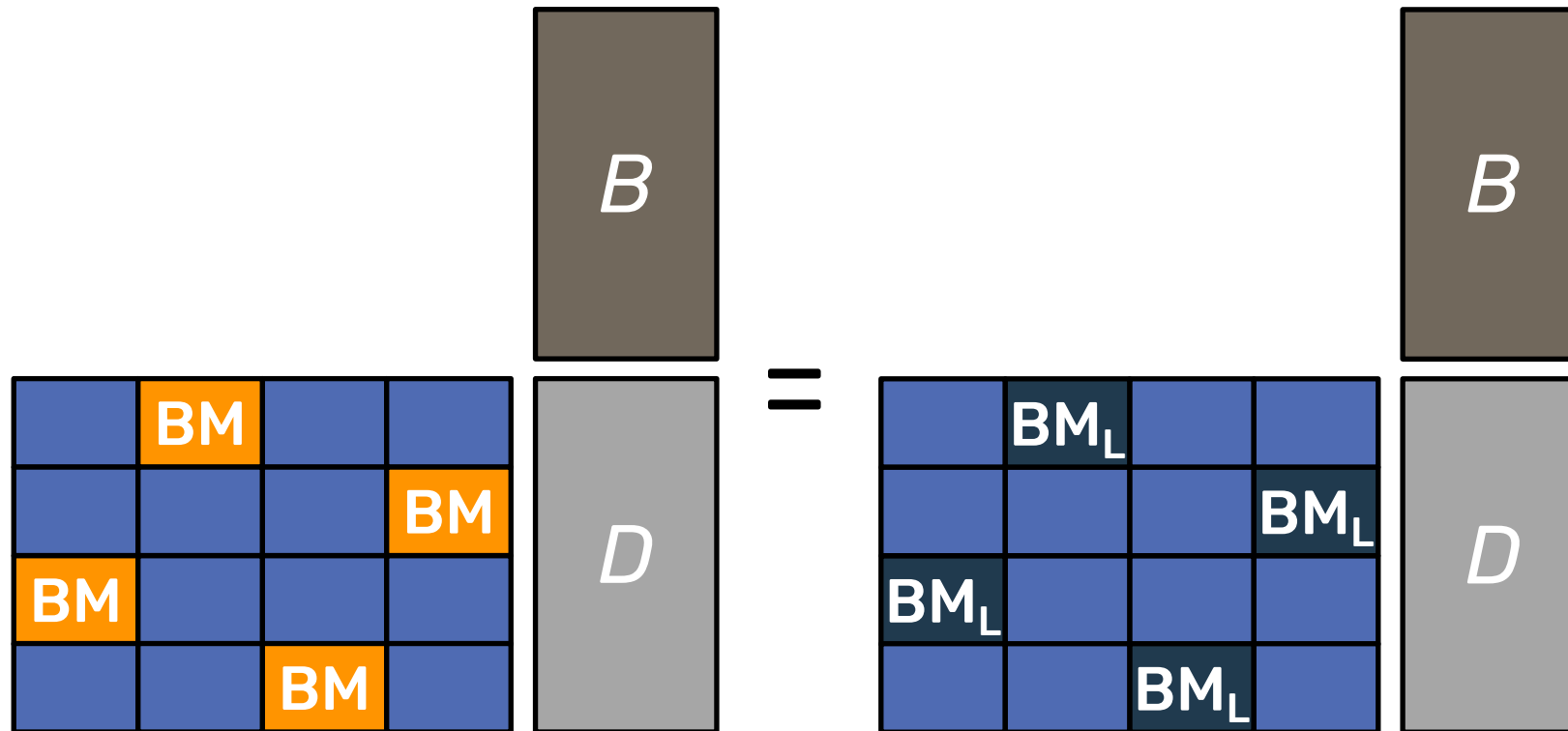
$$\boxed{\text{BM}} = \boxed{\text{BM}_L}^{FP4} + \boxed{\text{BM}_H}^{FP4}$$



Software Integration of MX+ on GPUs

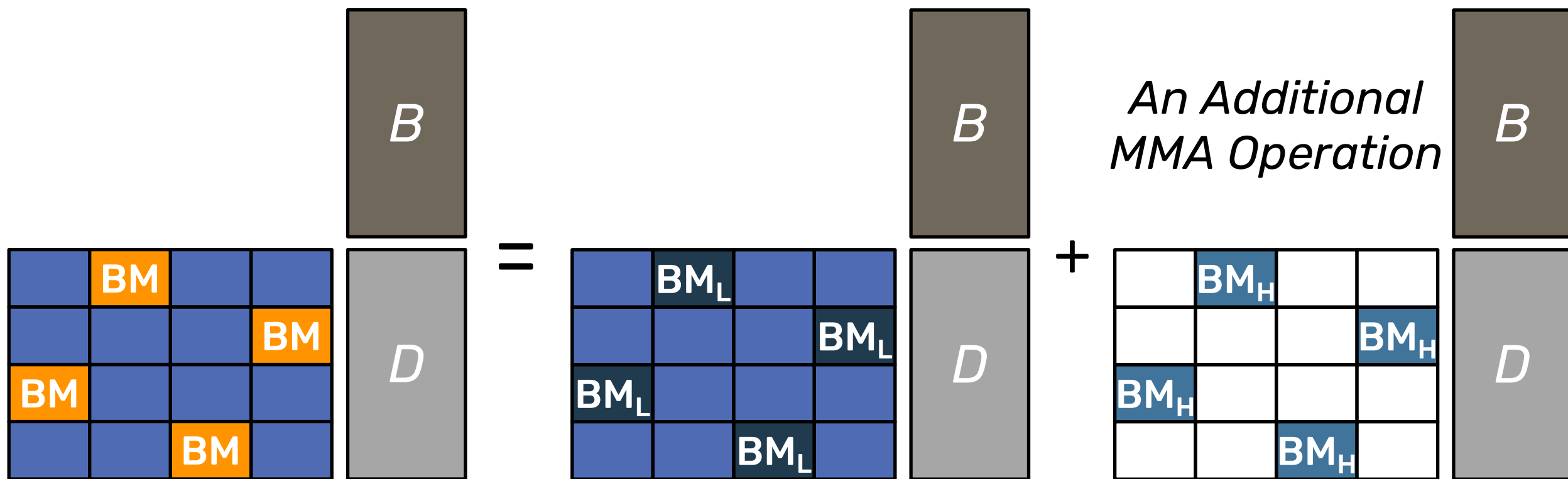
$$\text{BM} = \text{BM}_L + \text{BM}_H$$

FP4 *FP4*

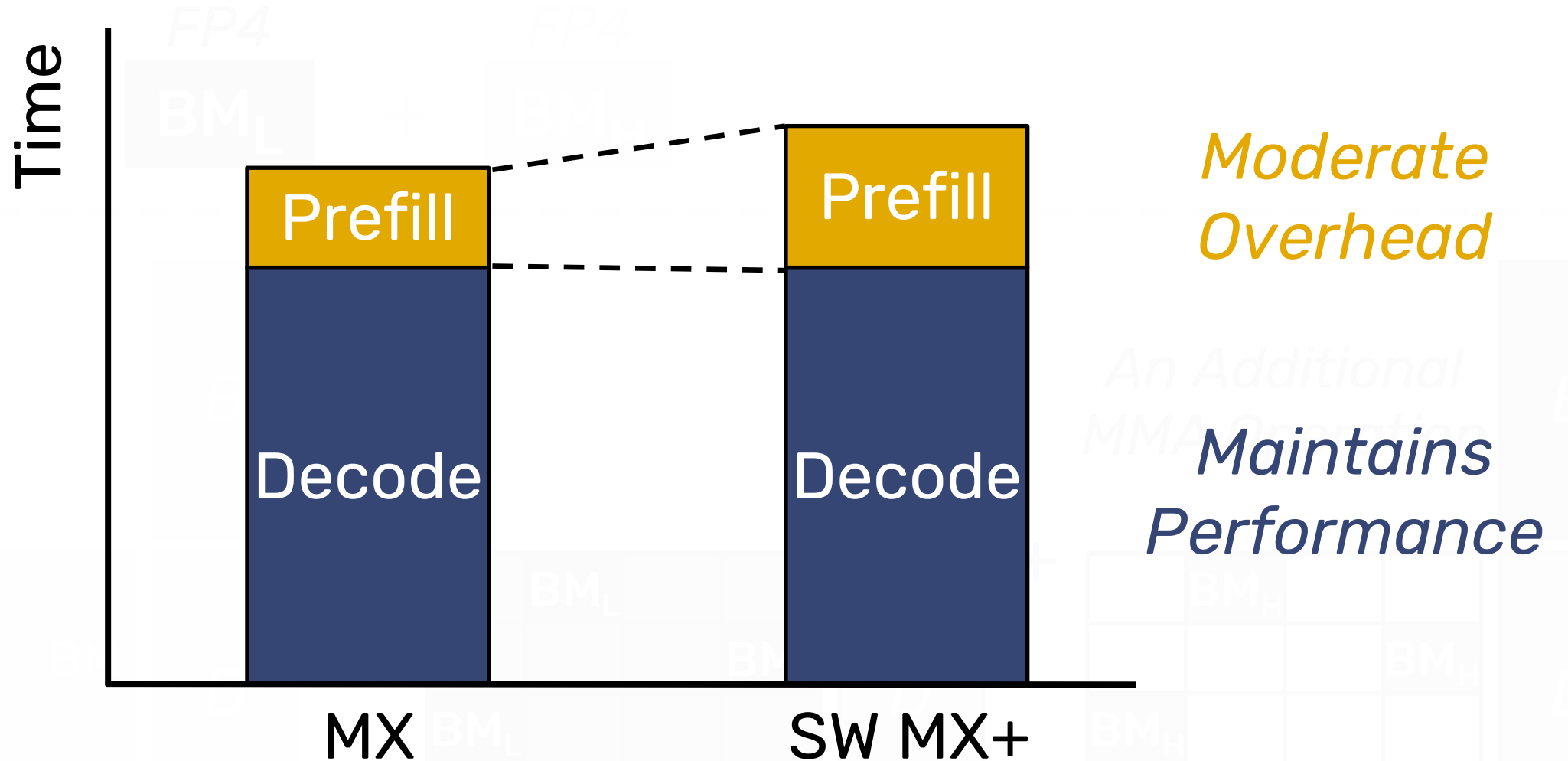


Software Integration of MX+ on GPUs

$$\text{BM} = \overset{FP4}{\text{BM}_L} + \overset{FP4}{\text{BM}_H}$$

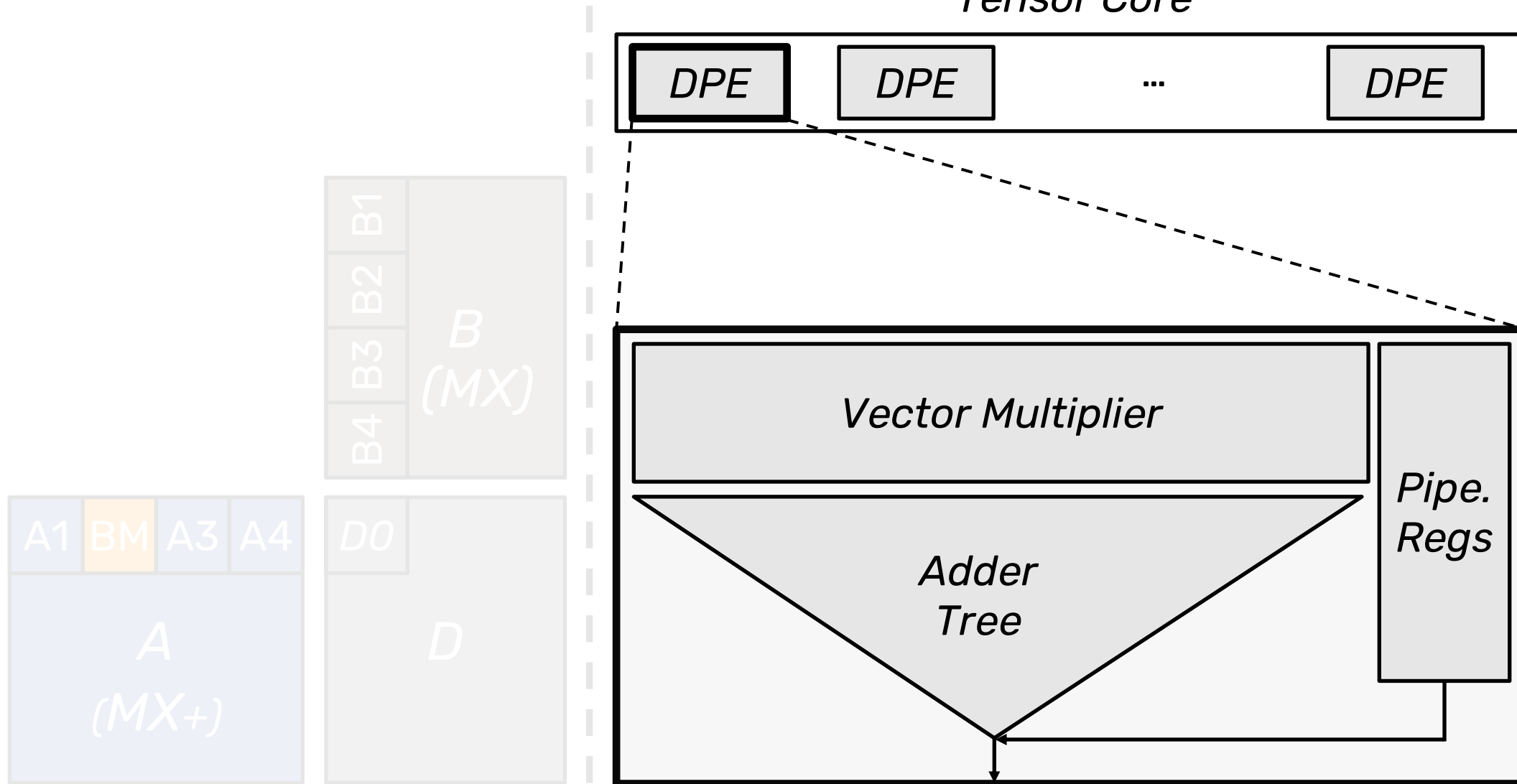


Software Integration of MX+ on GPUs

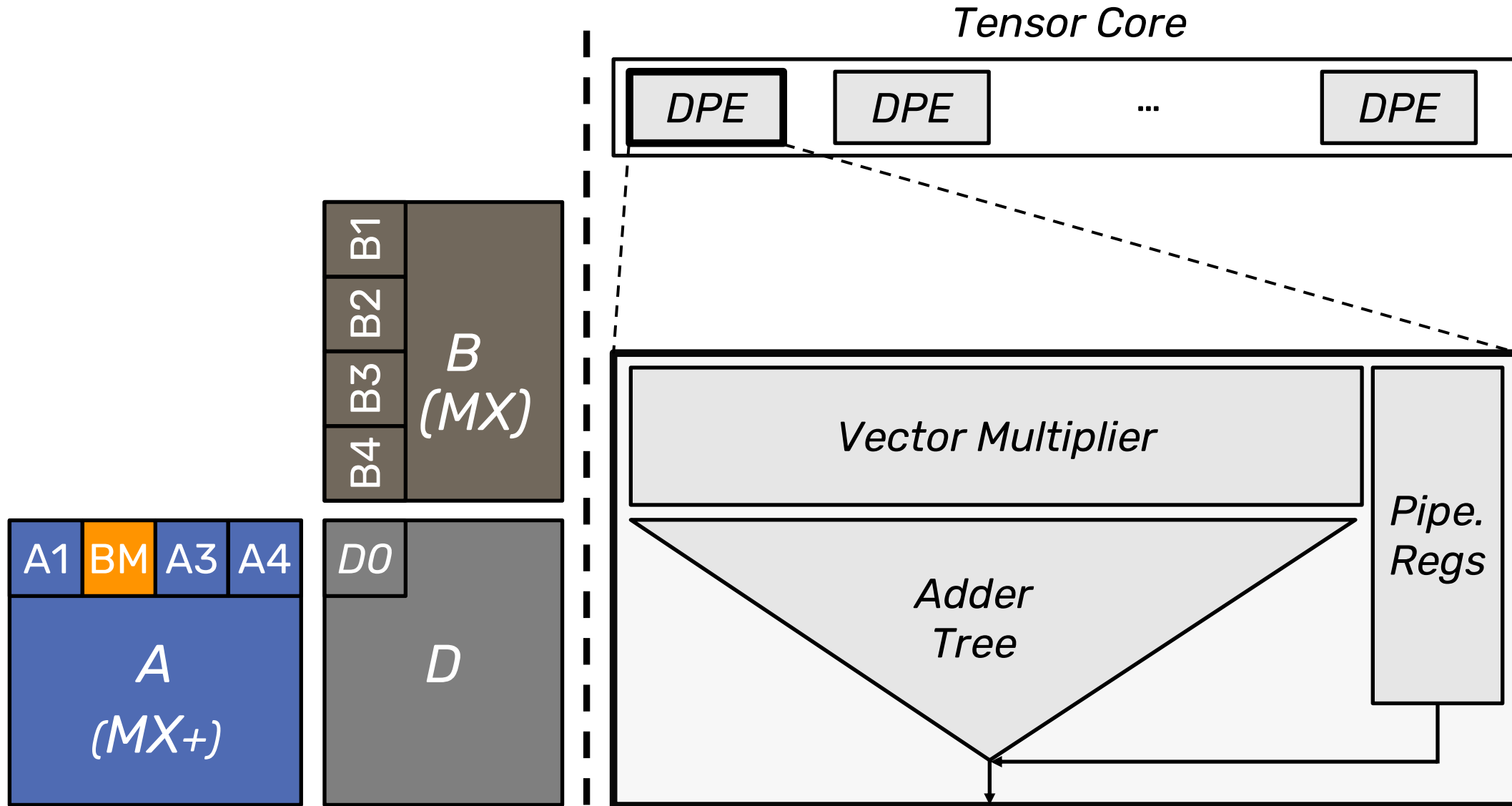


Architectural Support for MX+

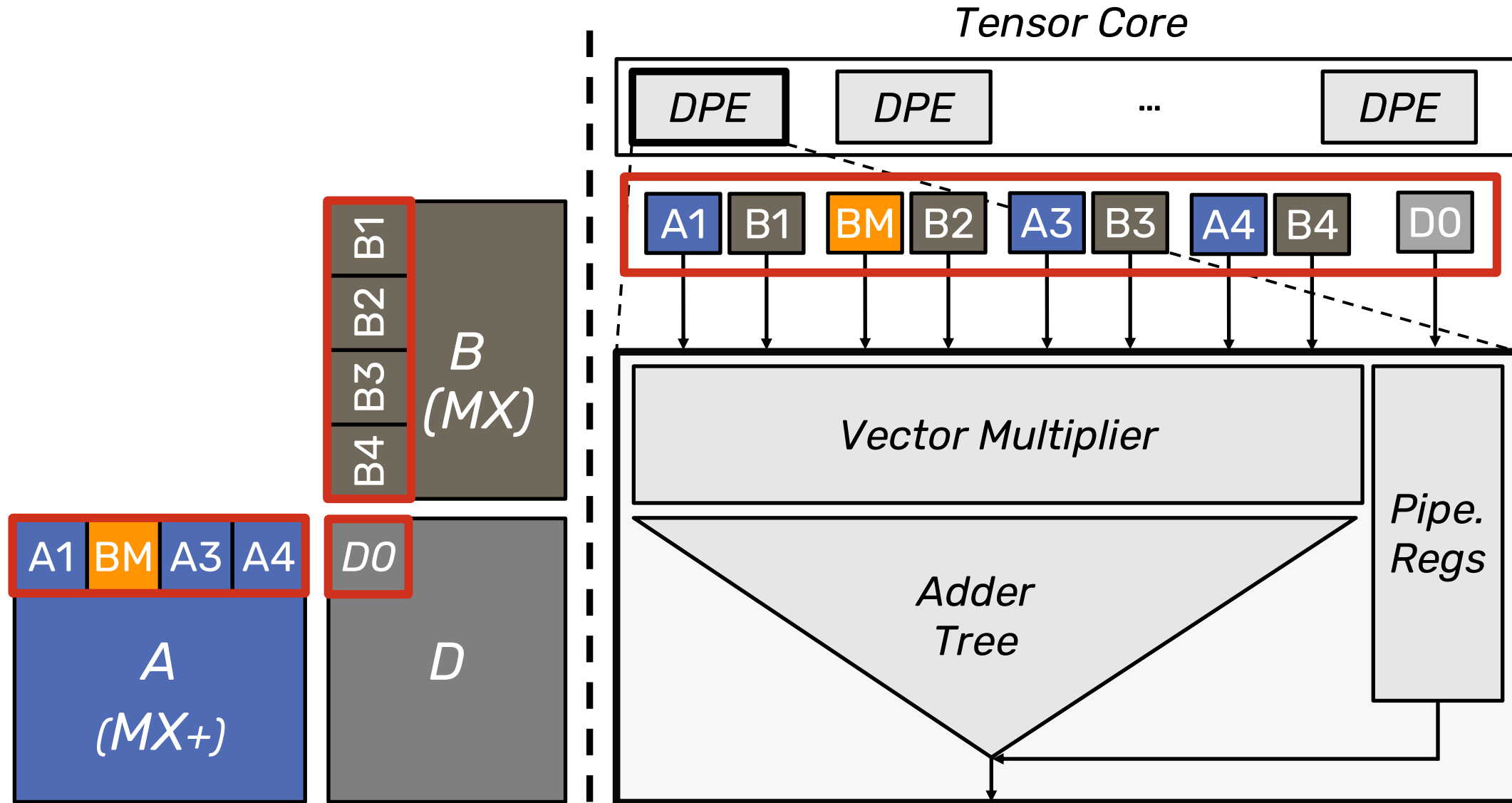
Tensor Core



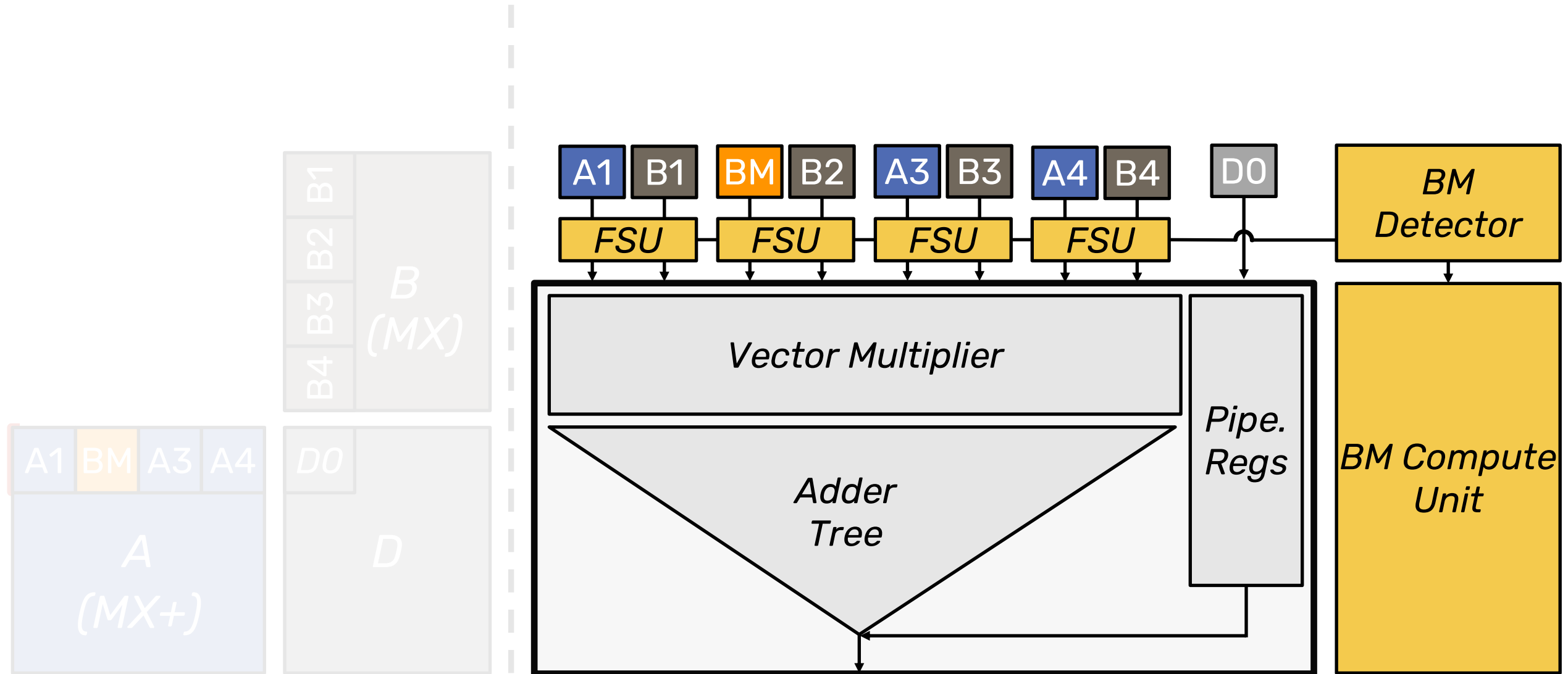
Architectural Support for MX+



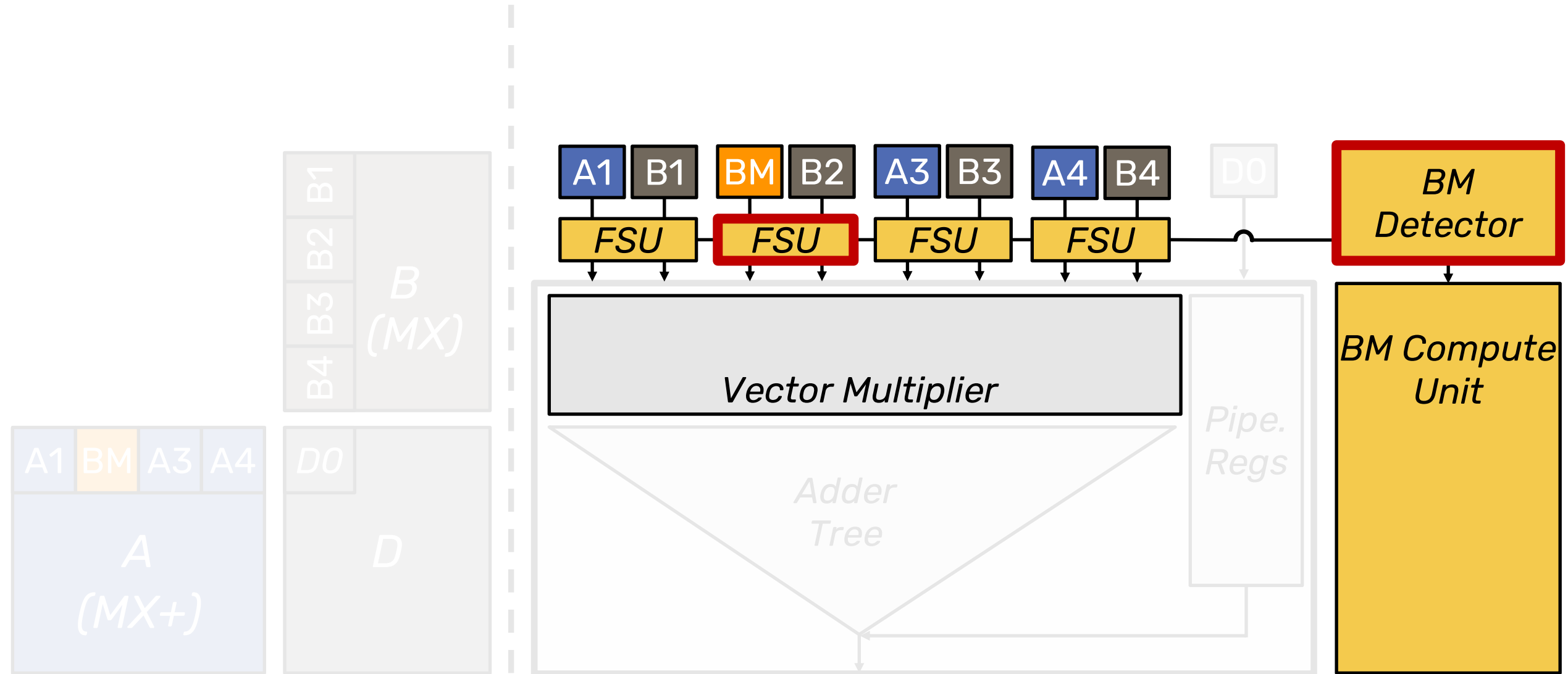
Architectural Support for MX+



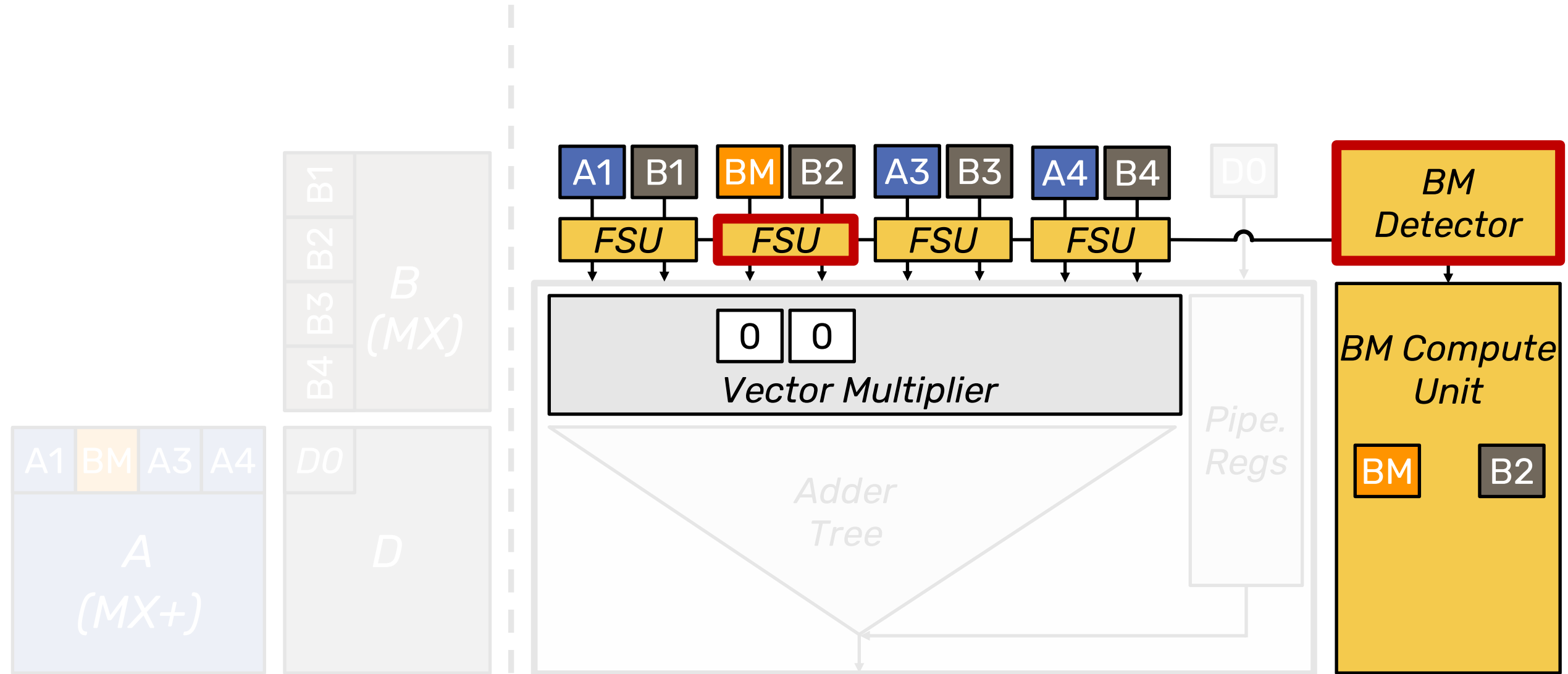
Architectural Support for MX+



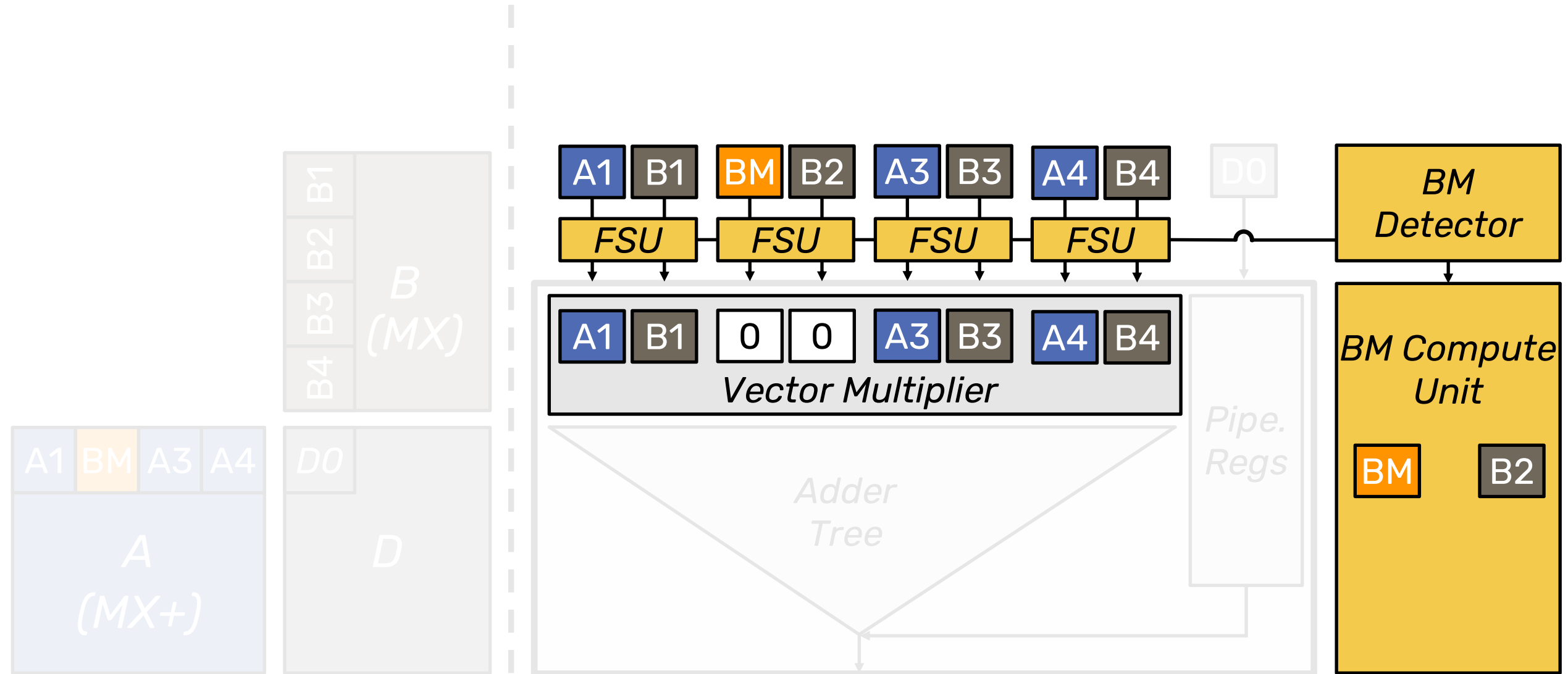
Architectural Support for MX+



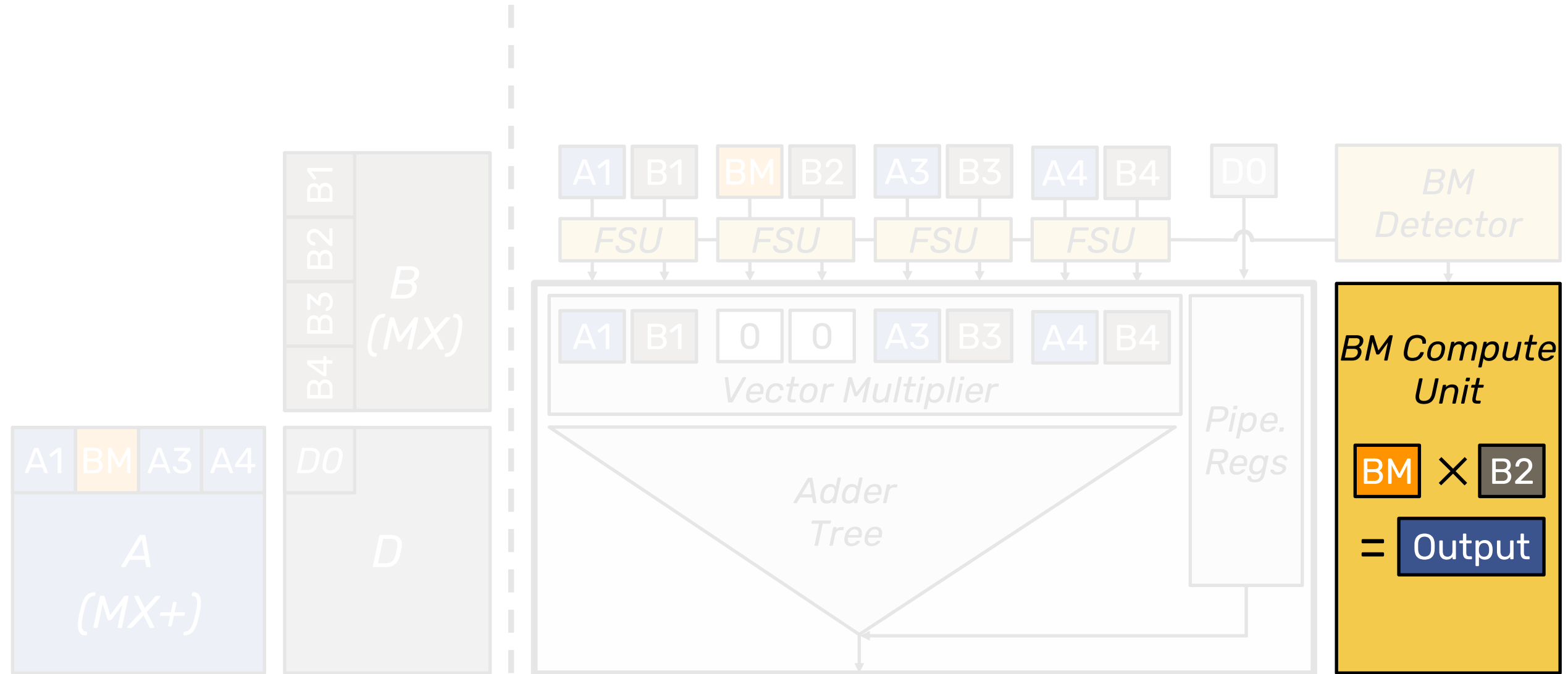
Architectural Support for MX+



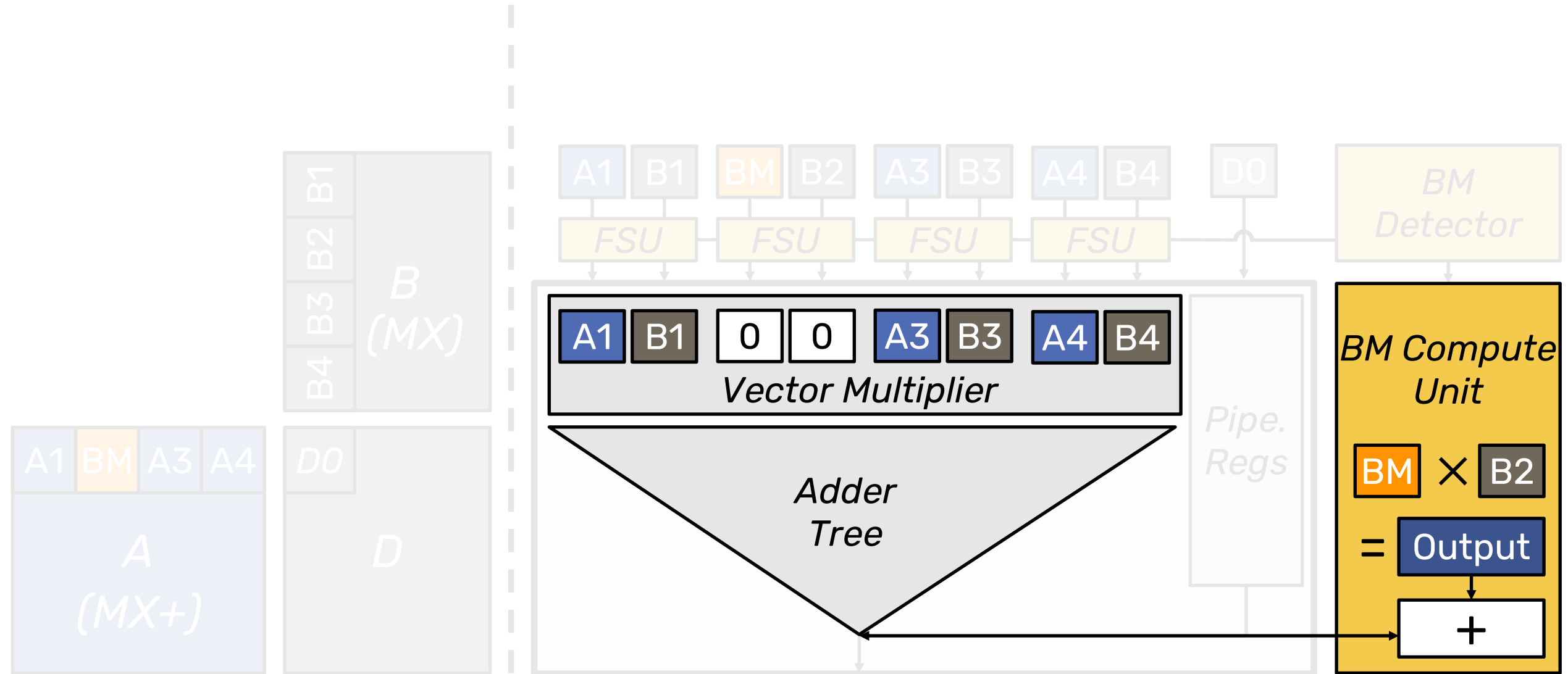
Architectural Support for MX+



Architectural Support for MX+



Architectural Support for MX+



Outline

- **Introduction to MX Formats**
- **Motivation**
 - Implications for Low-bit LLM Serving
 - Analysis on Low-bit MX Formats
- ***MX+*: Cost-Effective and Non-intrusive Extension to MX Formats**
 - Key Insight
 - Software & Hardware Integration on GPUs
- **Evaluation**
- **Conclusion**

Experimental Setup

Models

- OPT, Llama-2 & 3.1, Mistral, etc.

Accuracy

- MX PyTorch emulation library

Performance

- Software: CUTLASS library
- Hardware: NVIDIA RTX 5090 GPU

Workloads

- lm-evaluation-harness

MX+ Schemes

Scheme	Activation	Weight
A-MXFP4+	MXFP4+	MXFP4
MXFP4+	MXFP4+	MXFP4+

MX Schemes

Scheme	Activation	Weight
MXFP4	MXFP4	MXFP4

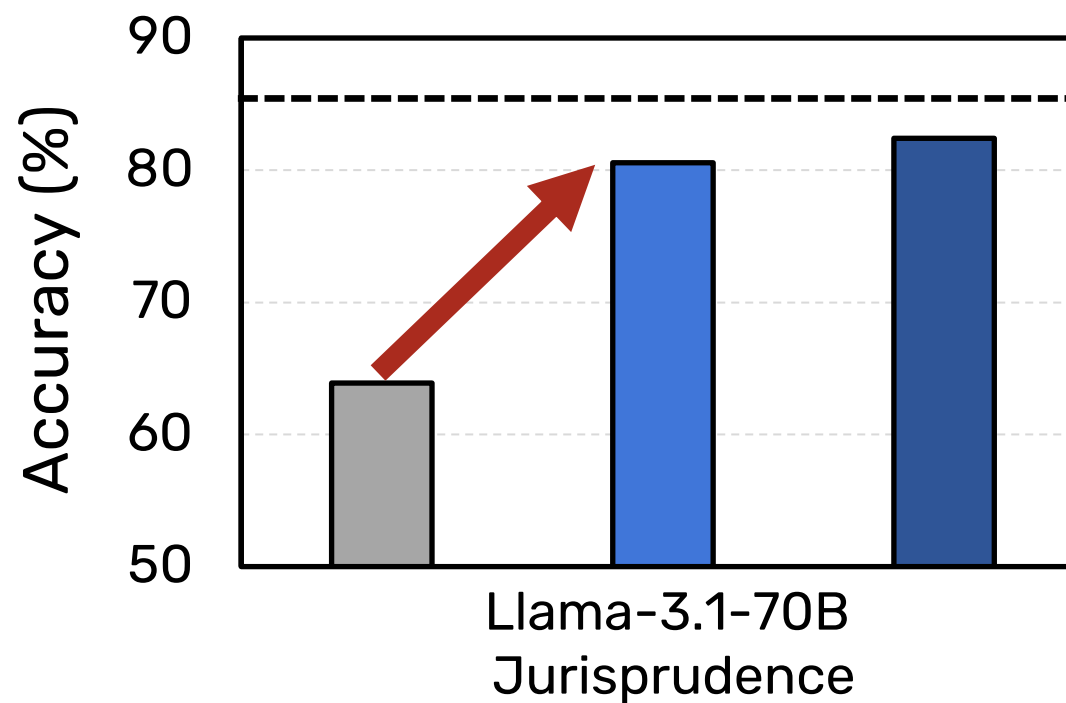
Accuracy

--- BF16

■ MXFP4

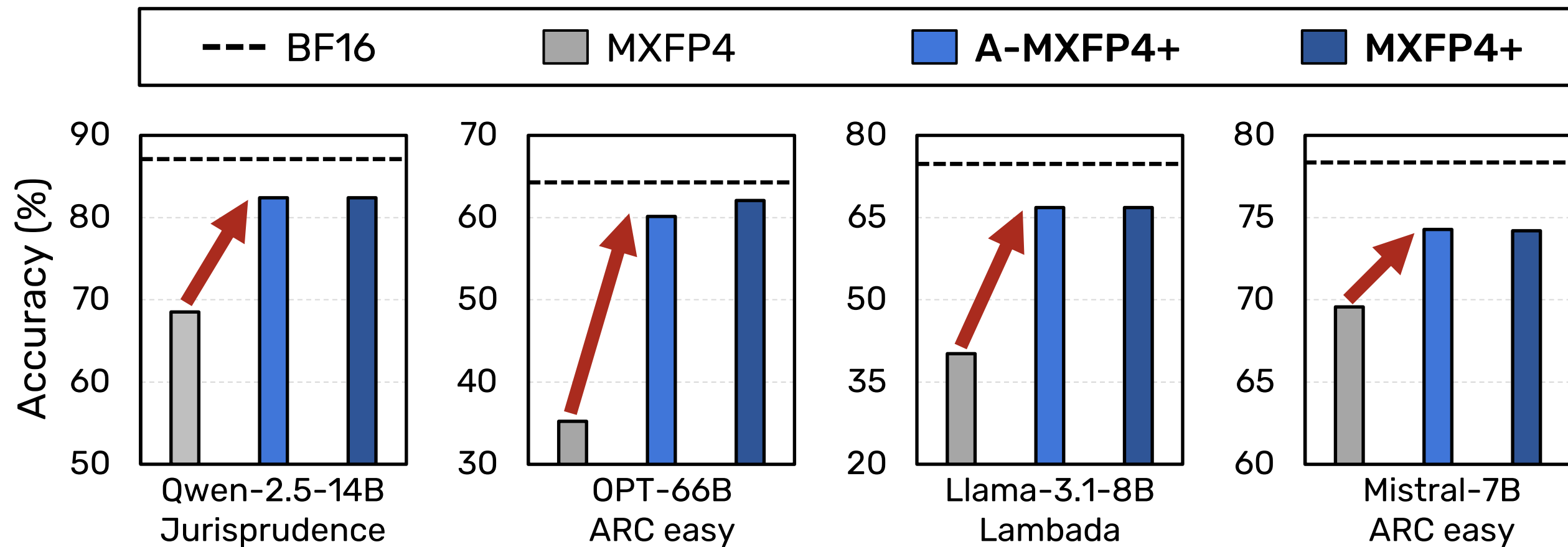
■ A-MXFP4+

■ MXFP4+



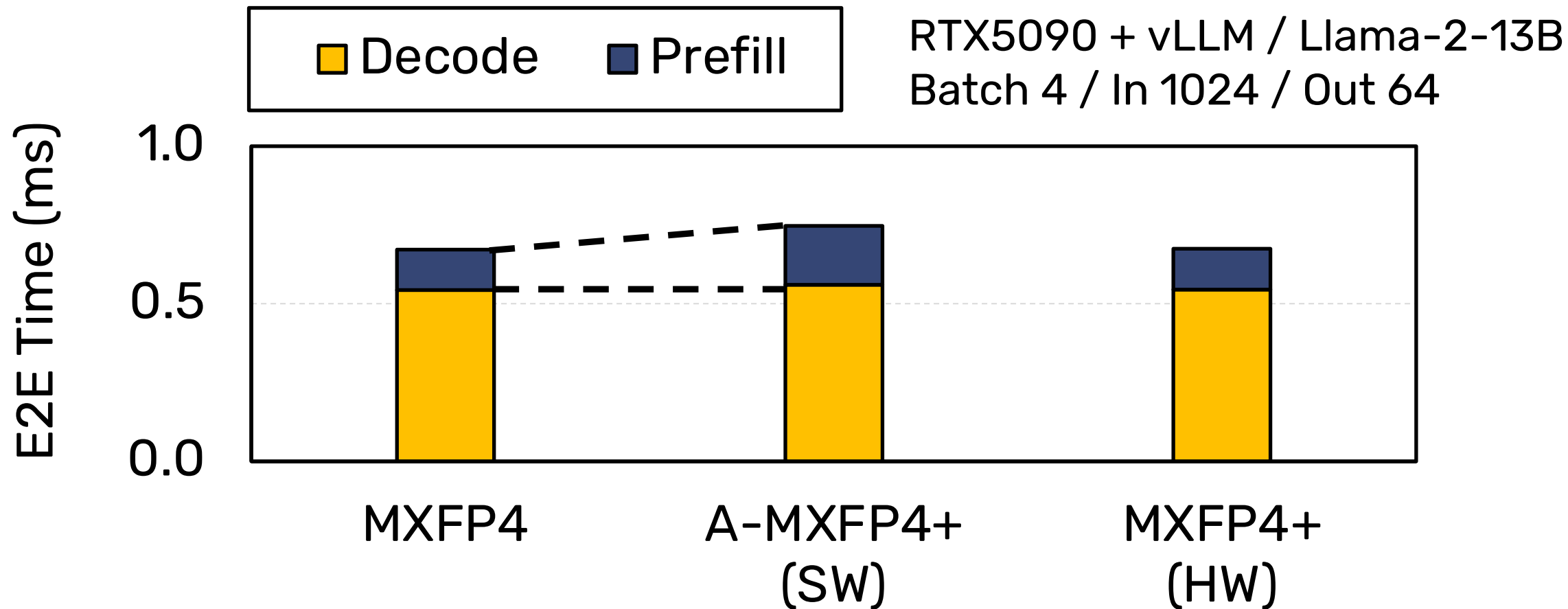
MX+ greatly improves accuracy over MX!

Accuracy



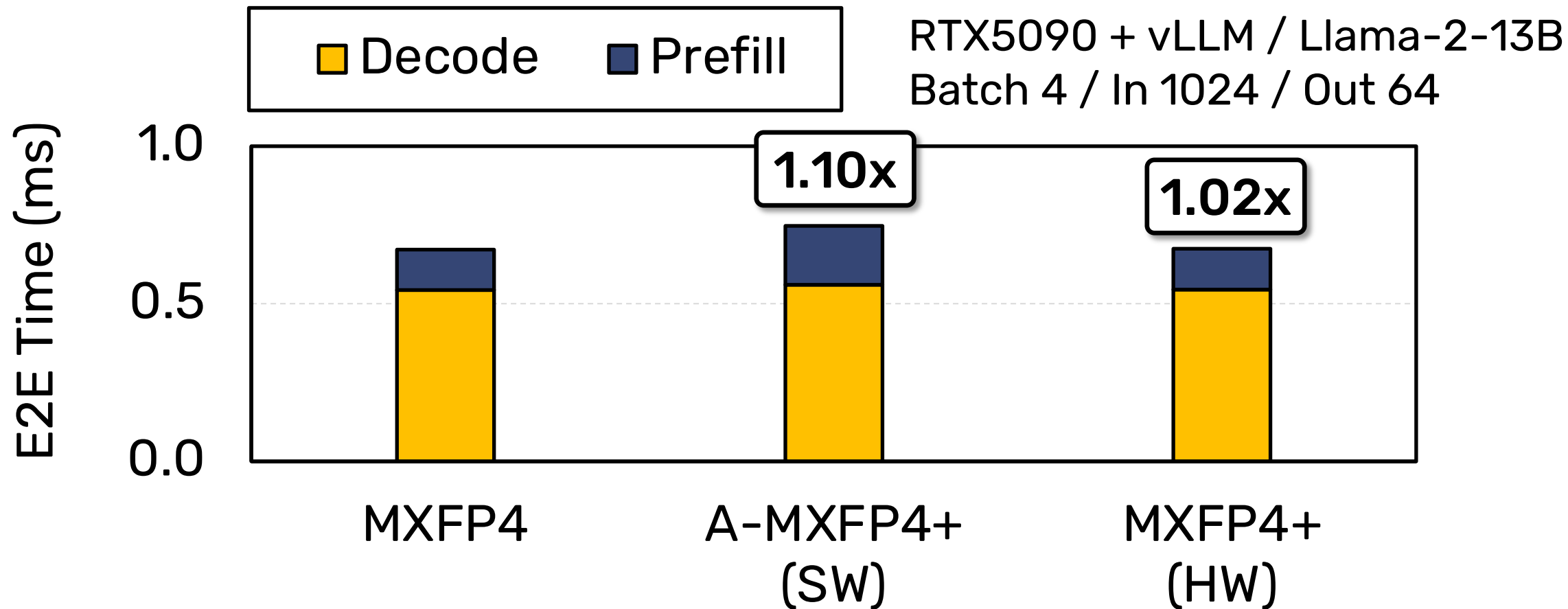
MX+ greatly improves accuracy over MX!

E2E Performance



MX+ maintains comparable performance to MX

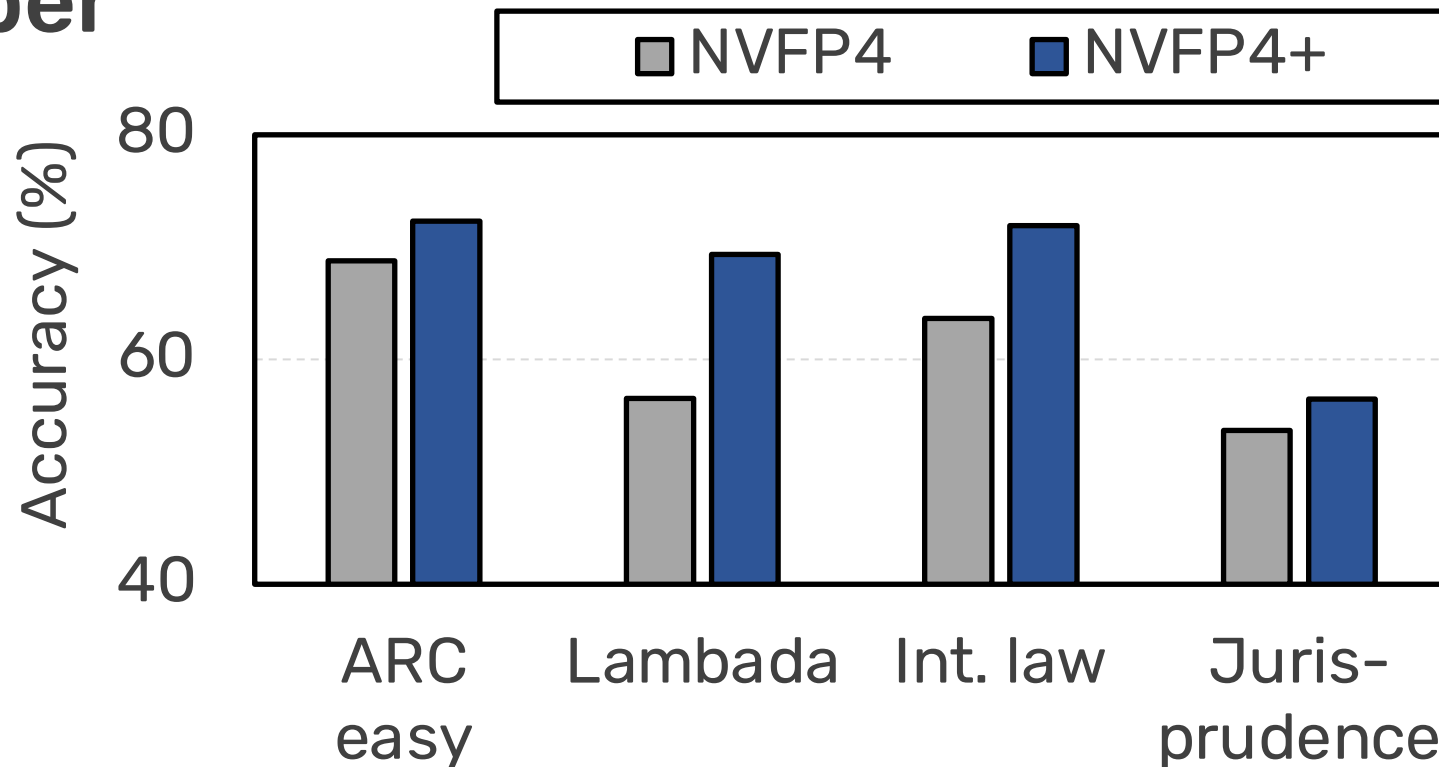
E2E Performance



MX+ maintains comparable performance to MX

More Details in Our Paper

- Application of MX+ for NVFP4
 - And comparison with MXFP4+



- Comparison with A8W4 and other quantization works
- Handling multiple outliers in a block
- Use of reserved bits
- MX+ for other DNN workloads (e.g., ViT and CNN)

Conclusion

Problem

- Low-Bit MX Format shows subpar accuracy

Solution: **MX+**, a cost-effective and non-intrusive extension to MX

- Higher precision for BMs by repurposing exponent bits as extended mantissa
- Integration to existing software and hardware system

Result

- MX+ improves **accuracy by up to +42.15%** compared to the MX format
- **Comparable performance** under software and hardware integration

Thank You!

MX+

Pushing the Limits of
Microscaling Formats for Efficient
Large Language Model Serving

Jungi Lee (jung.lee@snu.ac.kr)

MICRO'25 | October 2025

