

Adana: Accelerating Large Language Models via Addaptive Nonuniform Asymmetric Quantization

Xinkuang Geng¹, Siting Liu², Hui Wang¹, Leibo Liu³, Jie Han⁴, Honglan Jiang¹

¹ School of Integrated Circuits, Shanghai Jiao Tong University, Shanghai, China

² School of Information Science and Technology, ShanghaiTech University, Shanghai, China

³ School of Integrated Circuits, Tsinghua University, Beijing, China

⁴ Department of Electrical and Computer Engineering, University of Alberta, Alberta, Canada

xinkuang@sjtu.edu.cn, liust@shanghaitech.edu.cn, ihuiwang@sjtu.edu.cn, liulb@tsinghua.edu.cn, jhan8@ualberta.ca, honglan@sjtu.edu.cn



From FP16 Model to Low-Bit Quantization

- **LLM quantization** reduces memory footprint and inference cost [1].

FP16 LLM Inference

$$\begin{array}{c} \mathbf{X} \\ \hline \begin{array}{|c|} \hline 3.7 \\ \hline -0.1 \\ \hline 1.2 \\ \hline 2.8 \\ \hline \end{array} \end{array} \times \begin{array}{c} \mathbf{W} \\ \hline \begin{array}{|c|c|c|c|} \hline 1.2 & -1.1 & 0.2 & 2.2 \\ \hline 3.1 & 0.8 & -3.3 & -5.1 \\ \hline 3.5 & -2.1 & 2.9 & -1.8 \\ \hline 4.9 & 5.2 & 0.6 & 2.0 \\ \hline \end{array} \end{array} = \begin{array}{c} \mathbf{y} \\ \hline \begin{array}{|c|} \hline \\ \hline \\ \hline \\ \hline \\ \hline \end{array} \end{array}$$

- FP16 weights / activations / KV
- FP16 matrix multiplication (MM)

Low-Bit LLM Inference

$$\begin{array}{c} \mathbf{X} \\ \hline \begin{array}{|c|} \hline 4 \\ \hline 0 \\ \hline 1 \\ \hline 3 \\ \hline \end{array} \end{array} \times \begin{array}{c} \mathbf{W} \\ \hline \begin{array}{|c|c|c|c|} \hline 1 & -1 & 0 & 2 \\ \hline 3 & 1 & -3 & -5 \\ \hline 4 & -2 & 3 & -2 \\ \hline 5 & 5 & 1 & 2 \\ \hline \end{array} \end{array} = \begin{array}{c} \mathbf{y} \\ \hline \begin{array}{|c|} \hline \\ \hline \\ \hline \\ \hline \\ \hline \end{array} \end{array}$$

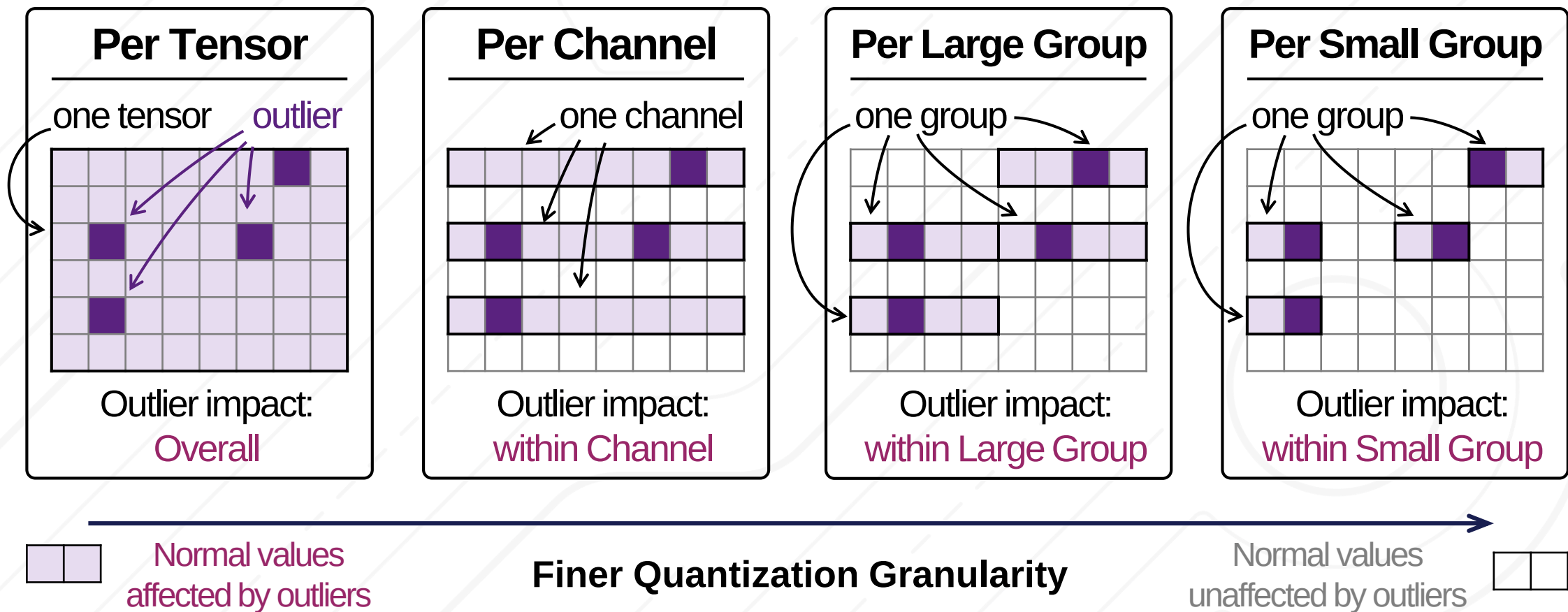
- 3-bit / 4-bit weights / activations / KV
- 3-bit / 4-bit matrix multiplication

[1] Guo, Cong, et al. "Olive: Accelerating large language models via hardware-friendly outlier-victim pair quantization." ISCA. 2023.



Toward Finer Quantization Granularity

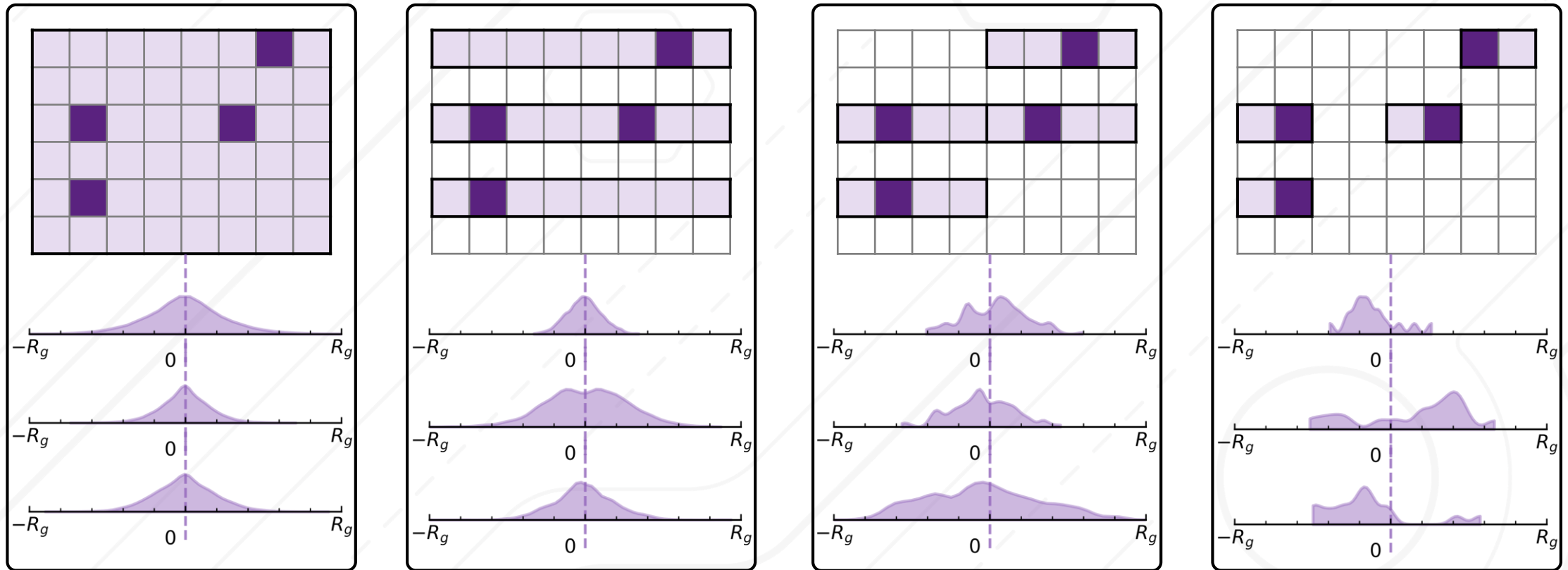
- **Smaller quantization groups** preserve more normal values.





Motivation & Problem Setting

Smaller Groups, Greater Distribution Variability

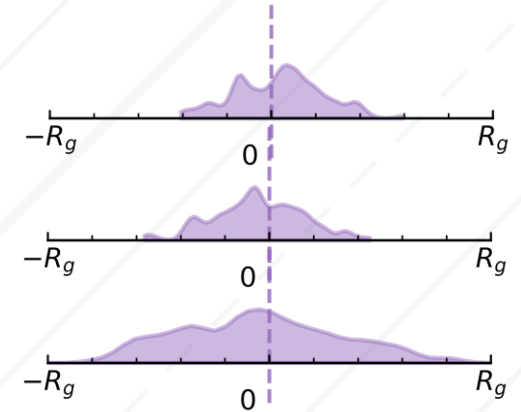


Finer Quantization Granularity

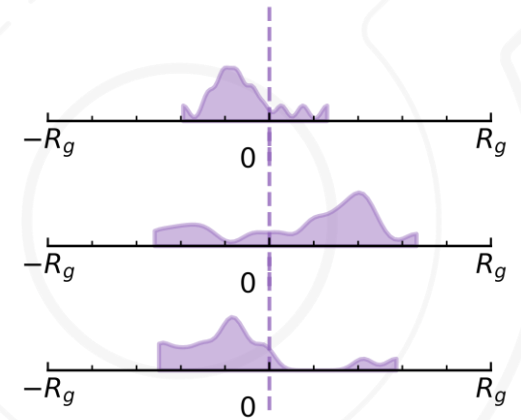


Gaps in Existing Quantization Methods

- **Insufficient exploitation of fine-grained group-wise quantization**
 - Smaller groups reduce quantization error, but existing numeric types [2-4] do not fully adapt to group sizes as small as 16.
- **Trade-off between accuracy and efficiency**
 - MSE-based search [2,4] improves accuracy but adds high hardware cost; simpler mapping [3] is efficient but less accurate.



Large Groups



Small Groups

[2] Guo, Cong, et al. "Ant: Exploiting adaptive numerical data type for low-bit deep neural network quantization." MICRO. IEEE, 2022

[3] Hu, Weiming, et al. "M-ANT: Efficient low-bit group quantization for LLMs via mathematically adaptive numerical type." HPCA. IEEE, 2025.

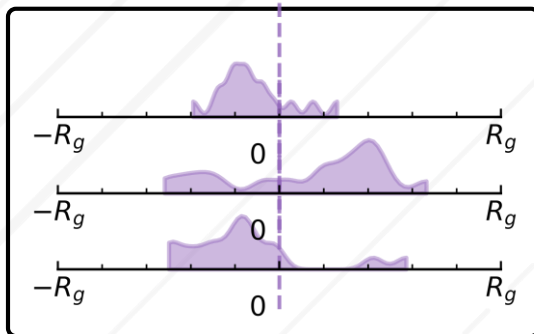
[4] Chen, Yuzong, et al. "Bitmod: Bit-serial mixture-of-datatype llm acceleration." HPCA. IEEE, 2025.



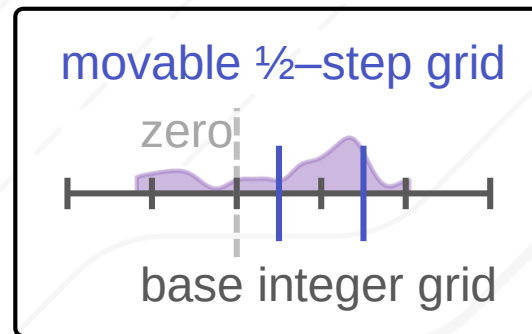
Our Approach

Algorithm side

- Adaptive nonuniform asymmetric (**Adana**) numeric type
- Efficient **online** activation / **KV** quantization



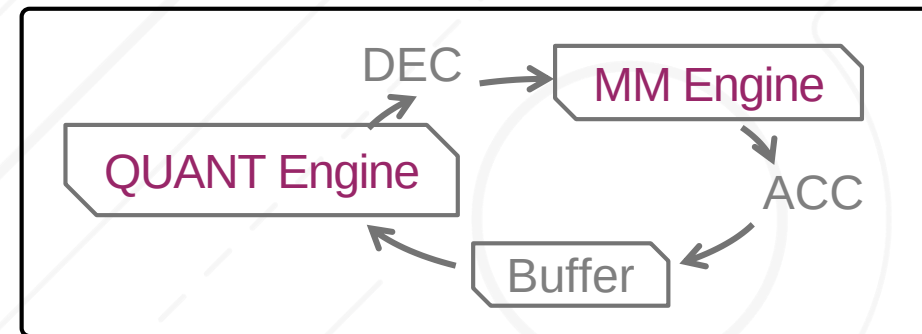
Fit Data Distributions



Efficient Quantization

Architecture side

- Low-bit group-wise **MM** microarchitecture and dataflow
- Hardware-efficient **QUANT** engine

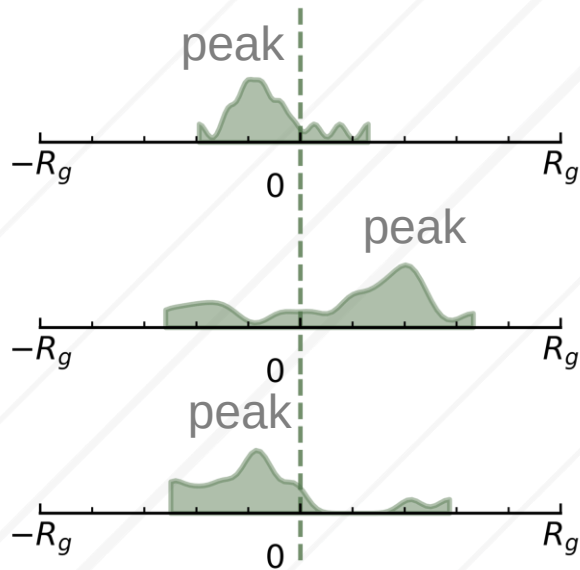


Efficient Computation



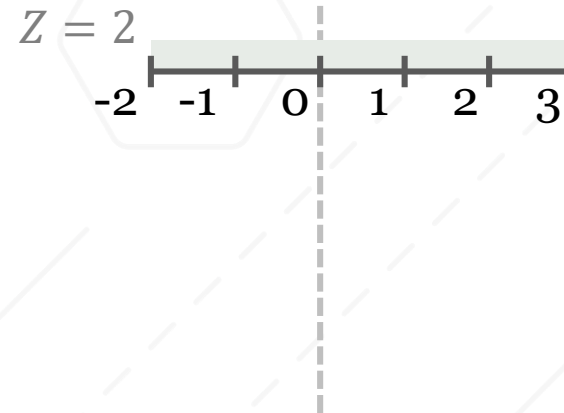
Numeric Type Design

diverse peak positions

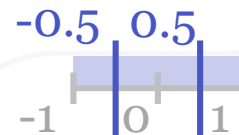


asymmetric data ranges

base integer grid



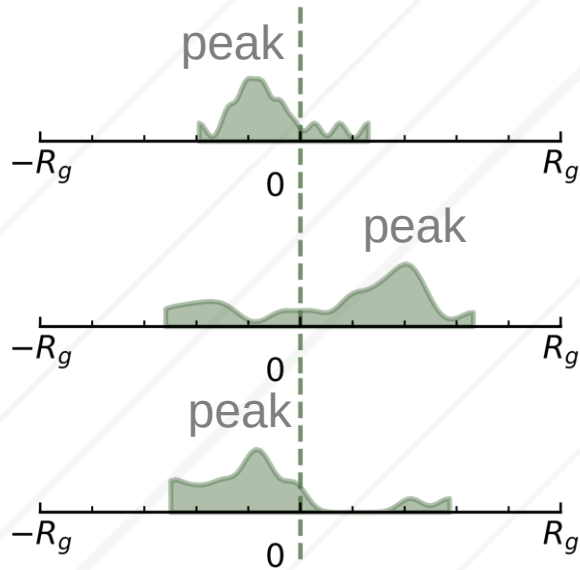
movable $\frac{1}{2}$ -step grid





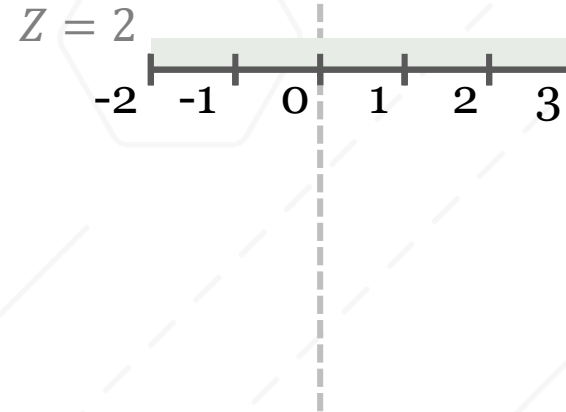
Numeric Type Design

diverse peak positions

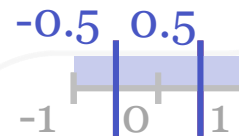


asymmetric data ranges

base integer grid



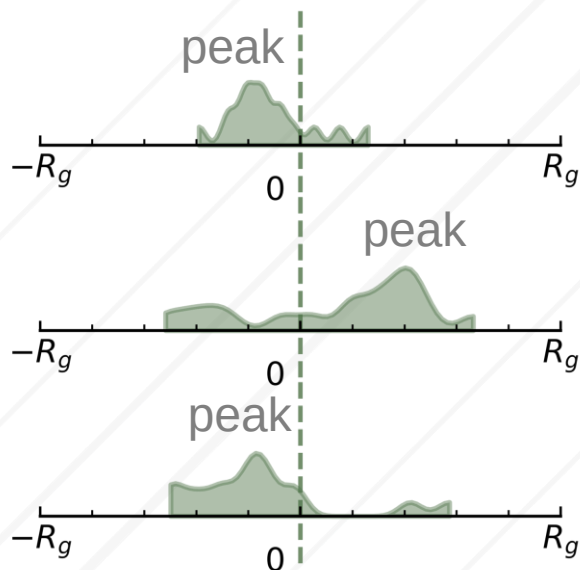
movable $\frac{1}{2}$ -step grid





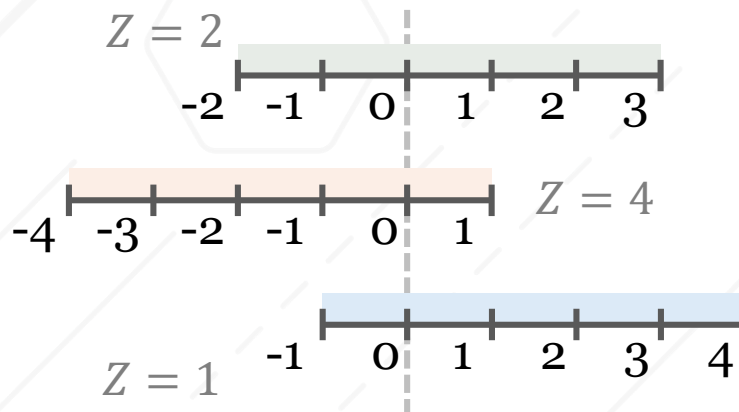
Numeric Type Design

diverse peak positions

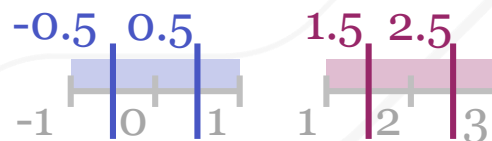


asymmetric data ranges

base integer grid



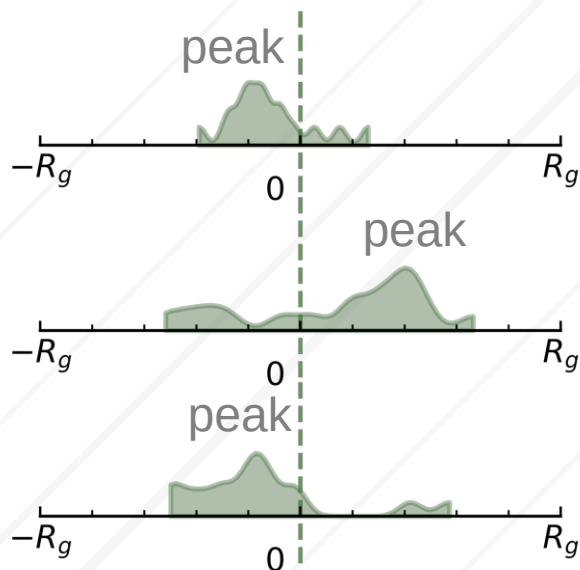
movable $\frac{1}{2}$ -step grid





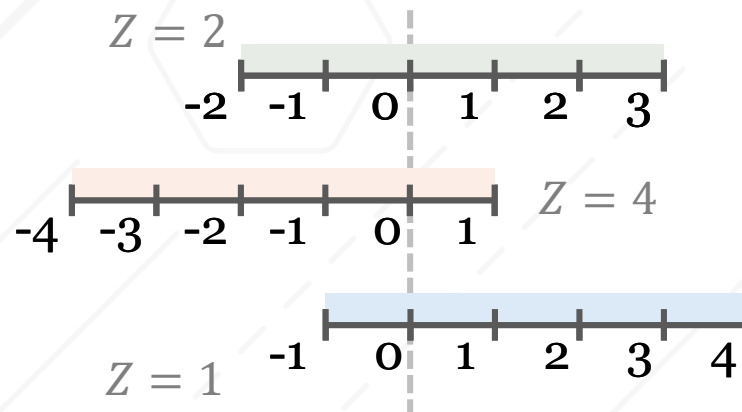
Numeric Type Design

diverse peak positions

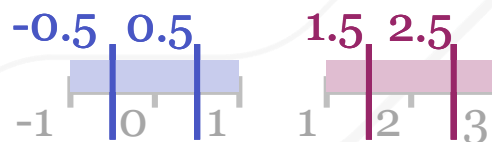


asymmetric data ranges

base integer grid



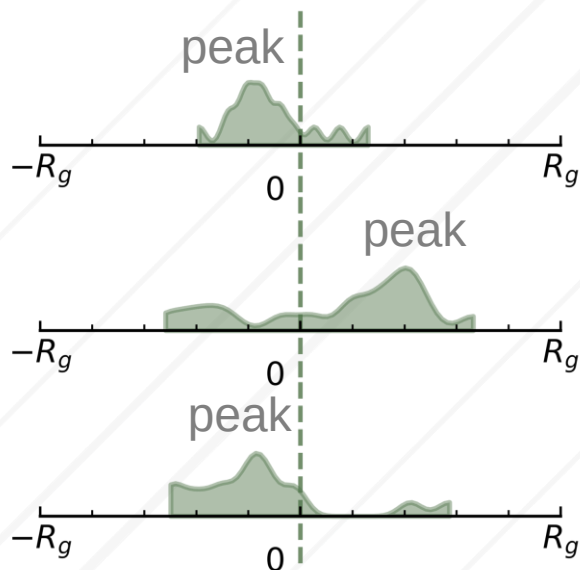
movable $\frac{1}{2}$ -step grid





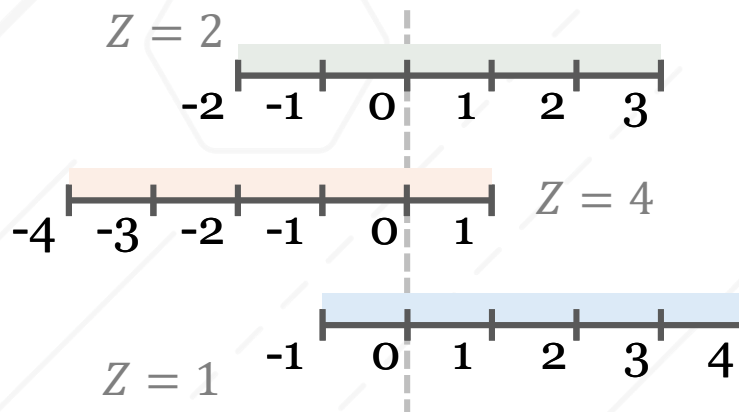
Numeric Type Design

diverse peak positions

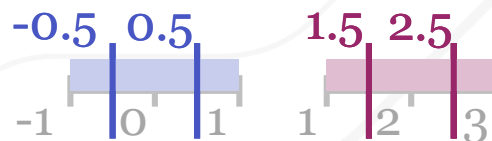


asymmetric data ranges

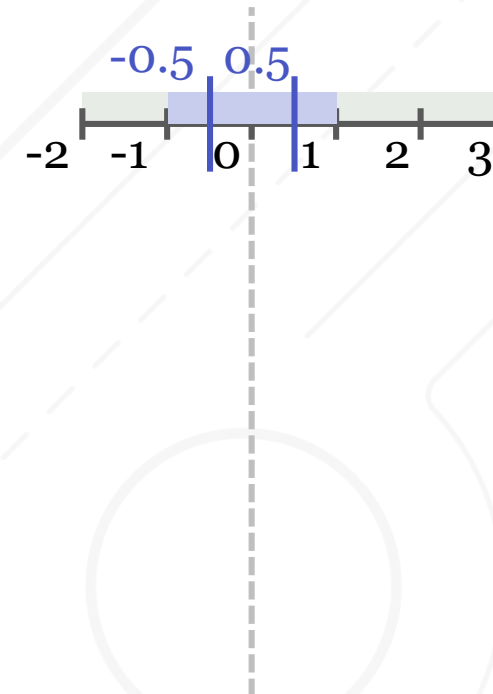
base integer grid



movable $\frac{1}{2}$ -step grid



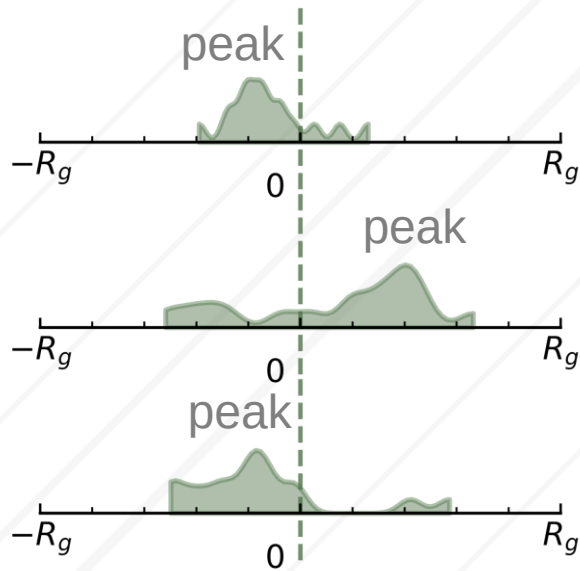
combine to form Adana grids





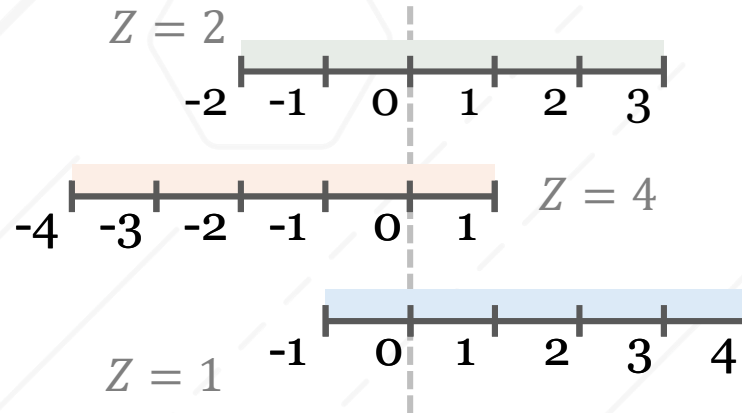
Numeric Type Design

diverse peak positions

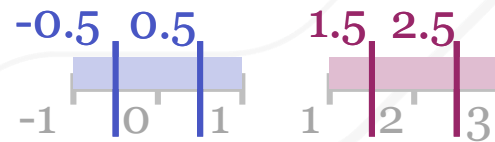


asymmetric data ranges

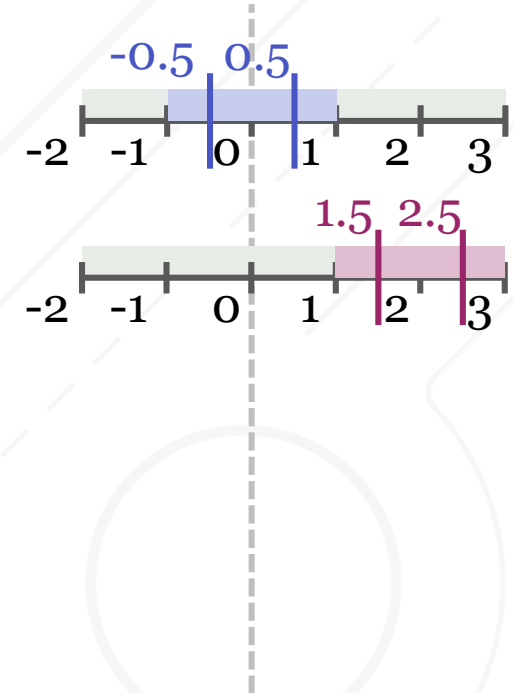
base integer grid



movable $\frac{1}{2}$ -step grid



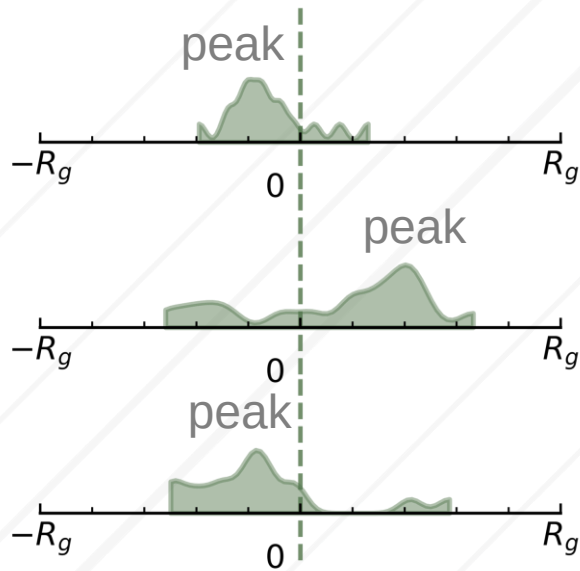
combine to form Adana grids





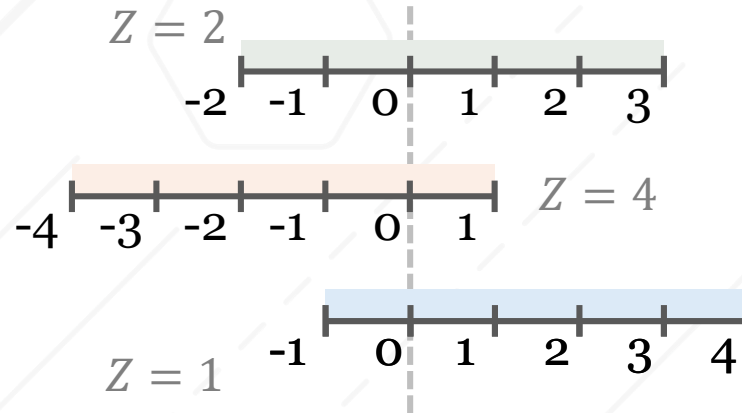
Numeric Type Design

diverse peak positions

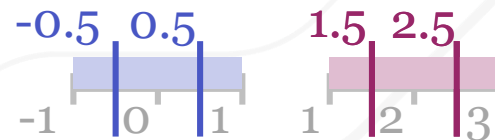


asymmetric data ranges

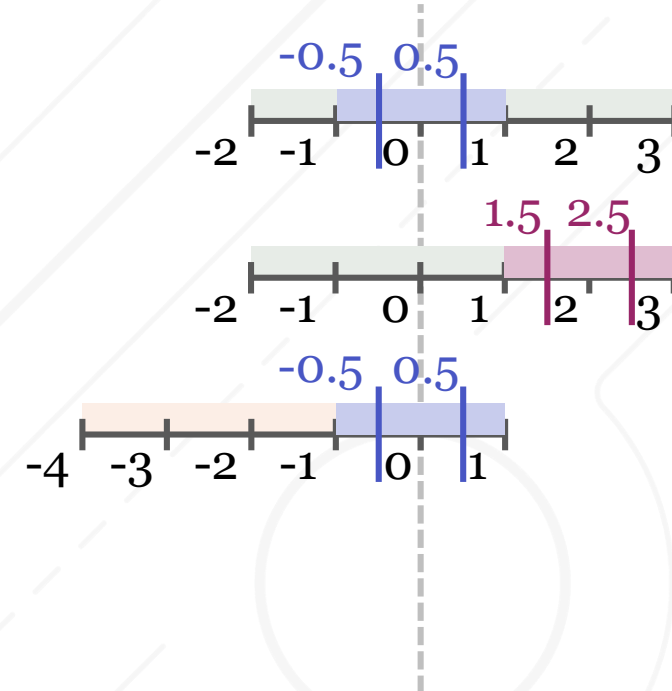
base integer grid



movable $\frac{1}{2}$ -step grid



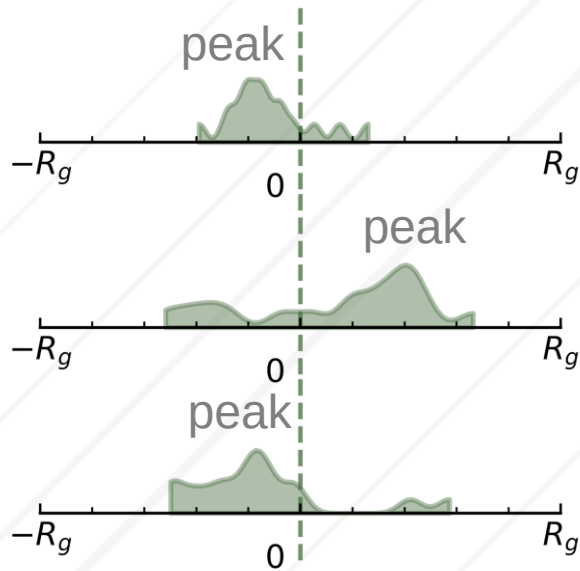
combine to form Adana grids





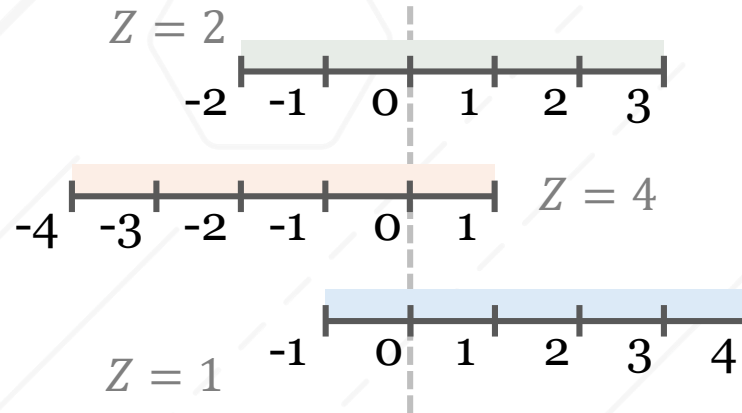
Numeric Type Design

diverse peak positions

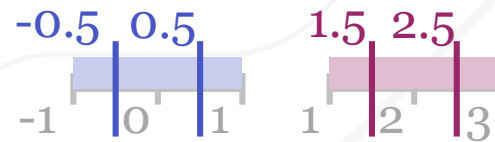


asymmetric data ranges

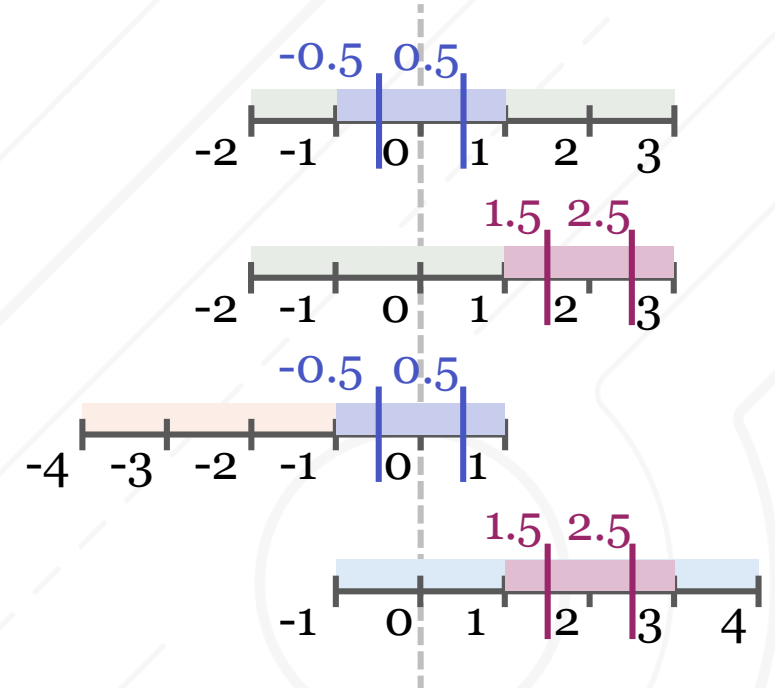
base integer grid



movable 1/2-step grid



combine to form Adana grids



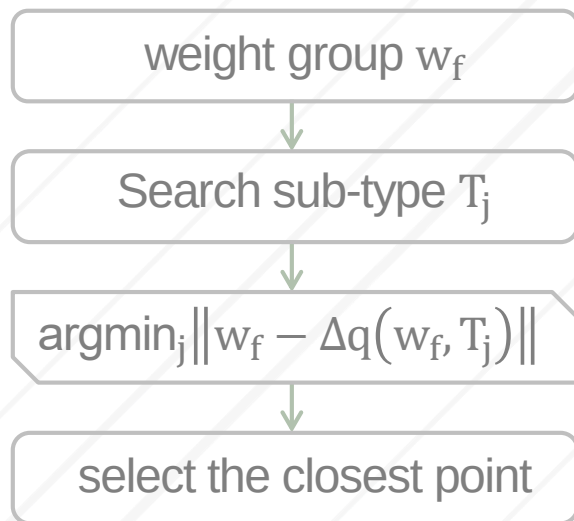
Zero point adapts to asymmetric data ranges.

Movable 1/2-step grid fits diverse peak positions.



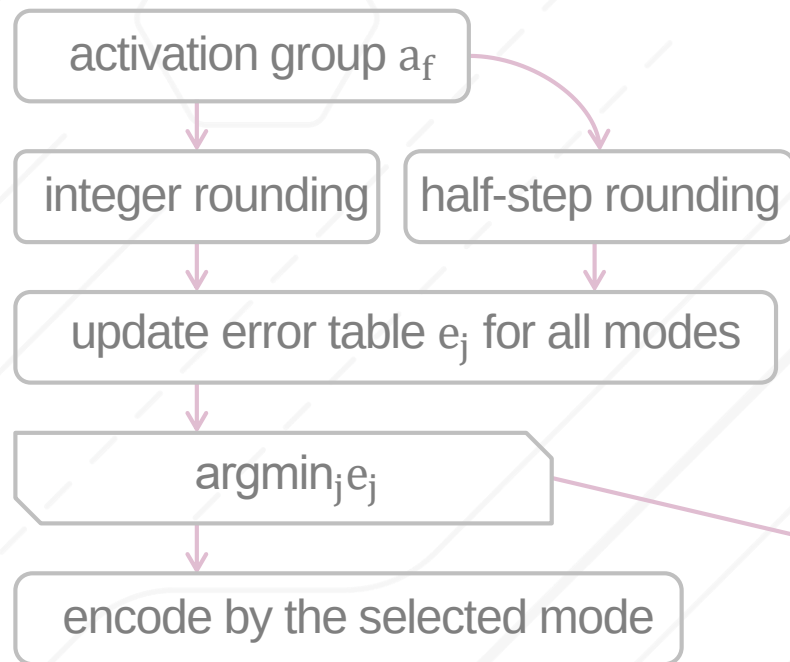
Quantization Framework

offline weight quantization



GPU-friendly / done once

online activation quantization



why it is efficient

No floating-point MSE computation

Shared base grid across sub-types

Only rounding & error counting

Approximate metric:
count mismatches to the union grid
instead of full floating-point MSE.

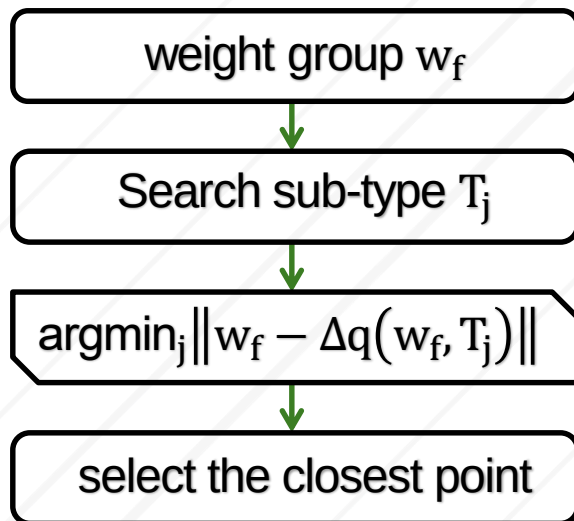
MSE-based search

lightweight rounding & error counting



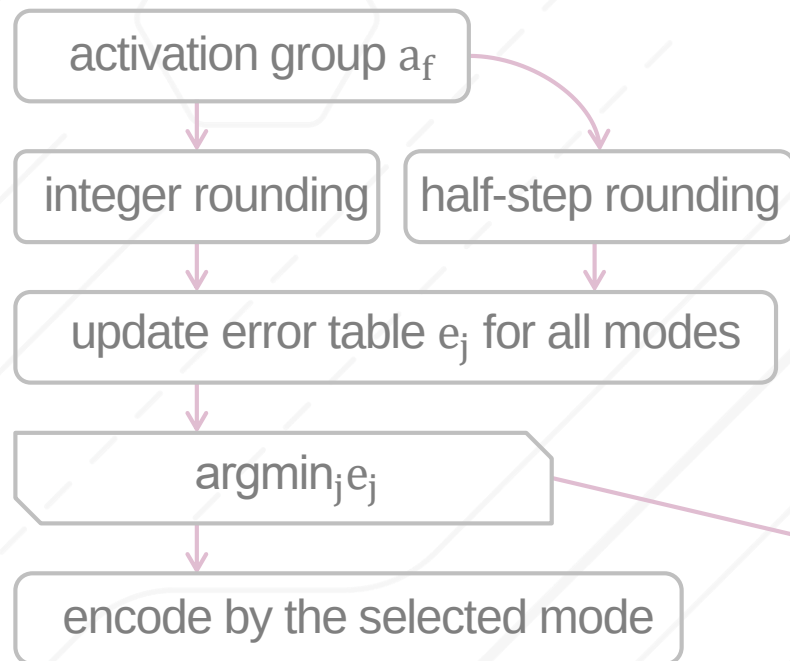
Quantization Framework

offline weight quantization



GPU-friendly / done once

online activation quantization



why it is efficient

No floating-point MSE computation

Shared base grid across sub-types

Only rounding & error counting

Approximate metric:
count mismatches to the union grid
instead of full floating-point MSE.

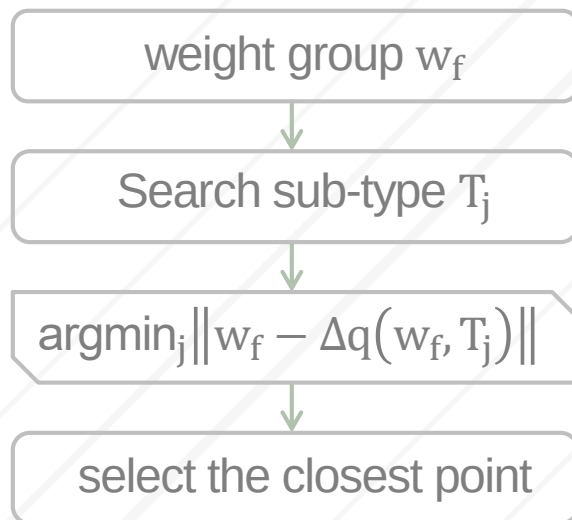
MSE-based search

lightweight rounding & error counting



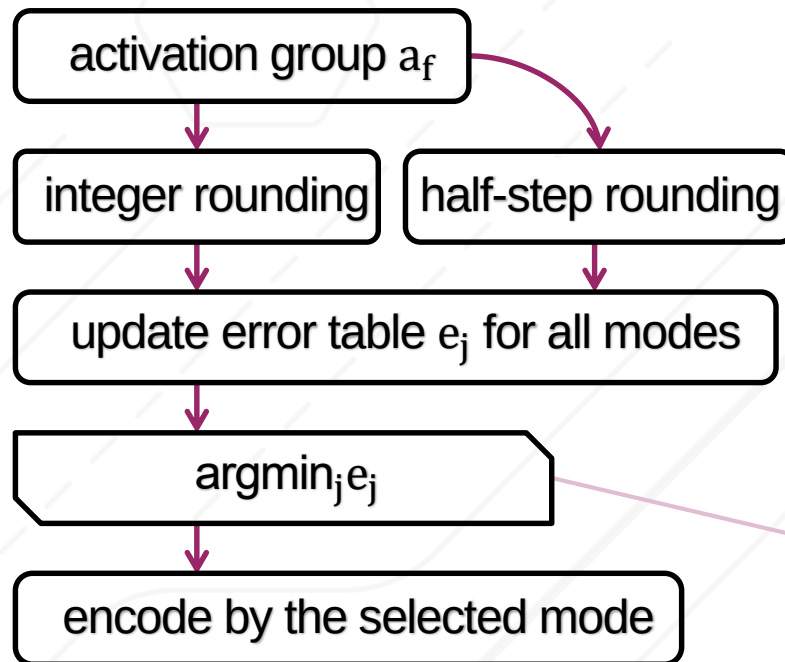
Quantization Framework

offline weight quantization



GPU-friendly / done once

online activation quantization



why it is efficient

No floating-point MSE computation

Shared base grid across sub-types

Only rounding & error counting

Approximate metric:
count mismatches to the union grid
instead of full floating-point MSE.

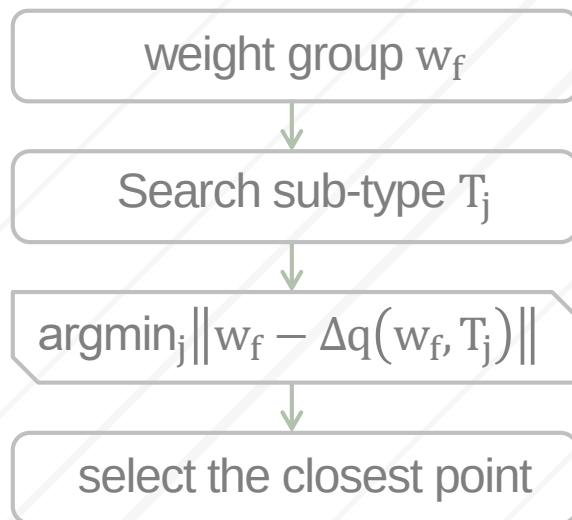
MSE-based search

lightweight rounding & error counting



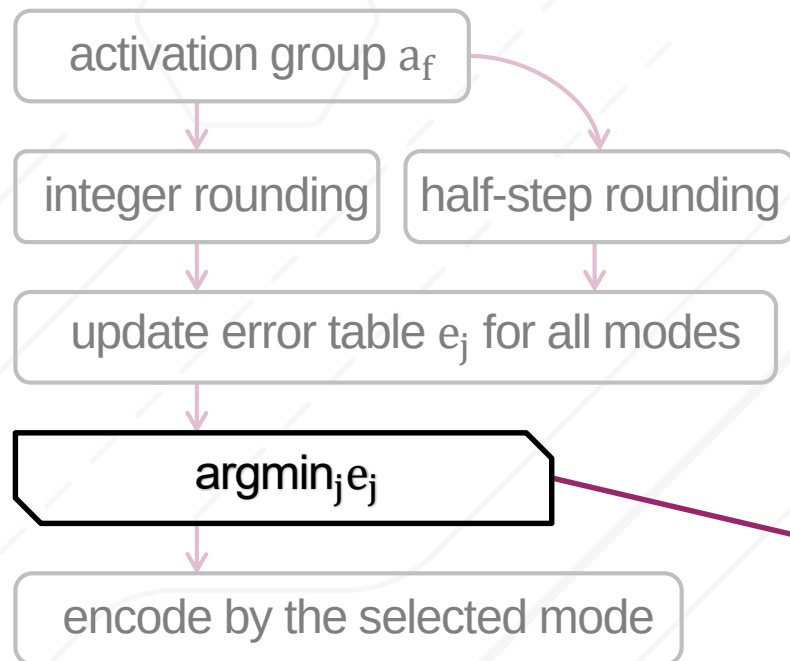
Quantization Framework

offline weight quantization



GPU-friendly / done once

online activation quantization



why it is efficient

No floating-point MSE computation

Shared base grid across sub-types

Only rounding & error counting

Approximate metric:
count mismatches to the union grid
instead of full floating-point MSE.

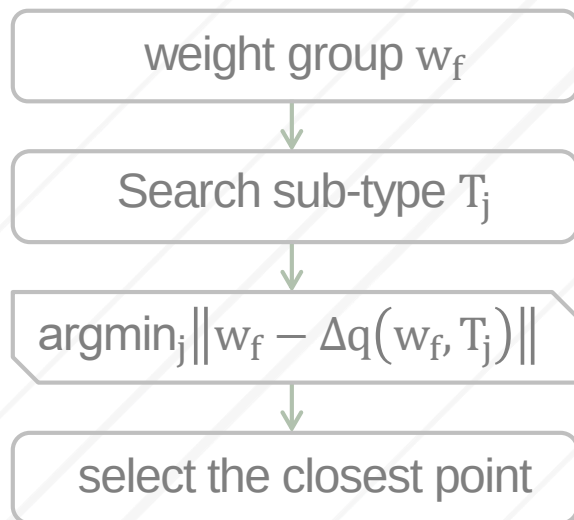
MSE-based search

lightweight rounding & error counting



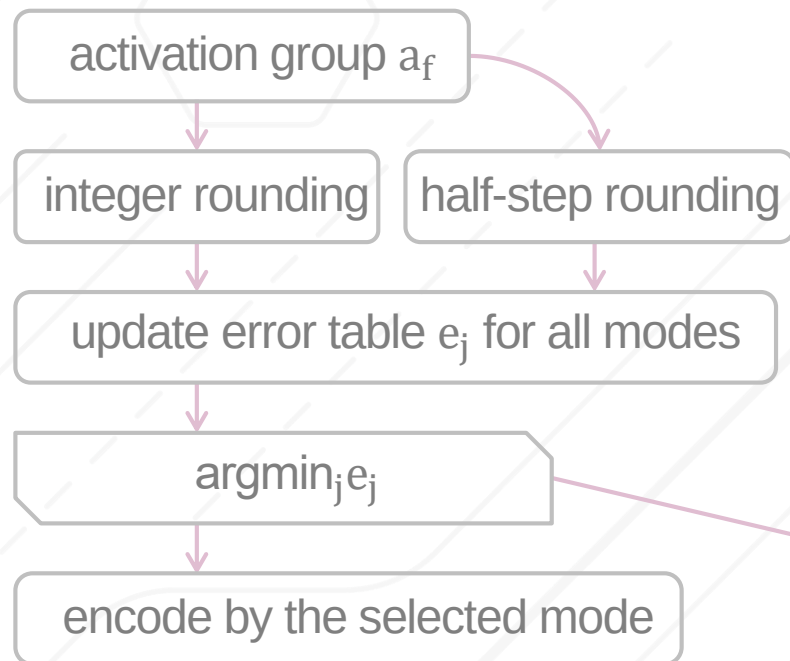
Quantization Framework

offline weight quantization



GPU-friendly / done once

online activation quantization



why it is efficient

No floating-point MSE computation

Shared base grid across sub-types

Only rounding & error counting

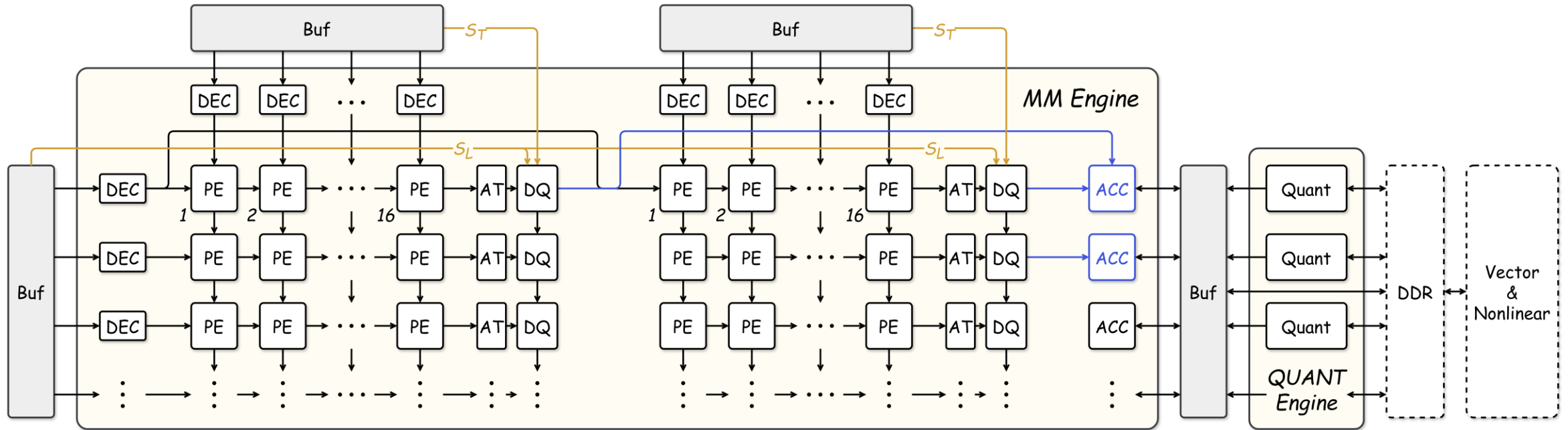
Approximate metric:
count mismatches to the union grid
instead of full floating-point MSE.

MSE-based search

lightweight rounding & error counting

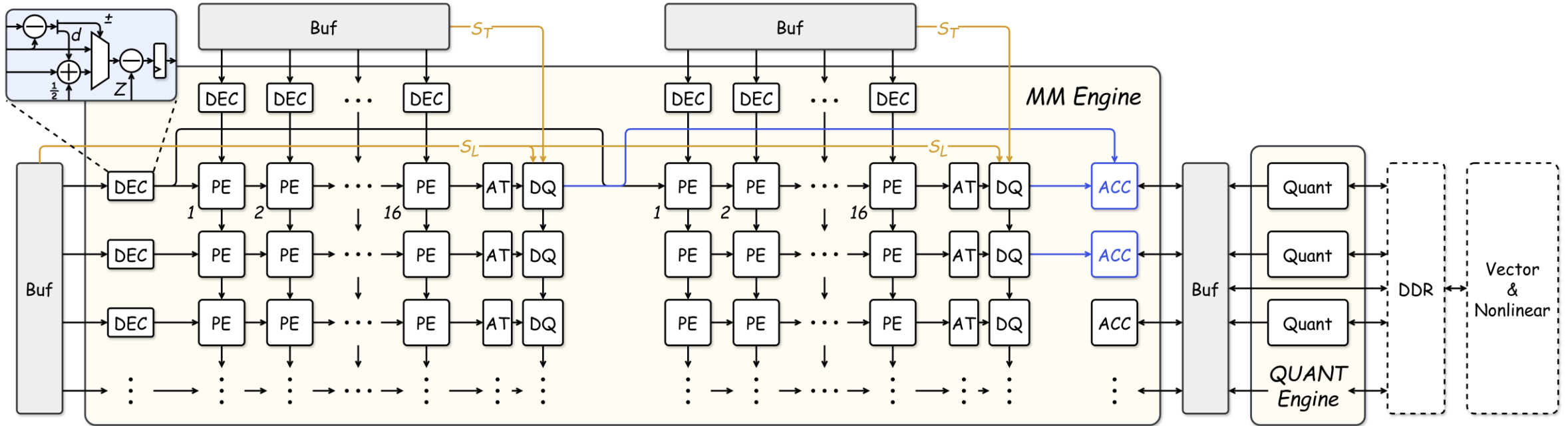


Accelerator Implementation





Accelerator Implementation

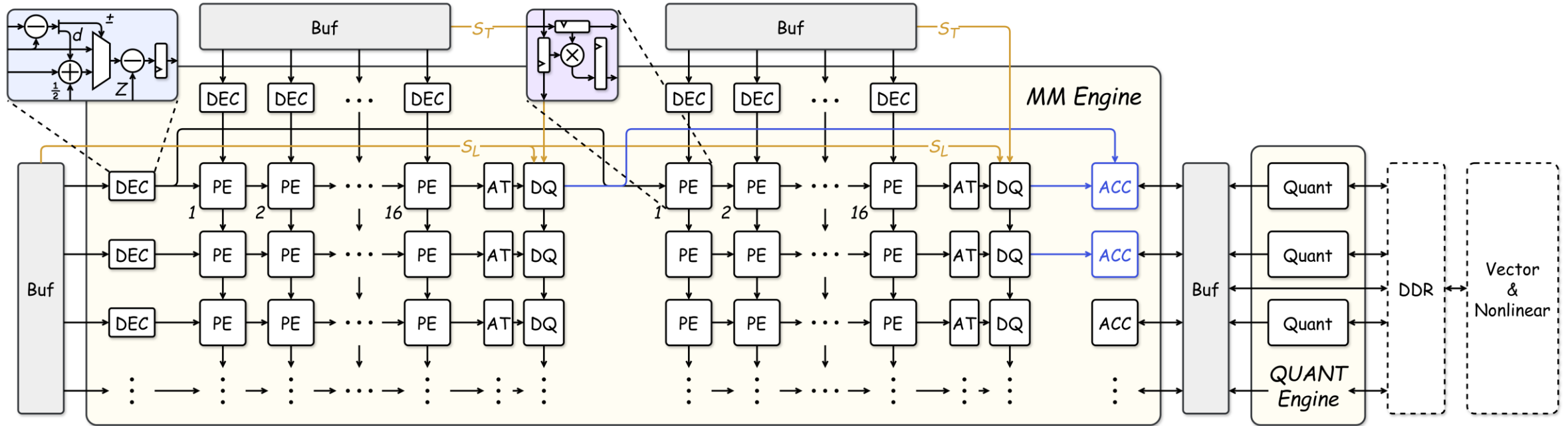


64 + 64 Decoders



Algorithm & Architecture

Accelerator Implementation



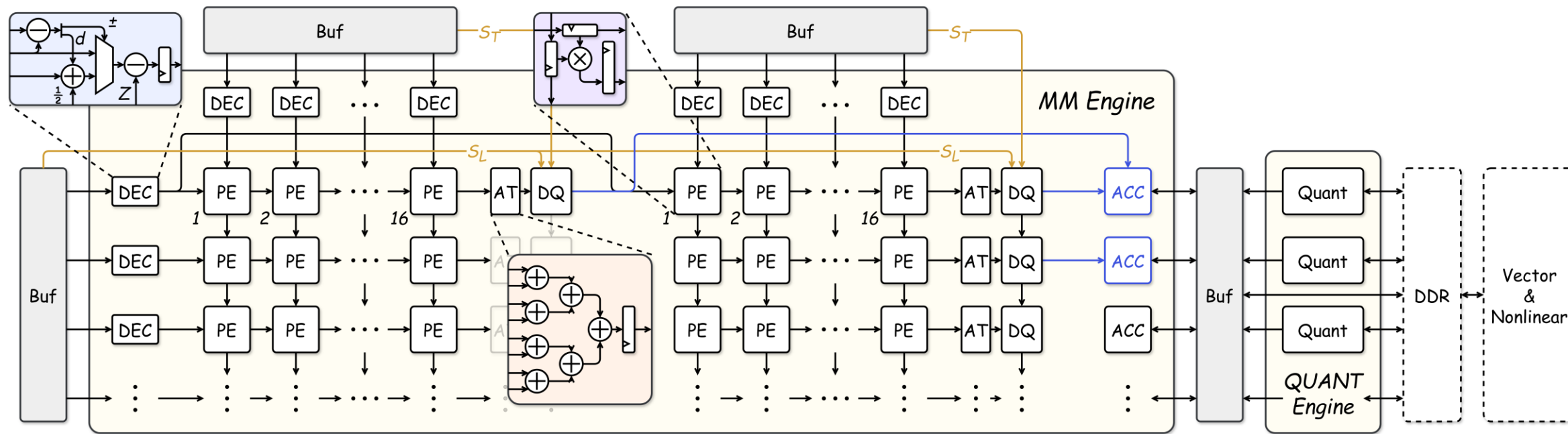
64 + 64 Decoders

$4 \times 64 \times 16$ Processing Elements

Low-bit INT Multiplication



Accelerator Implementation



64 + 64 Decoders

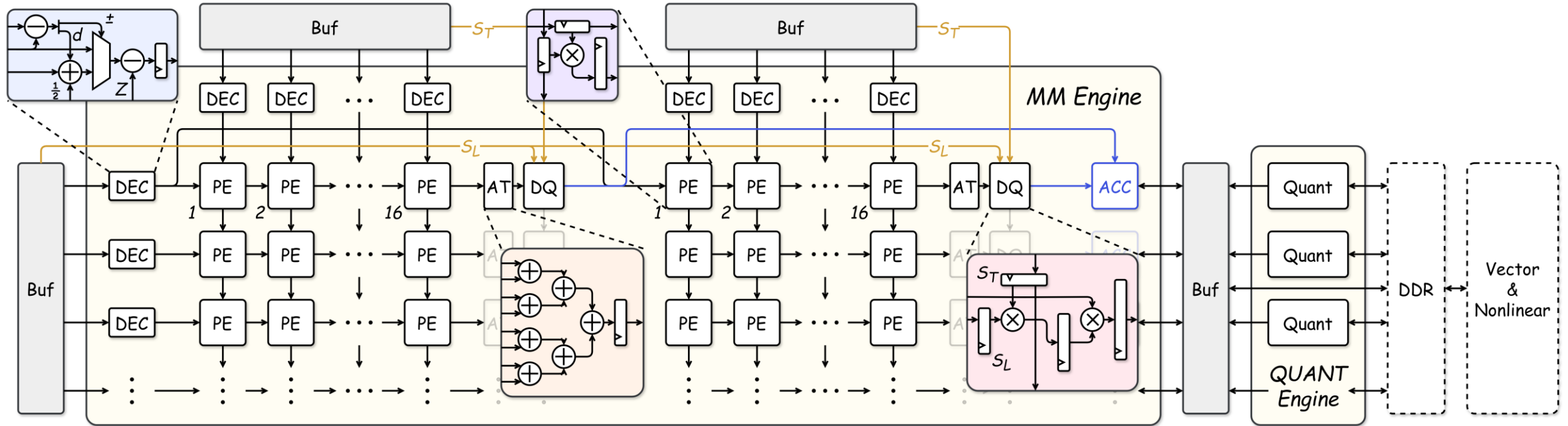
$4 \times 64 \times 16$ Processing Elements

Low-bit INT Multiplication

4×64 Adder Trees



Accelerator Implementation



64 + 64 Decoders

$4 \times 64 \times 16$ Processing Elements

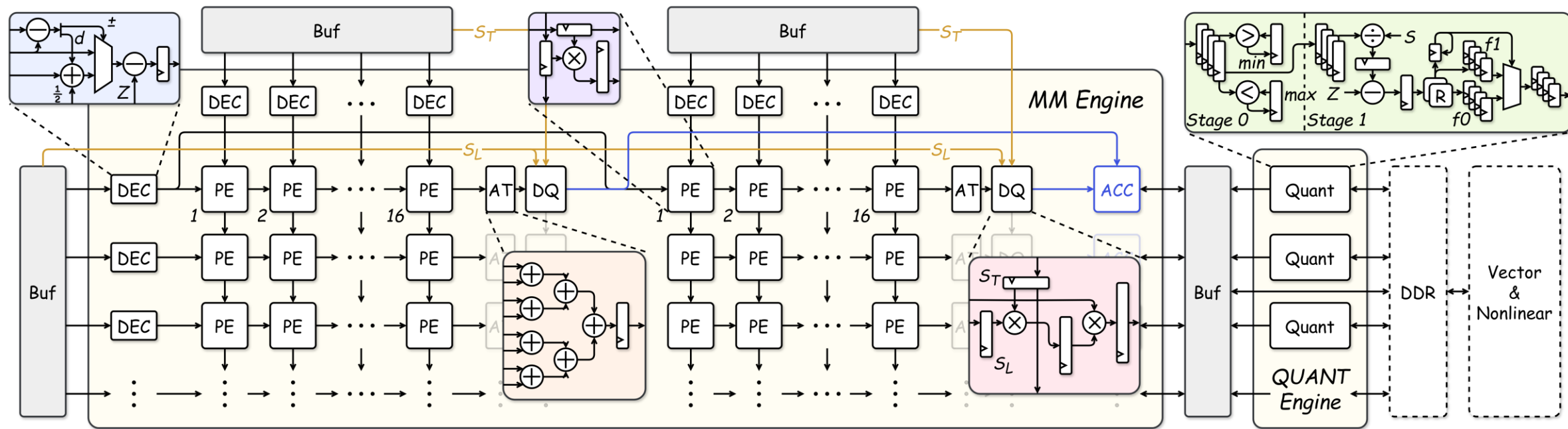
Low-bit INT Multiplication

4×64 Adder Trees

4×64 De-quantizers



Accelerator Implementation



64 + 64 Decoders

$4 \times 64 \times 16$ Processing Elements

Low-bit INT Multiplication

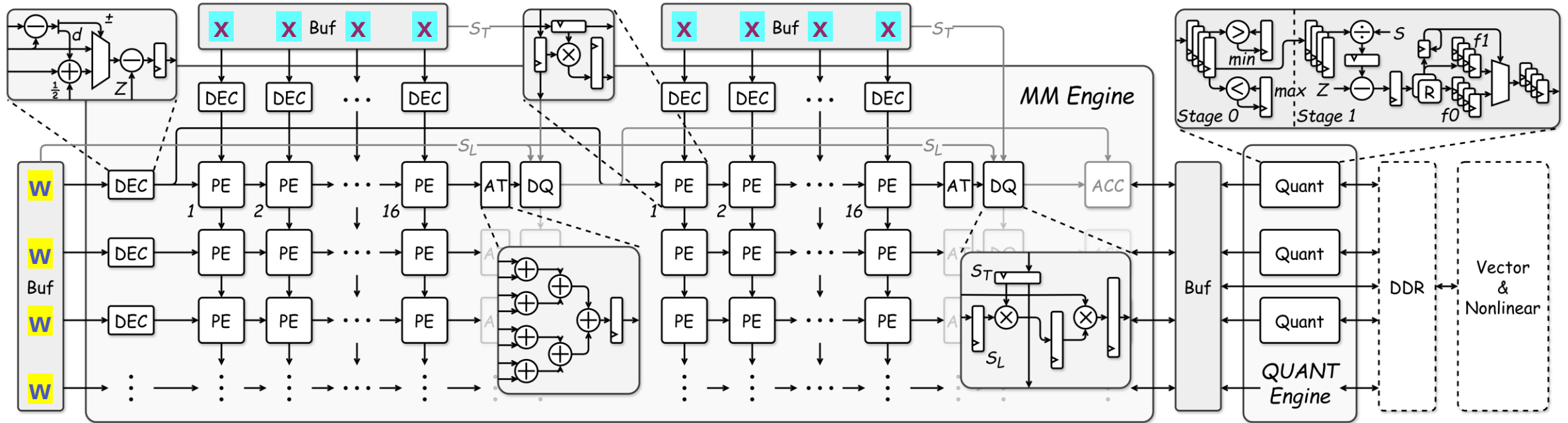
4×64 Adder Trees

4×64 De-quantizers

64 Quantizers



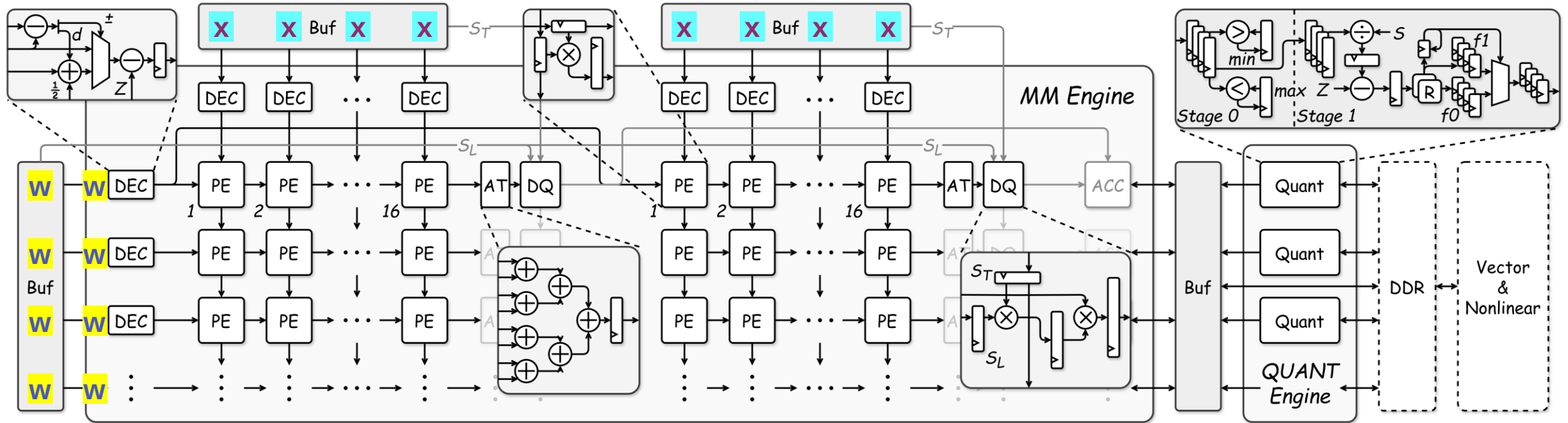
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



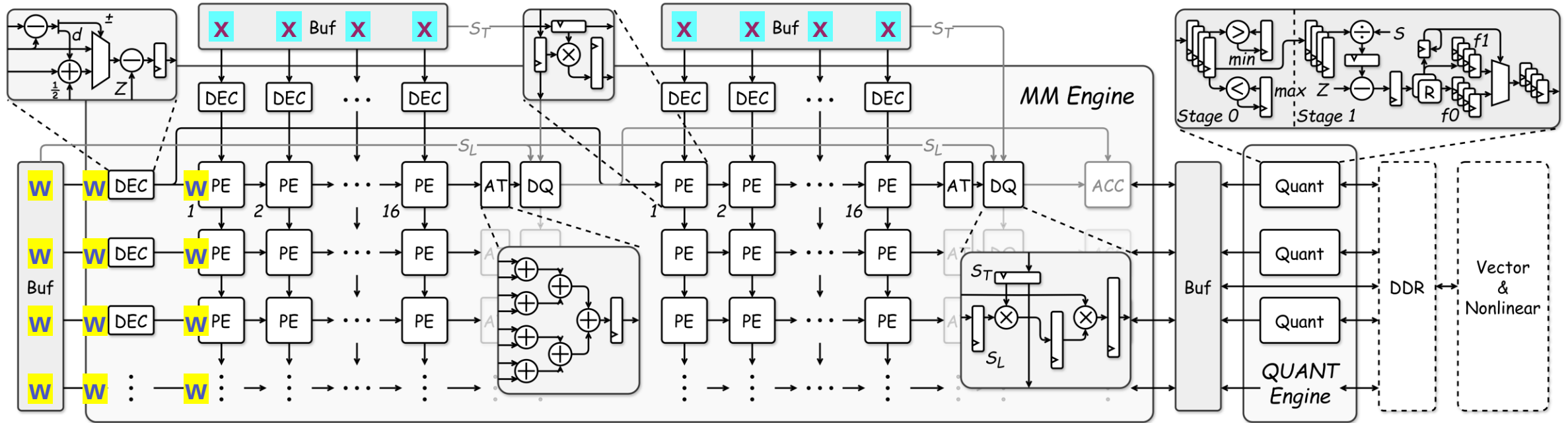
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



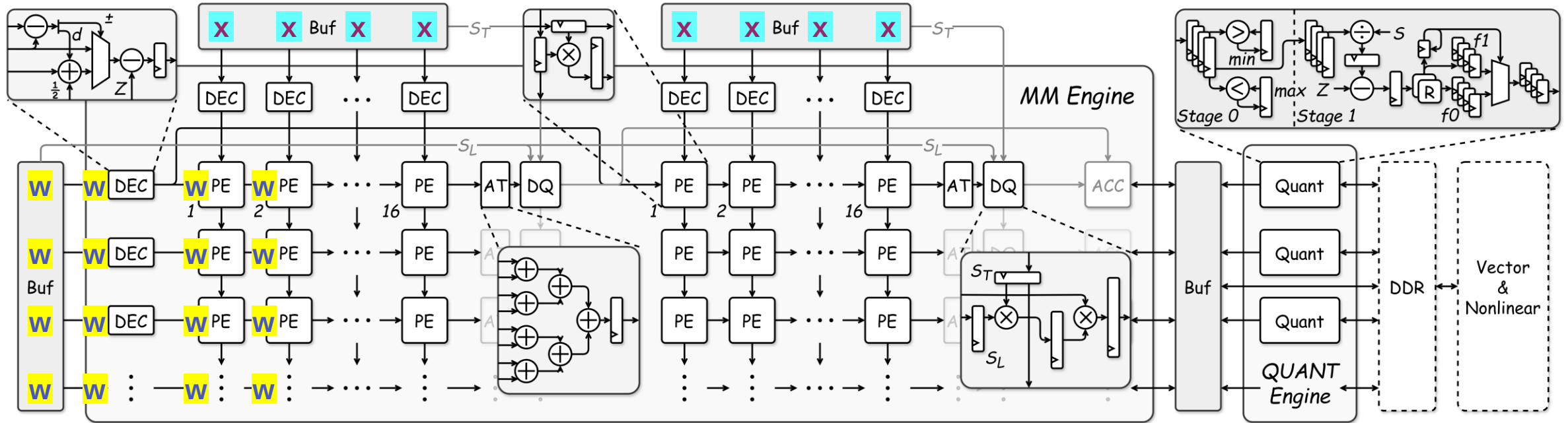
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.

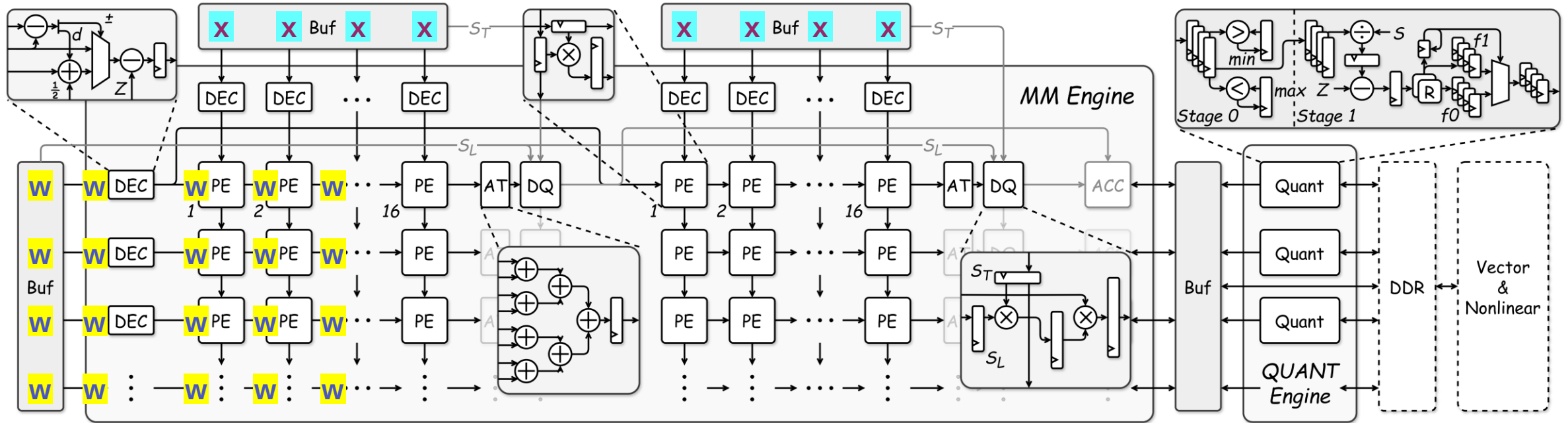


Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.

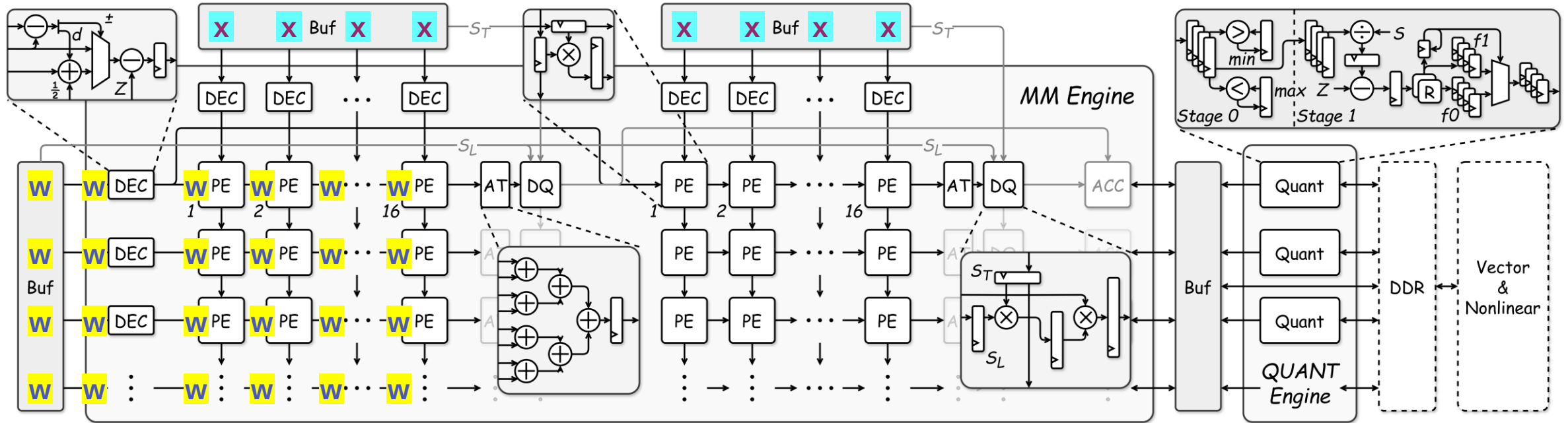
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum.
This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



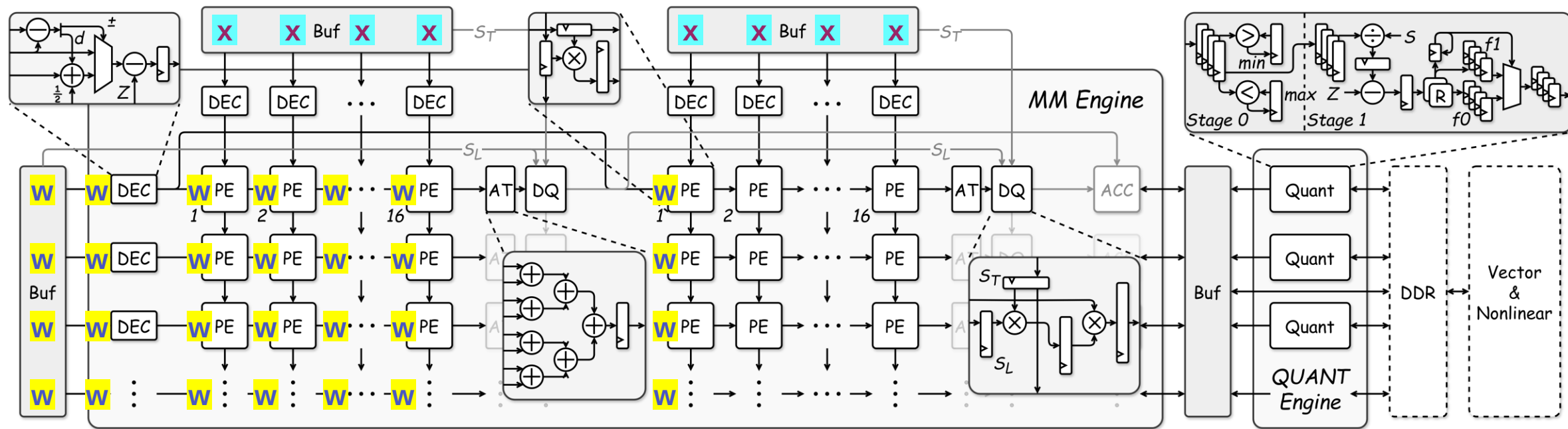
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



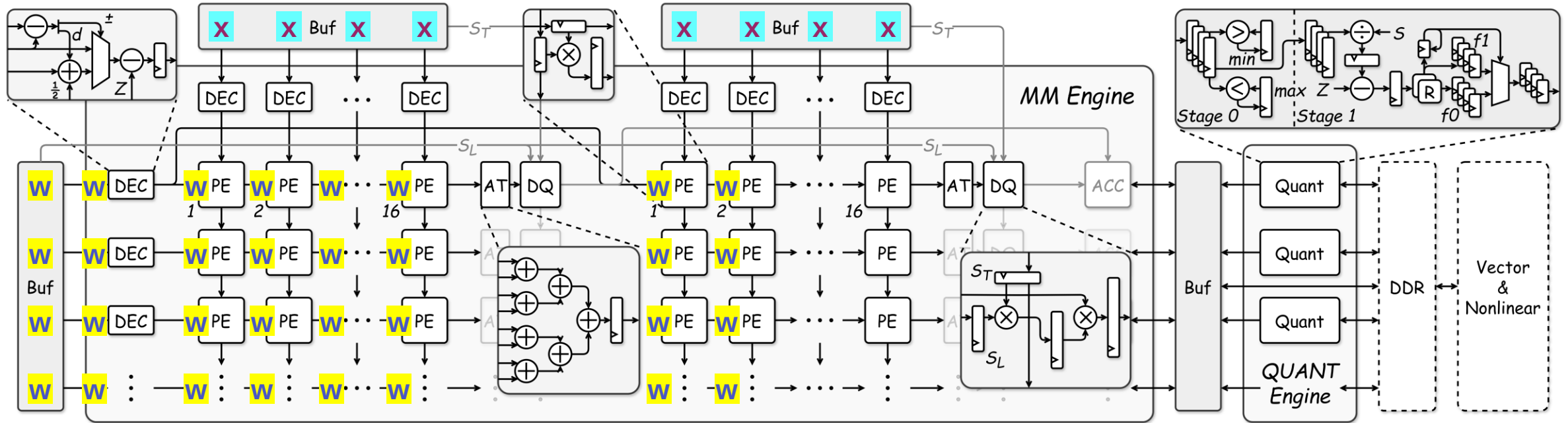
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



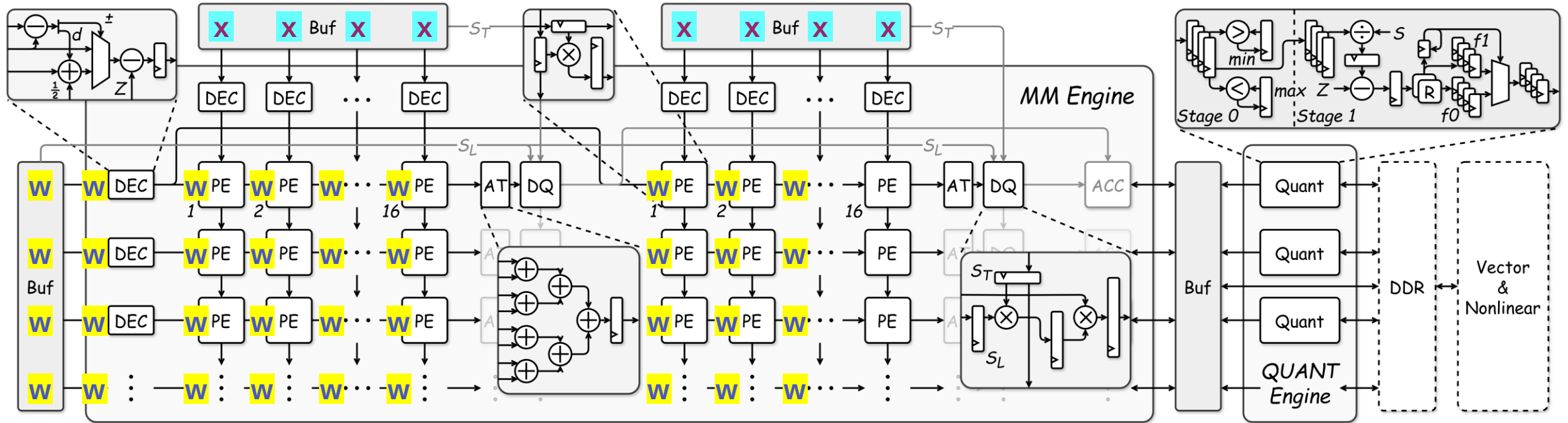
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



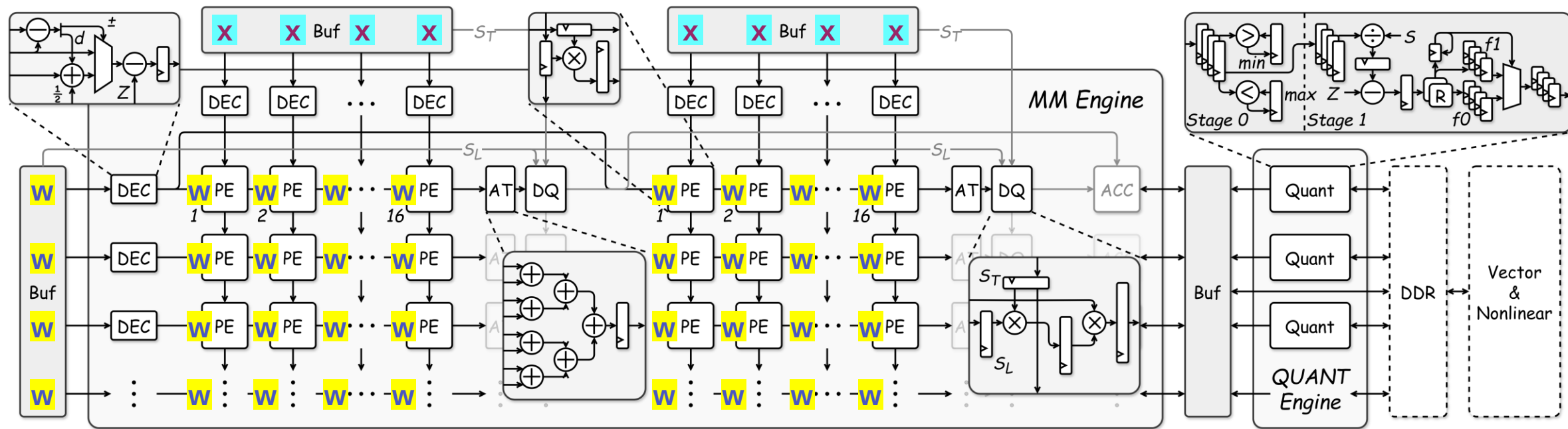
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



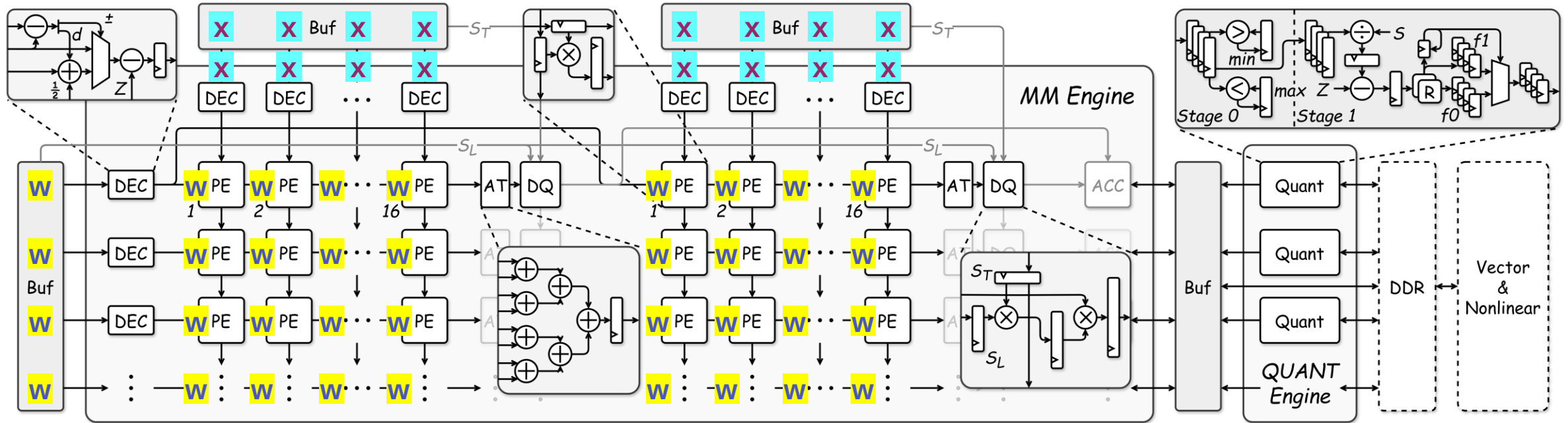
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



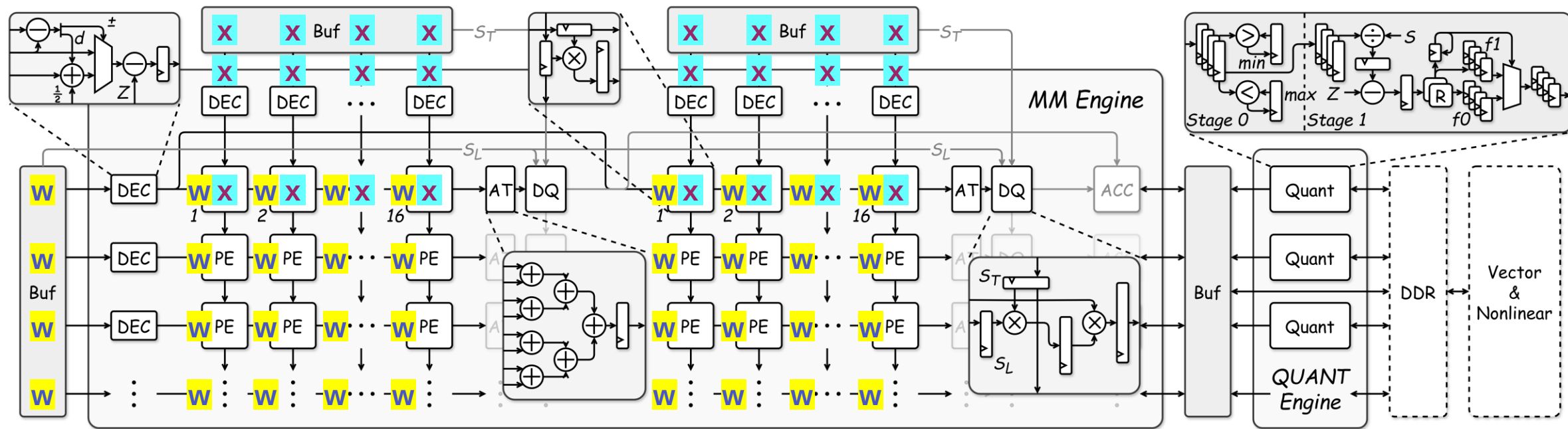
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



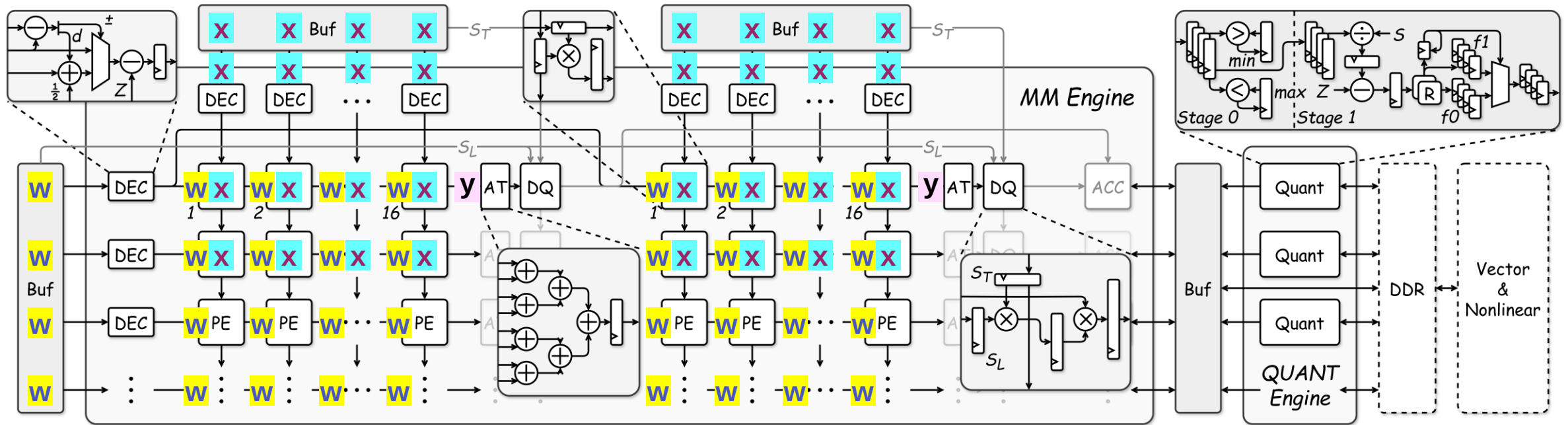
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



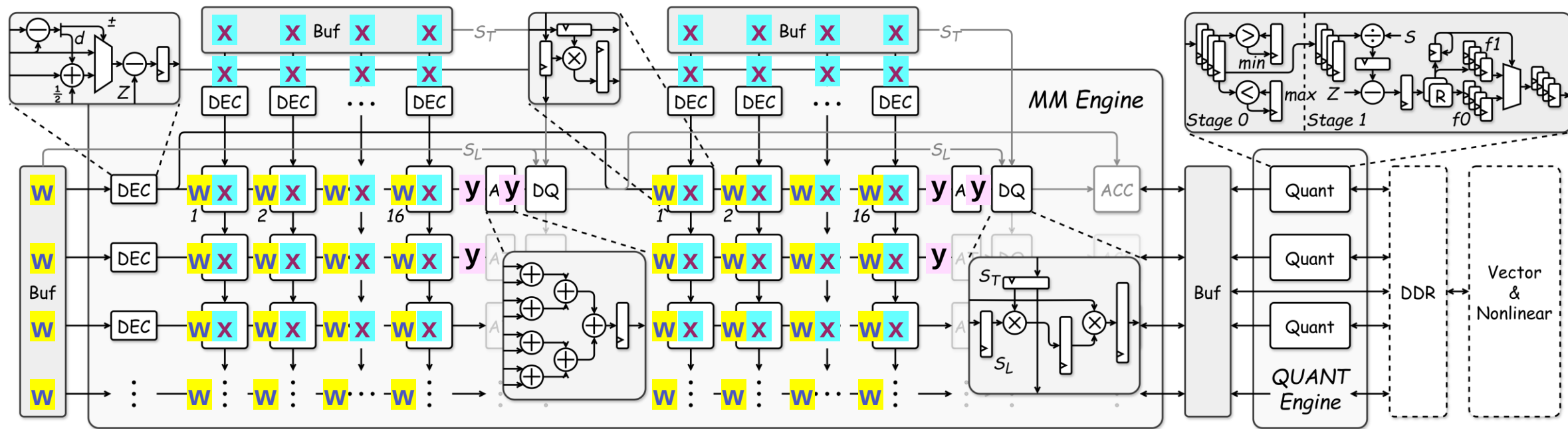
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



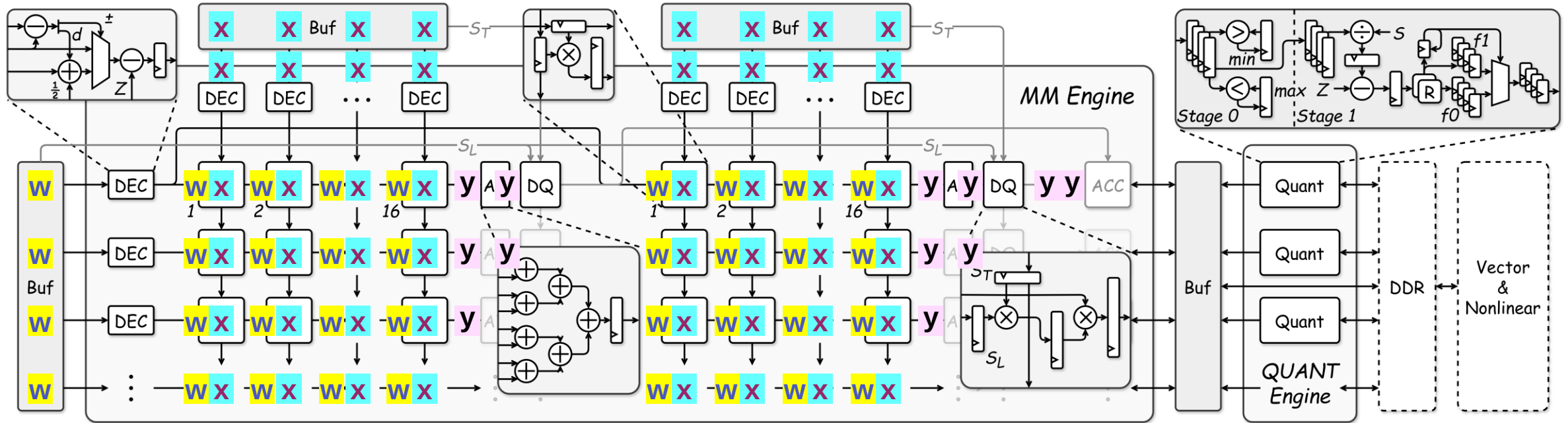
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum.
This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



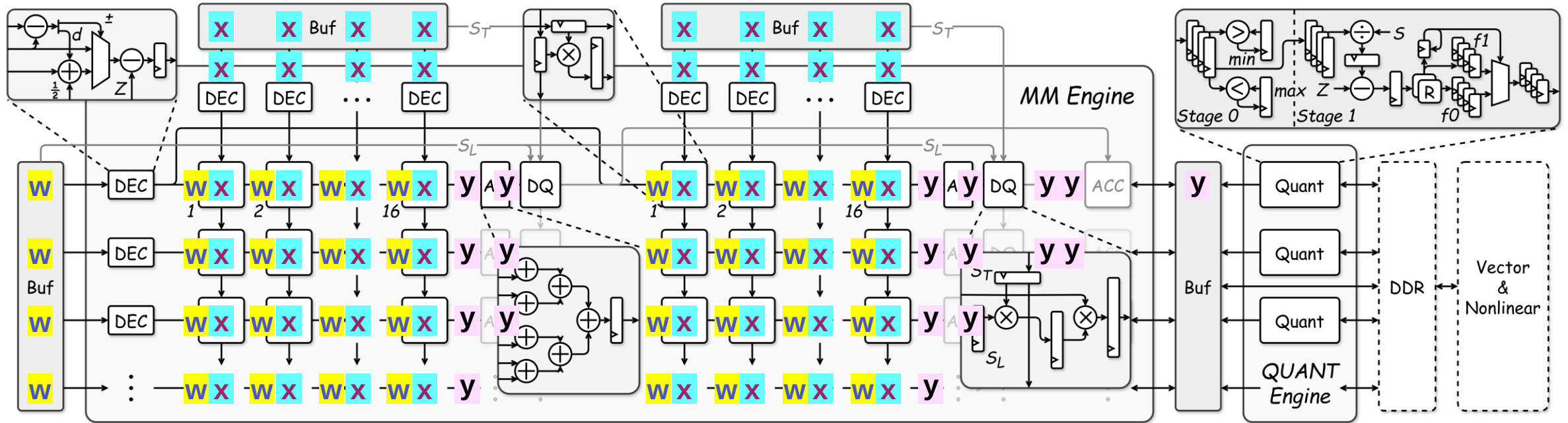
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



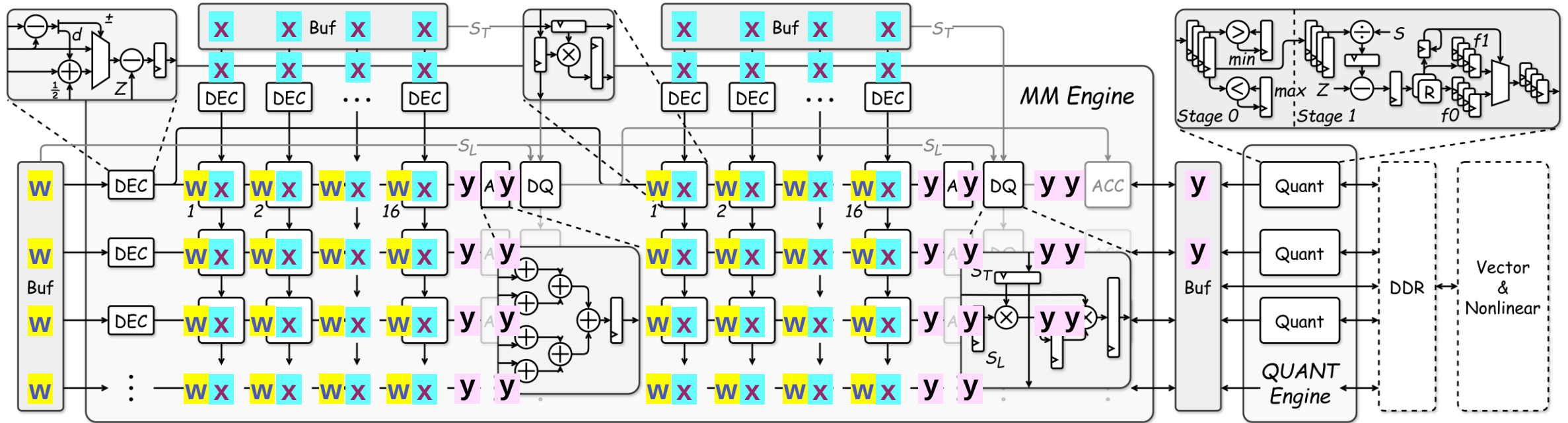
Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



Accelerator Implementation



Decoded values enter **INT-style multipliers**; scaling is deferred until the completion of dot-product sum. This dataflow keeps most computation in **low precision** and avoids unnecessary high-precision operations.



Results & Conclusion

Comparison of LLM Quantization Accelerators

	Prior works			This work
Venue	MICRO'22	HPCA'25	HPCA'25	DAC'26
Method	ANT	M-ANT	BitMod	Adana
Avg. 4-bit PPL↓	9.35	9.45	9.28	9.00
Avg. 3-bit PPL↓	60.39	103.76	19.79	15.17
QU Area (μm^2)	14,328.1	7,828.9	14,931.3	5,165.9
QU Power (mW)	8.22	3.90	8.42	2.80

ANT

M-ANT

BitMod

Adana

0.60

0.48

0.71

1.00

Speedup

1.49

1.85

1.22

1.00

Norm. Energy

Avg. Speedup: 1.68× / 2.10× / 1.42×

Avg. Energy Reduction: 32.81% / 45.85% / 17.91%

- Avg. PPL is averaged over 15 LLMs: Llama-3.1/3.2-1B/3B/8B, Phi-1.5/2/3-1.3B/2.7B/4B, Yi-1.5-6B/9B/34B, Gemma-2/3-1B/2B/4B, and Qwen-3-4B/8B/14B.
- Avg. Speedup and energy results are evaluated on the full accelerator and averaged over four representative models: Llama-3.1-8B, Phi-2-2.7B, Gemma-3-4B, and Qwen-3-14B.



Comparison of LLM Quantization Accelerators

	Prior works			This work
Venue	MICRO'22	HPCA'25	HPCA'25	DAC'26
Method	ANT	M-ANT	BitMod	Adana
Avg. 4-bit PPL↓	9.35	9.45	9.28	9.00
Avg. 3-bit PPL↓	60.39	103.76	19.79	15.17
QU Area (μm²)	14,328.1	7,828.9	14,931.3	5,165.9
QU Power (mW)	8.22	3.90	8.42	2.80

ANT

M-ANT

BitMod

Adana

Speedup

Norm. Energy

0.60

0.48

0.71

1.00

1.49

1.85

1.22

1.00

Avg. Speedup: 1.68x / 2.10x / 1.42x

Avg. Energy Reduction: 32.81% / 45.85% / 17.91%

- Avg. PPL is averaged over 15 LLMs: Llama-3.1/3.2-1B/3B/8B, Phi-1.5/2/3-1.3B/2.7B/4B, Yi-1.5-6B/9B/34B, Gemma-2/3-1B/2B/4B, and Qwen-3-4B/8B/14B.
- Avg. Speedup and energy results are evaluated on the full accelerator and averaged over four representative models: Llama-3.1-8B, Phi-2-2.7B, Gemma-3-4B, and Qwen-3-14B.



Comparison of LLM Quantization Accelerators

	Prior works			This work
Venue	MICRO'22	HPCA'25	HPCA'25	DAC'26
Method	ANT	M-ANT	BitMod	Adana
Avg. 4-bit PPL↓	9.35	9.45	9.28	9.00
Avg. 3-bit PPL↓	60.39	103.76	19.79	15.17
QU Area (μm²)	14,328.1	7,828.9	14,931.3	5,165.9
QU Power (mW)	8.22	3.90	8.42	2.80

ANT

M-ANT

BitMod

Adana

0.60

0.48

0.71

1.00

Speedup

1.49

1.85

1.22

1.00

Norm. Energy

Avg. Speedup: 1.68× / 2.10× / 1.42×

Avg. Energy Reduction: 32.81% / 45.85% / 17.91%

- Avg. PPL is averaged over 15 LLMs: Llama-3.1/3.2-1B/3B/8B, Phi-1.5/2/3-1.3B/2.7B/4B, Yi-1.5-6B/9B/34B, Gemma-2/3-1B/2B/4B, and Qwen-3-4B/8B/14B.
- Avg. Speedup and energy results are evaluated on the full accelerator and averaged over four representative models: Llama-3.1-8B, Phi-2-2.7B, Gemma-3-4B, and Qwen-3-14B.



Comparison of LLM Quantization Accelerators

	Prior works			This work
Venue	MICRO'22	HPCA'25	HPCA'25	DAC'26
Method	ANT	M-ANT	BitMod	Adana
Avg. 4-bit PPL↓	9.35	9.45	9.28	9.00
Avg. 3-bit PPL↓	60.39	103.76	19.79	15.17
QU Area (μm²)	14,328.1	7,828.9	14,931.3	5,165.9
QU Power (mW)	8.22	3.90	8.42	2.80

ANT

M-ANT

BitMod

Adana

0.60

0.48

0.71

1.00

Speedup

1.49

1.85

1.22

1.00

Norm. Energy

Avg. Speedup: 1.68× / 2.10× / 1.42×

Avg. Energy Reduction: 32.81% / 45.85% / 17.91%

- Avg. PPL is averaged over 15 LLMs: Llama-3.1/3.2-1B/3B/8B, Phi-1.5/2/3-1.3B/2.7B/4B, Yi-1.5-6B/9B/34B, Gemma-2/3-1B/2B/4B, and Qwen-3-4B/8B/14B.
- Avg. Speedup and energy results are evaluated on the full accelerator and averaged over four representative models: Llama-3.1-8B, Phi-2-2.7B, Gemma-3-4B, and Qwen-3-14B.



Conclusion

- **Adana Numeric Type**
 - A novel numeric type that effectively captures **nonuniformity** and **asymmetry** in small groups.
- **Quantization Framework**
 - An **efficient online algorithm** that uses rounding-based error counting, avoiding full MSE search.
- **Dedicated Microarchitecture**
 - A **scalable architecture** with MM and QUANT engines, tailored for group-wise quantization.
- **Superior Performance**
 - Adana delivers 1.42× - 2.10× speedups, and 17.9% - 45.9% power savings with consistently higher accuracy compared to **SOTA** works.



**THE CHIPS
TO SYSTEMS
CONFERENCE**

SPONSORED BY:



**Adana: Accelerating Large Language Models via
Addaptive Nonuniform Asymmetric Quantization**

Thank you for your attention.