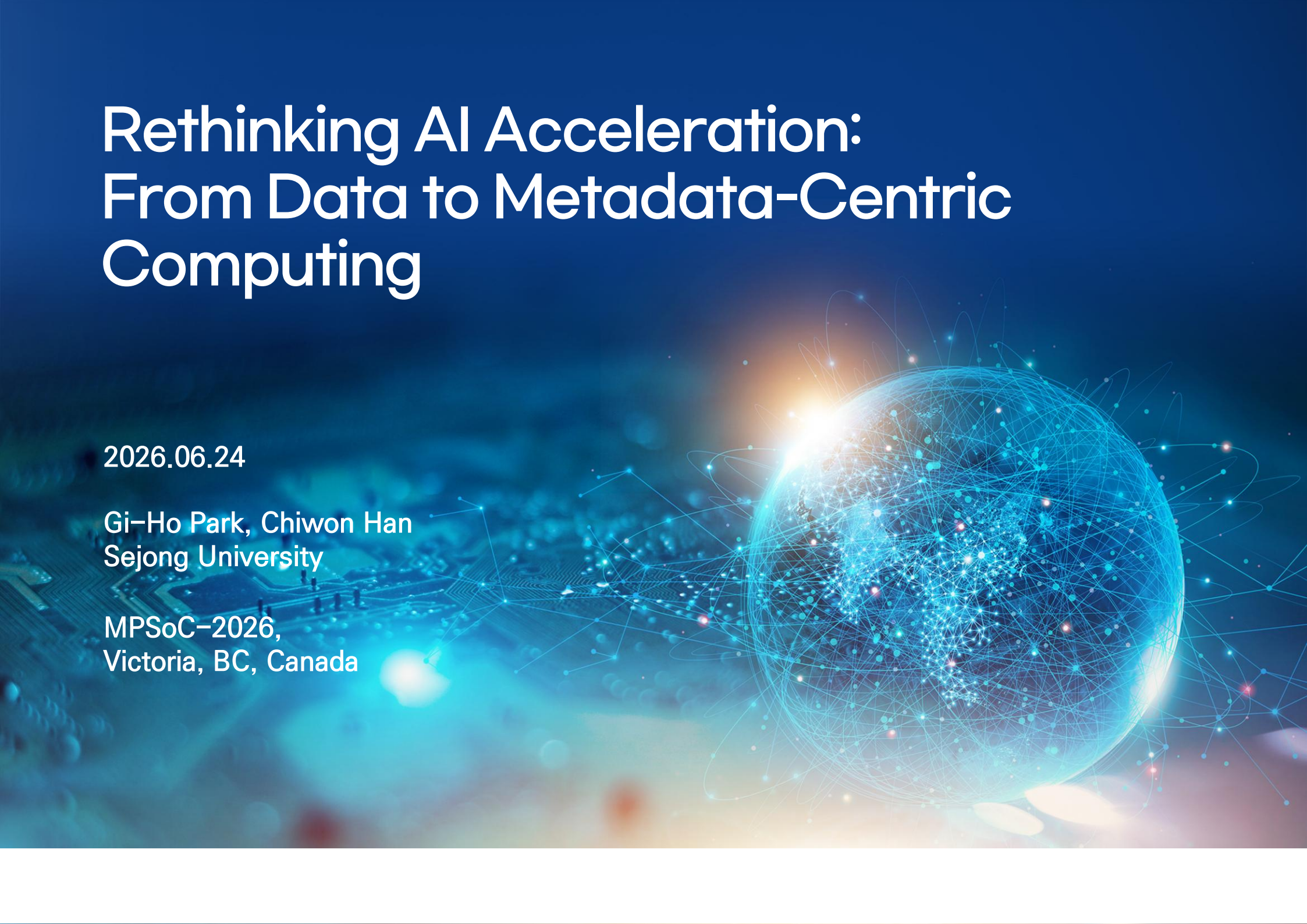


Rethinking AI Acceleration: From Data to Metadata-Centric Computing

2026.06.24

Gi-Ho Park, Chiwon Han
Sejong University

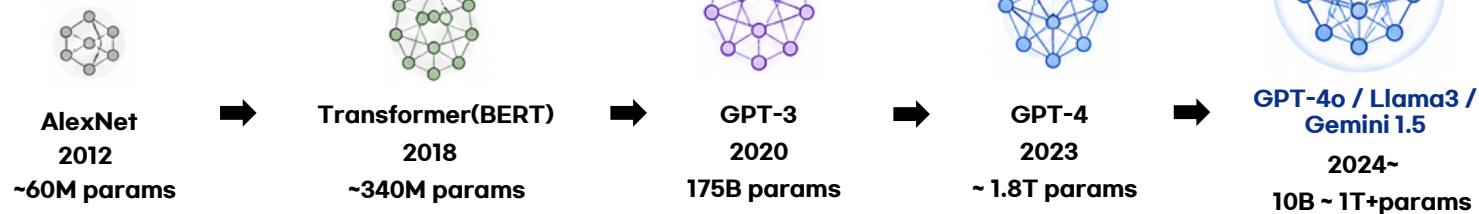
MPSoC-2026,
Victoria, BC, Canada



Models are getting bigger, and compression techniques are evolving

1. Models keep getting larger

Driven by better capability and performance



Bigger models

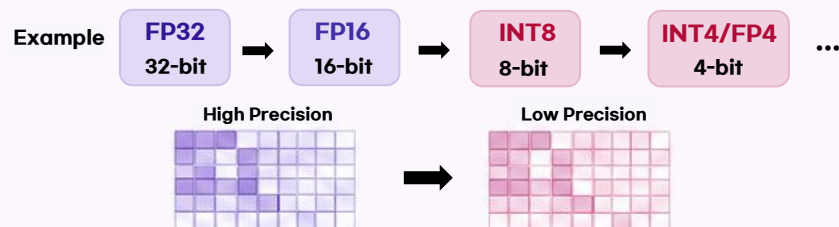
- Higher accuracy
- Higher compute & memory cost

2. Compression techniques have advanced

Two major directions

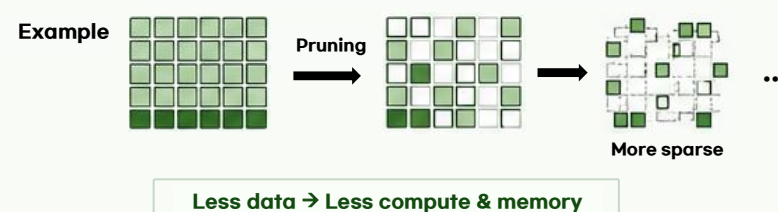
Quantization (Lower precision)

Represent weights/activations with fewer bits



Pruning / Sparsity (Remove unimportant data)

Remove less important weights or connections



3. However, these techniques come with a trade-off

reduce computation and memory requirements VS inevitably reduce accuracy

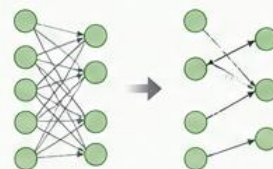
Quantization

Lower precision → less information
→ Quantization Error



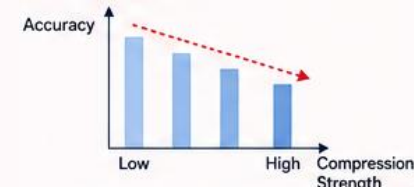
Pruning / Sparsity

Removing weights
→ Loss of useful
information



Common Outcome

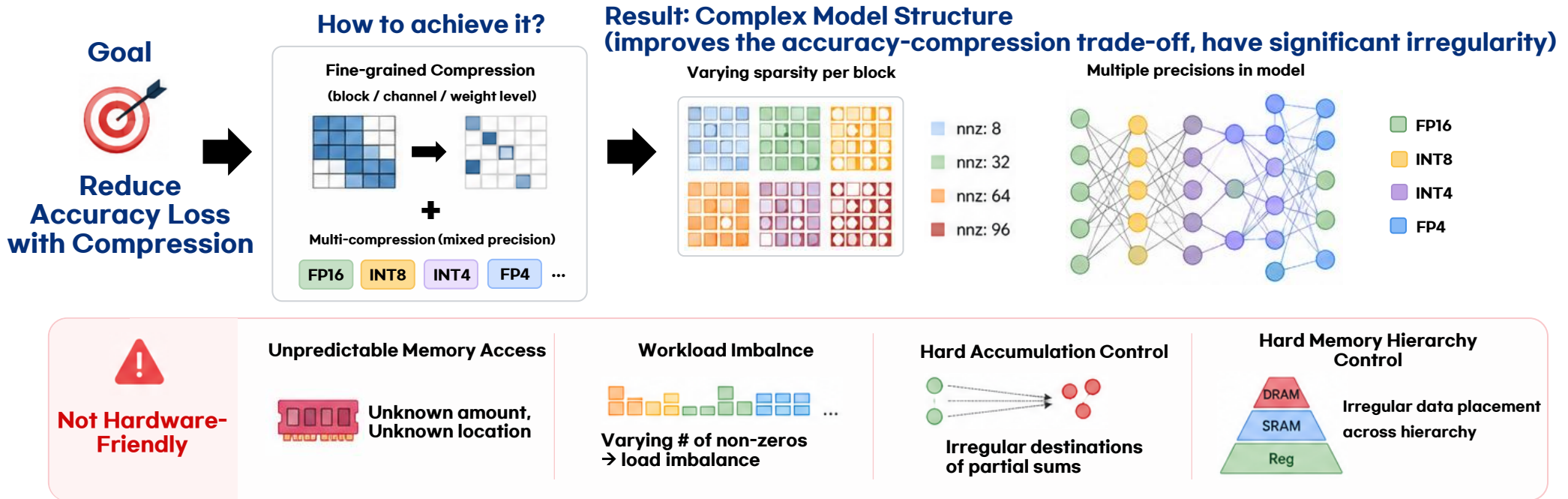
Accuracy degradation



Challenge

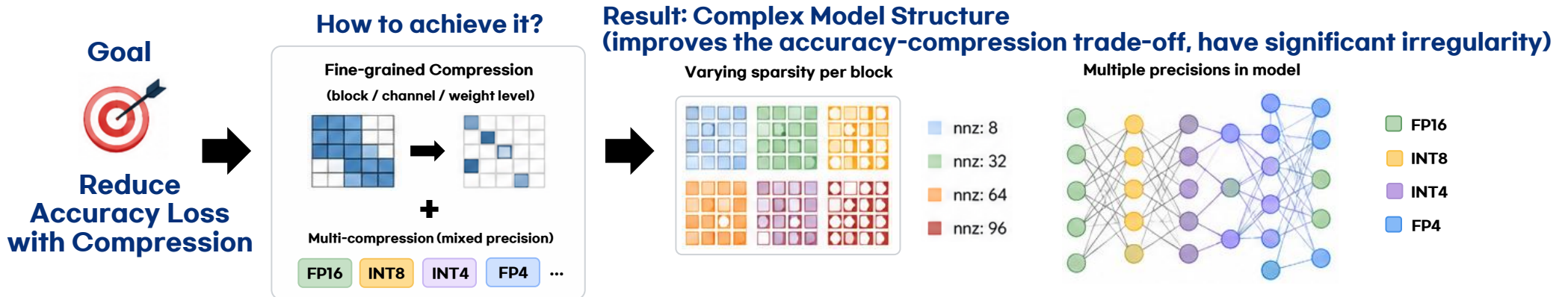
How to maintain
high accuracy
while reducing
computation and
memory footprint?

High accuracy requires fine-grained, multi-compression → complexity and irregularity is inevitable



Achieving high accuracy often means introducing complexity that is not very hardware-friendly.

High accuracy requires fine-grained, multi-compression → complexity is inevitable → Metadata-guided Execution



Not Hardware-Friendly

Unpredictable Memory Access



Unknown amount, Unknown location

Workload Imbalance



Varying # of non-zeros
→ load imbalance

Hard Accumulation Control



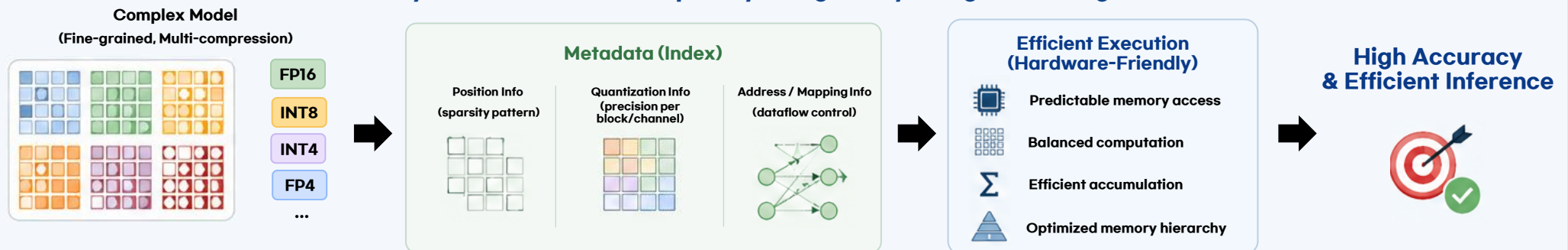
Irregular destinations of partial sums

Hard Memory Hierarchy Control



Irregular data placement across hierarchy

Key is to control the complexity/irregularity using well-designed metadata (index)

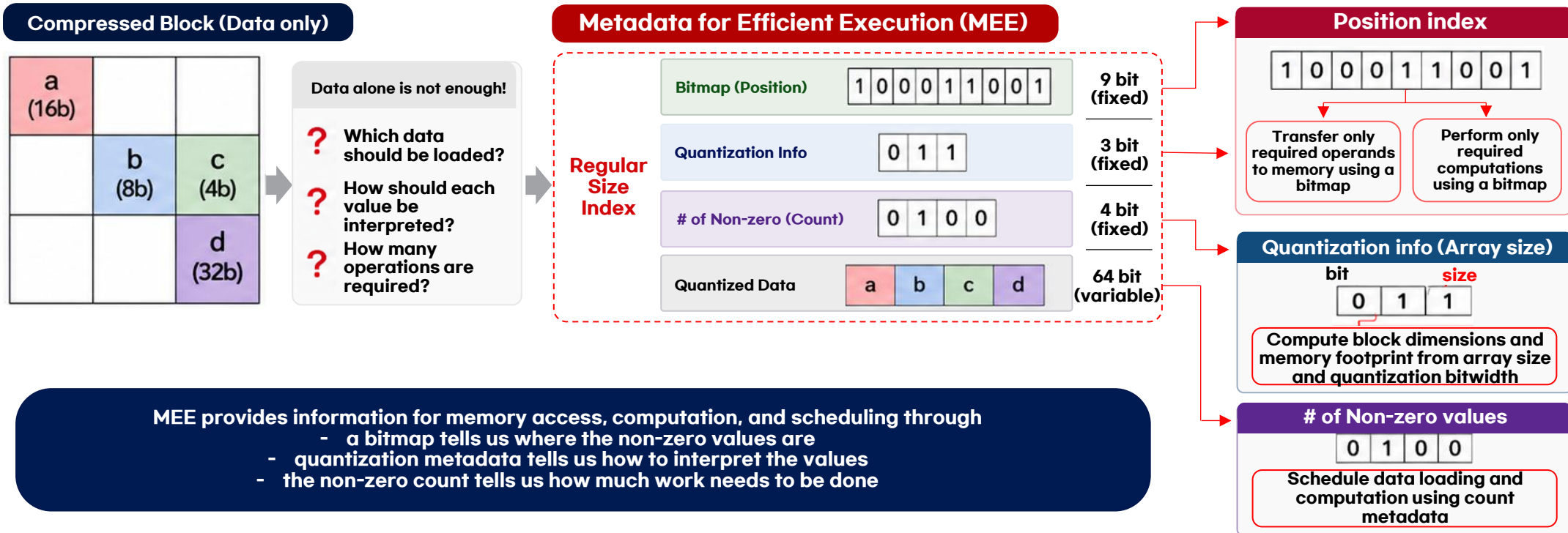


What if we use metadata to help hardware understand and manage complexity? (Instead of treating metadata as an overhead)

→ This idea led us to metadata-guided execution

Metadata for Efficient Execution (MEE)

Providing Useful Information for Hardware via Metadata



Metadata for Efficient Execution (MEE): Metadata for Memory Access, Computing, and Scheduling

Compressed Block (Data only)

a (16b)		
	b (8b)	c (4b)
		d (32b)

Data alone is not enough!

- ? Which data should be loaded?
- ? How should each value be interpreted?
- ? How many operations are required?

Metadata for Efficient Execution (MEE)

Regular
Size
Index

Bitmap (Position)	1 0 0 0 1 1 0 0 1	9 bit (fixed)
Quantization Info	0 1 1	3 bit (fixed)
# of Non-zero (Count)	0 1 0 0	4 bit (fixed)
Quantized Data	a b c d	64 bit (variable)

Position index

1 0 0 0 1 1 0 0 1

Transfer only
required operands
to memory using a
bitmap

Perform only
required
computations
using a bitmap

Quantization info (Array size)

bit size
0 1 1

Compute block dimensions and
memory footprint from array size
and quantization bitwidth

of Non-zero values

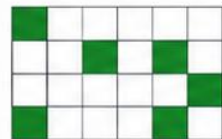
0 1 0 0

Schedule data loading and
computation using count
metadata

1. Position Metadata (Where)

Bitmap (Mask bit)

1 0 0 0 1 1 0 0 1



- ✓ Load only valid operands
- ✓ Skip useless memory access



Where is the useful data?

Load only useful data
(Bandwidth & Energy saving)

2. Value Metadata (What)

Quantization Info (bit table + precision)

Bit pattern	Precision
00	32bit
01	16bit
10	8bit
11	4bit

- ✓ Decode with correct precision
- ✓ Control PE utilization



What does the
value represent?

Decode correctly &
Utilize proper precision

3. Control Metadata (How)

of Non-zero (Count)

0 1 0 0

= 4 (non-zero)

- ✓ Know how many operations
- ✓ Schedule workload in advance



How much work is
required?

Schedule & balance
workload efficiently

MEE provides information for memory access, computation, and scheduling

Metadata for Efficient Execution (MEE): Predictable Execution through Metadata

Compressed Block (Data only)

a (16b)		
	b (8b)	c (4b)
		d (32b)

Data alone is not enough!

- ? Which data should be loaded?
- ? How should each value be interpreted?
- ? how many operations are required?

Metadata for Efficient Execution (MEE)

Regular
Size
Index

Bitmap (Position)	1 0 0 0 1 1 0 0 1	9 bit (fixed)
Quantization Info	0 1 1	3 bit (fixed)
# of Non-zero (Count)	0 1 0 0	4 bit (fixed)
Quantized Data	a b c d	64 bit (variable)

Position index

1 0 0 0 1 1 0 0 1

Transfer only
required operands
to memory using a
bitmap

Perform only
required
computations
using a bitmap

Quantization info (Array size)

bit size
0 1 1

Compute block dimensions and
memory footprint from array size
and quantization bitwidth

of Non-zero values

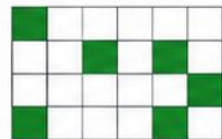
0 1 0 0

Schedule data loading and
computation using count
metadata

1. Position Metadata (Where)

Bitmap (Mask bit)

1 0 0 0 1 1 0 0 1



- ✓ Load only valid operands
- ✓ Skip useless memory access



Where is the useful data?

Load only useful data
(Bandwidth & Energy saving)

2. Value Metadata (What)

Quantization Info (bit table + precision)

Bit pattern	Precision
00	32bit
01	16bit
10	8bit
11	4bit

- ✓ Decode with correct precision
- ✓ Control PE utilization



What does the
value represent?

Decode correctly &
Utilize proper precision

3. Control Metadata (How)

of Non-zero (Count)

0 1 0 0

= 4 (non-zero)

- ✓ Know how many operations
- ✓ Schedule workload in advance



How much work is
required?

Schedule & balance
workload efficiently

Result: Predictable/Efficient Execution



Predictable
Memory Access
(Where to load)



Predictable
Computation
(How many ops)



Predictable
Decoding
(How to decode)



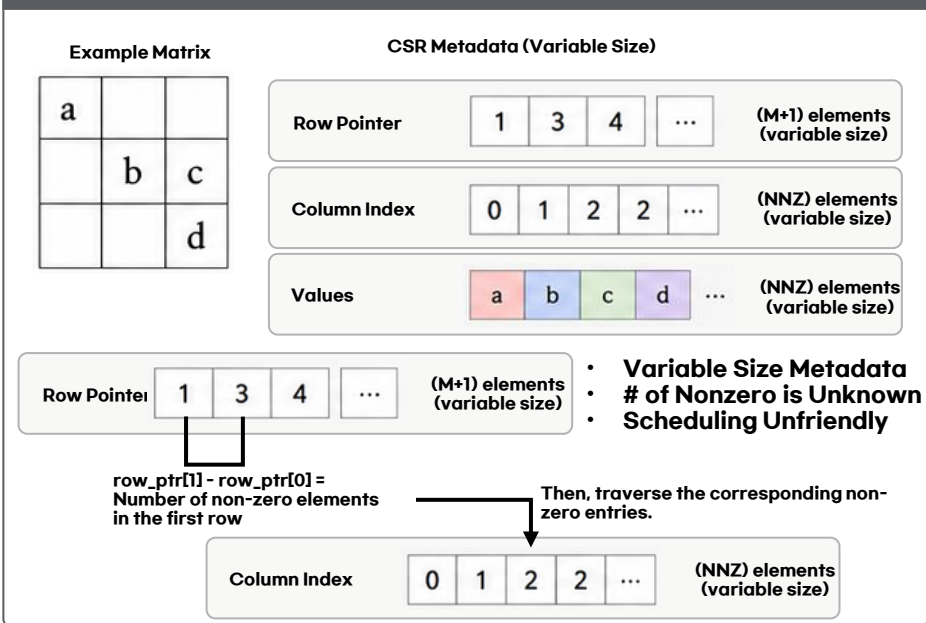
High Accuracy
& High Efficiency

MEE provides information for memory access, computation, and scheduling

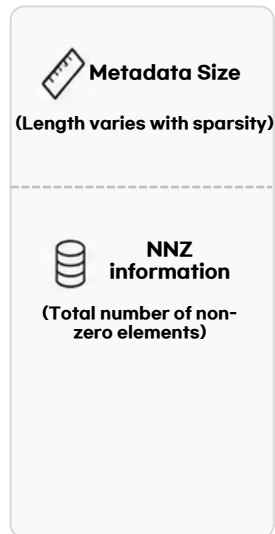
Fixed-Size Metadata itself for Efficient Processing

Predictable and Regular Fetching for Metadata Itself

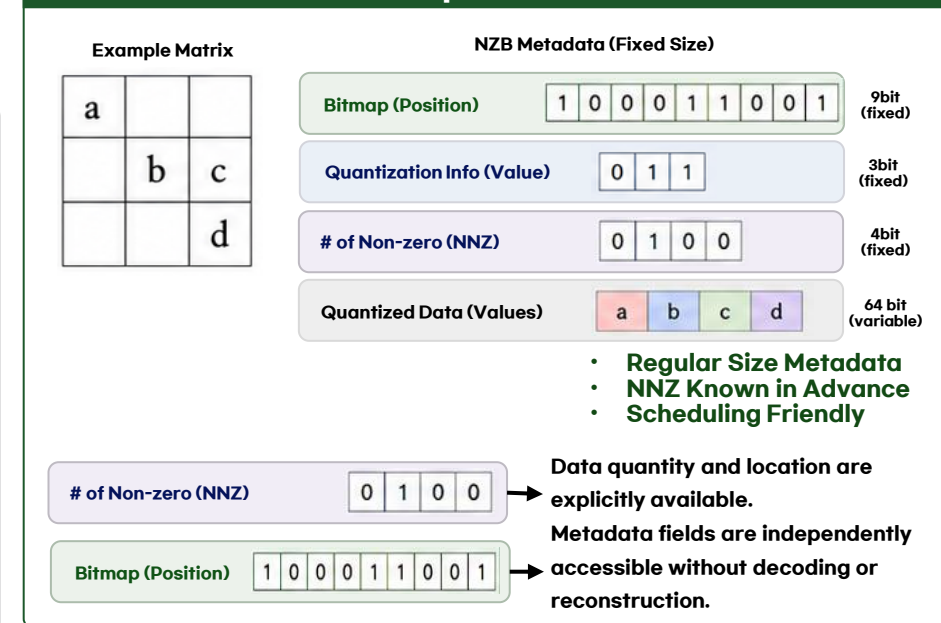
CSR (Compressed Sparse Row)



VS



MEE (Bitmap + Quant + NNZ)

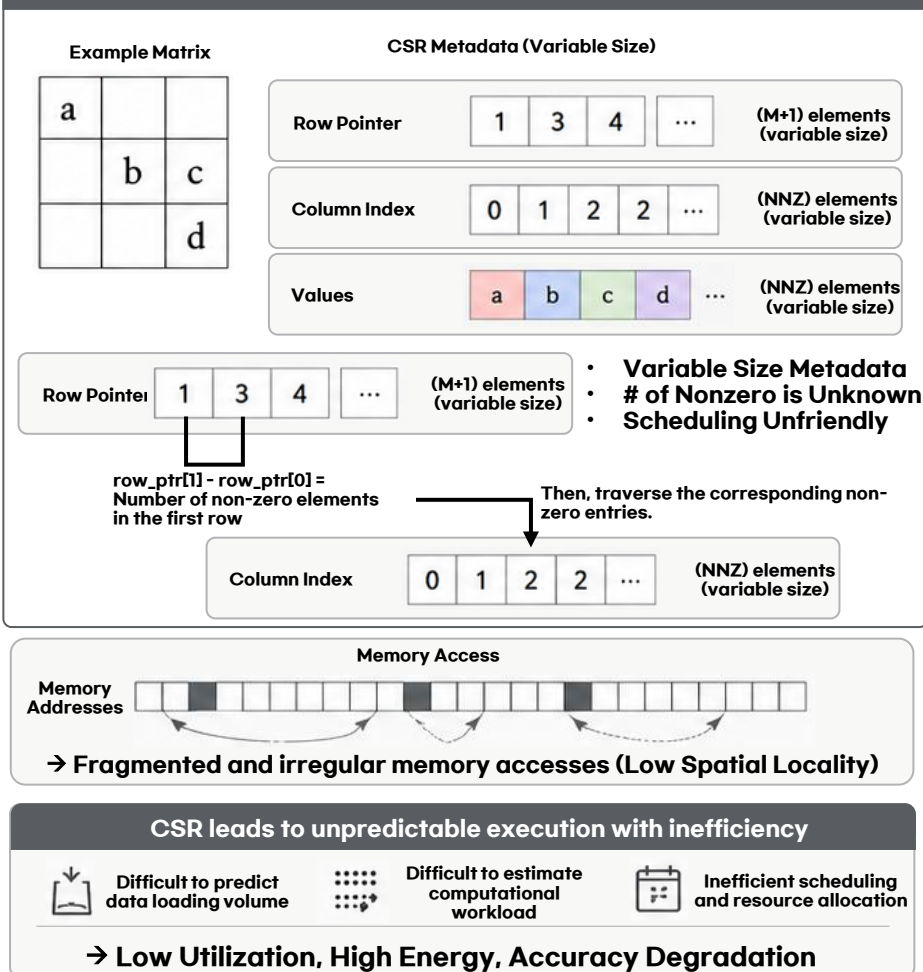


MEE itself is regular. Unlike CSR, the metadata size does not change with sparsity. bitmap, quantization information, and NNZ count all have fixed sizes.

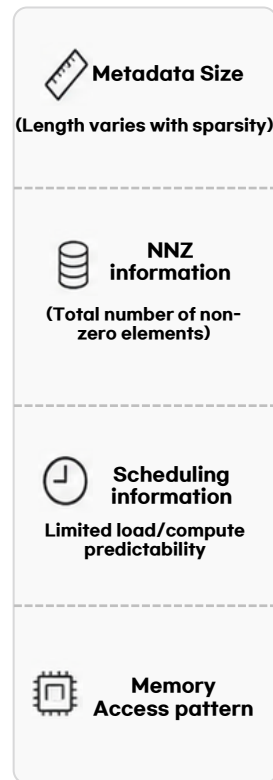
Fixed-Size Metadata for Efficient Processing

Predictable Execution Enables High Efficiency

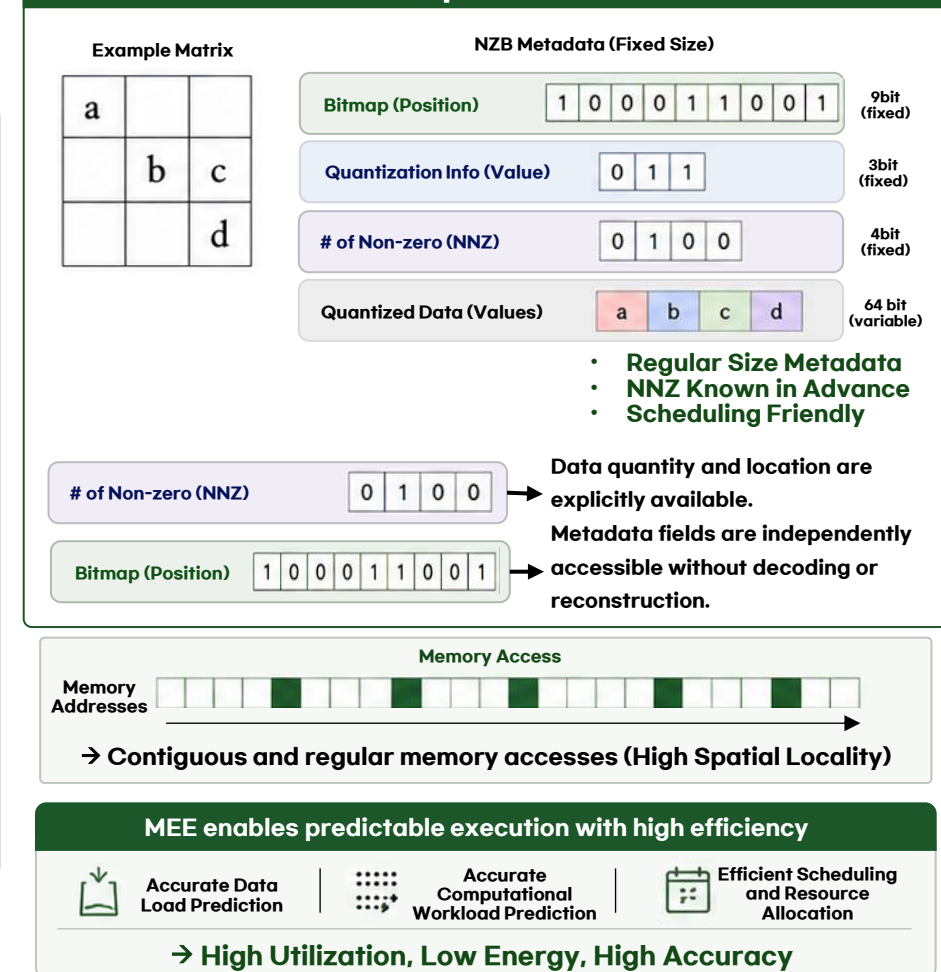
CSR (Compressed Sparse Row)



VS



MEE (Bitmap + Quant + NNZ)



💡 MEE leverages fixed-size metadata and provide explicit NNZ information to determine Where, What, and How Many before execution.

→ Predictable Memory Access, Efficient Computation, and Effective Scheduling

VS

CSR only tells us where the non-zero values are.

Three Approaches to Processing Compressed Matrices

Processing Compressed Matrices : From Decompression → Irregular Access → Regular Metadata Access

1. Conventional (Decompress & Compute)

Decompress to Dense, Then Compute

Compressed Matrix
(Example)

a b c d

Decompress

Dense Matrix

a	0	0
0	b	c
0	0	d

Data
Representation

2. CSR/COO/CSC (Compressed Sparse Format)

Compute on Compressed Data,
But Irregular Access & Index Overhead

Compressed Matrix
(Example)

a b c d

Row Pointer 1 3 4 ... (M+1)

Column Index 0 1 2 ... (NNZ)

Values a b c d (NNZ)

3. MEE (New Metadata Format)

Compute on Compressed Data,
With Regular Metadata Access

MEE Metadata (Fixed Size per Block/Row)

Bitmap (Position) 1 0 0 0 1 1 0 0 1 (9bit)

Quantization Info 0 1 1 (3bit)

NNZ Count (Control) 0 1 0 0 (4bit)

Quantized Data (values) a b c d (64bit)

Fixed
size

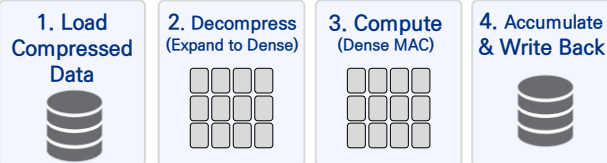
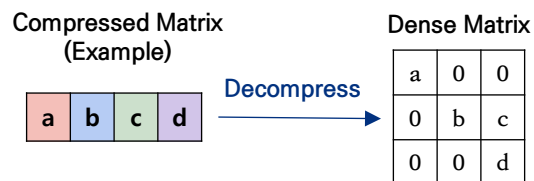
Three Approaches to Processing Compressed Matrices

Execution Flow Comparison for Compressed Matrices

Processing Compressed Matrices : From Decompression → Irregular Access → Regular Metadata Access

1. Conventional (Decompress & Compute)

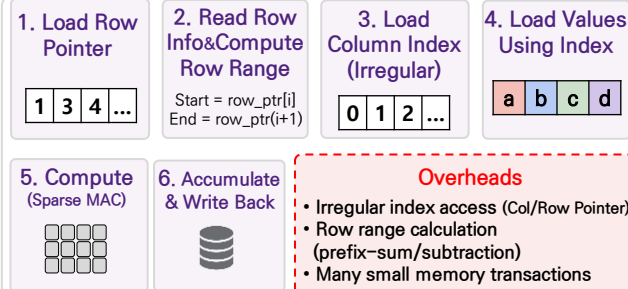
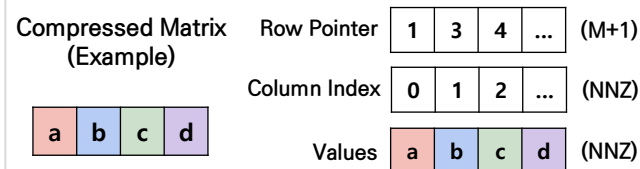
Decompress to Dense, Then Compute



Decompression Overhead
(Extra Memory, Bandwidth, Energy)

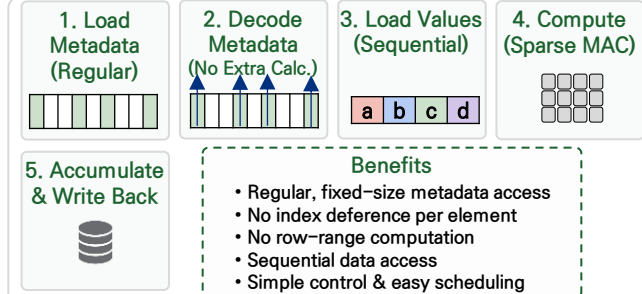
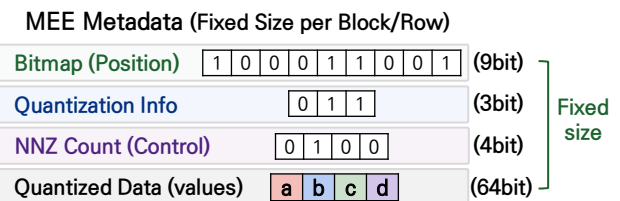
2. CSR/COO/CSC (Compressed Sparse Format)

Compute on Compressed Data,
But Irregular Access & Index Overhead



3. MEE (New Metadata Format)

Compute on Compressed Data,
With Regular Metadata Access



Three Approaches to Processing Compressed Matrices

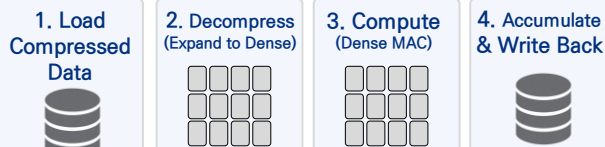
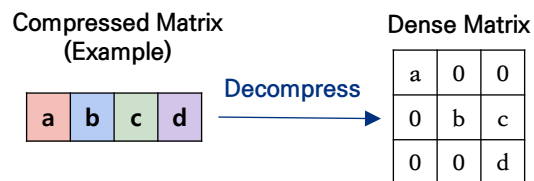
Regular Metadata Enables Predictable Execution

Processing Compressed Matrices : From Decompression → Irregular Access → Regular Metadata Access

Data Representation

1. Conventional (Decompress & Compute)

Decompress to Dense, Then Compute



Decompression Overhead
(Extra Memory, Bandwidth, Energy)

Hardware Execution Path

Characteristics

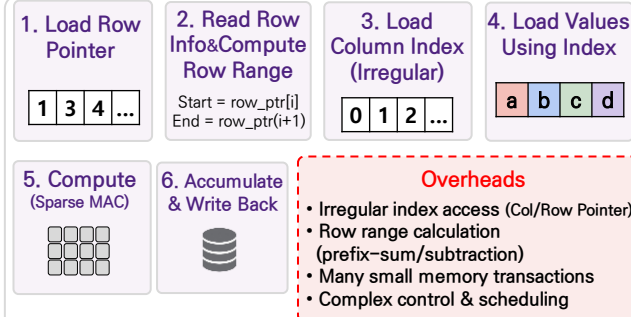
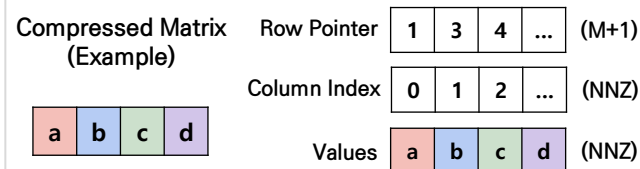
- Must restore to dense format
- Additional decompression cost (time, energy, memory)
- High memory traffic (read compressed + write dense)
- Simple compute, but inefficient overall

Key Issue

Decompression overhead and unnecessary memory traffic for zeros

2. CSR/COO/CSC (Compressed Sparse Format)

Compute on Compressed Data, But Irregular Access & Index Overhead



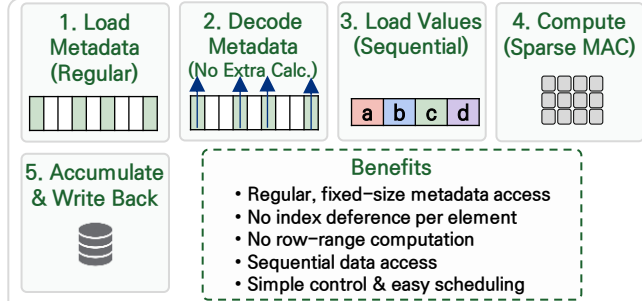
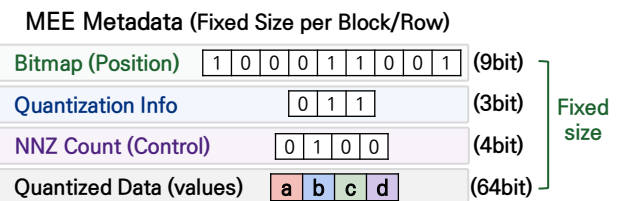
- Compute on compressed data (no decompression)
- But requires index access for every non-zero element
- Row pointer calculation needed for each row
- Irregular memory access → low locality, hard to predict
- Additional control logic and complex scheduling

Key Issue

Irregular metadata access and index computation cause low predictability and inefficiency

3. MEE (New Metadata Format)

Compute on Compressed Data, With Regular Metadata Access



- Compute on compressed data without decompression
- Metadata has fixed size and regular structure
- No per-element idx access, no row range calculation
- Sequential memory access → high locality
- Highly predictable, simple control, efficient scheduling

Key Issue

Regular metadata access enables predictable, efficient and high-performance execution

Summary

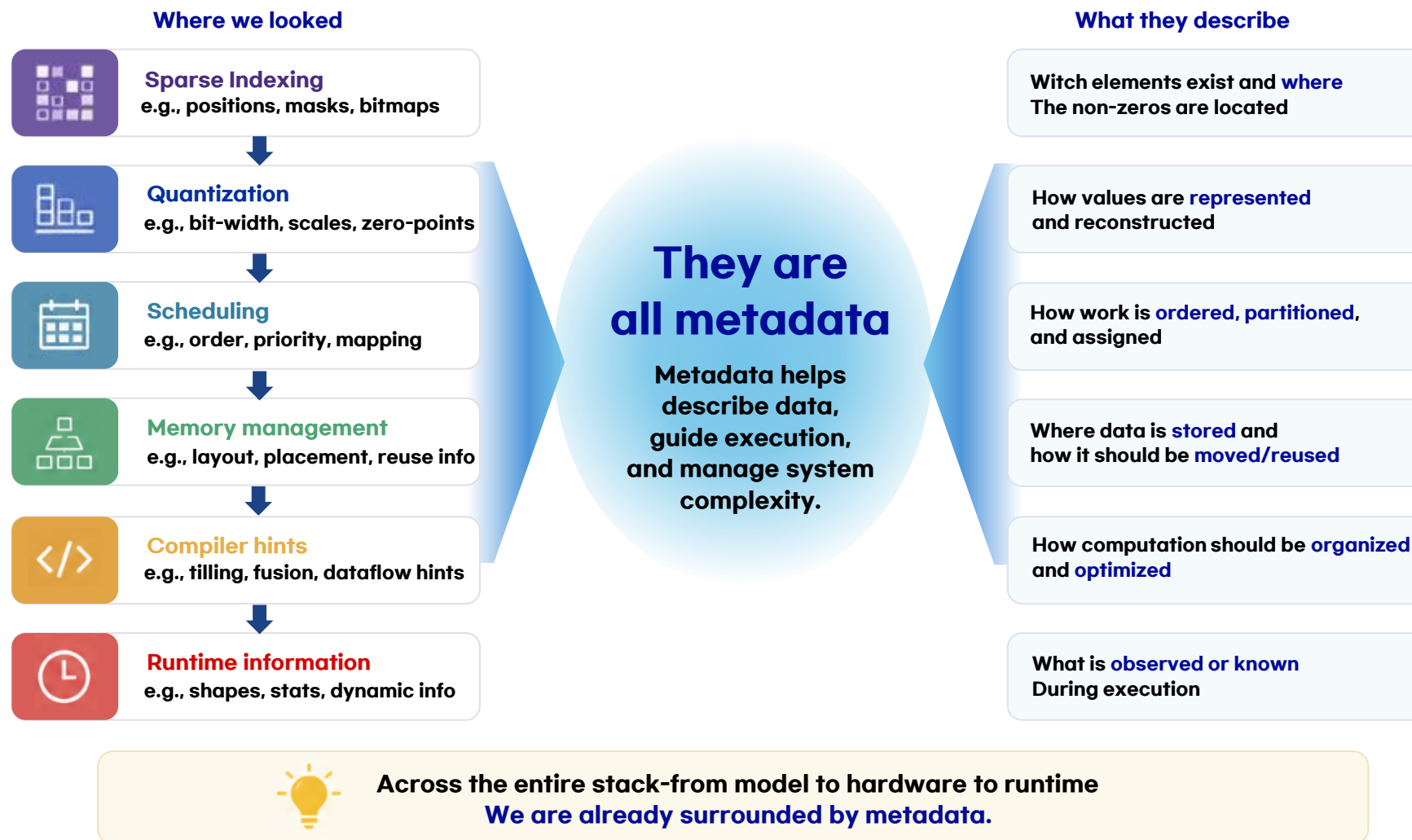
1) Decompression cost is high

2) Still compressed, but irregular access & index computation hurt performance

3) Compressed + Regular Metadata → Predictable, Efficient and Scalable

What surprised us?

As we continued working on compression-related research, including sparse representations and quantization, we started noticing a common pattern.



What is Metadata?

Metadata is information **ABOUT** data that that helps efficient and correct execution.

Metadata Type	What it Describes	Example	Illustration (Example)																															
 Sparse Indexing Metadata	Where non-zero values exist	CSR / COO / CSC (Index Format) (e.g., CSR)	Sparse Matrix <table><tr><td>1</td><td>0</td><td>0</td><td>2</td><td>0</td></tr><tr><td>0</td><td>3</td><td>0</td><td>0</td><td>0</td></tr><tr><td>4</td><td>0</td><td>5</td><td>0</td><td>6</td></tr></table> → CSR Representation Values <table><tr><td>1</td><td>2</td><td>3</td><td>4</td><td>5</td><td>6</td></tr></table> (non-zero values) Col Indices <table><tr><td>0</td><td>3</td><td>1</td><td>0</td><td>2</td><td>4</td></tr></table> (column indices) Row Ptr <table><tr><td>0</td><td>2</td><td>3</td><td>6</td></tr></table> (row pointers)	1	0	0	2	0	0	3	0	0	0	4	0	5	0	6	1	2	3	4	5	6	0	3	1	0	2	4	0	2	3	6
1	0	0	2	0																														
0	3	0	0	0																														
4	0	5	0	6																														
1	2	3	4	5	6																													
0	3	1	0	2	4																													
0	2	3	6																															
 Quantization Metadata	How values should be interpreted	Quantization (Scale, Zero-point, Bit-width)	Quantization parameters for a block/channel Scale = 0.03125 Zero-point = 128 Bit-width = 4-bit																															
 Control Metadata	How work should be assigned and executed.	NNZ Count (Non-Zero Count)	Number of non-zero elements 128 64 256 ...																															
 Scheduling Metadata	When, where, and in what order work should be executed	Task Mapping (Priority, Mapping, Dependency)	Mapping of tasks to processing elements (PEs) Task 1 → PE3 Task 2 → PE7 Task 3 → PE1 Task 4 → PE5 ...																															
 Memory Metadata	Where determine where data resides and how it should move	Placement / Location (Reuse, Layout)	Data placement across memory hierarchy DRAM → SRAM → Buffer → Register																															

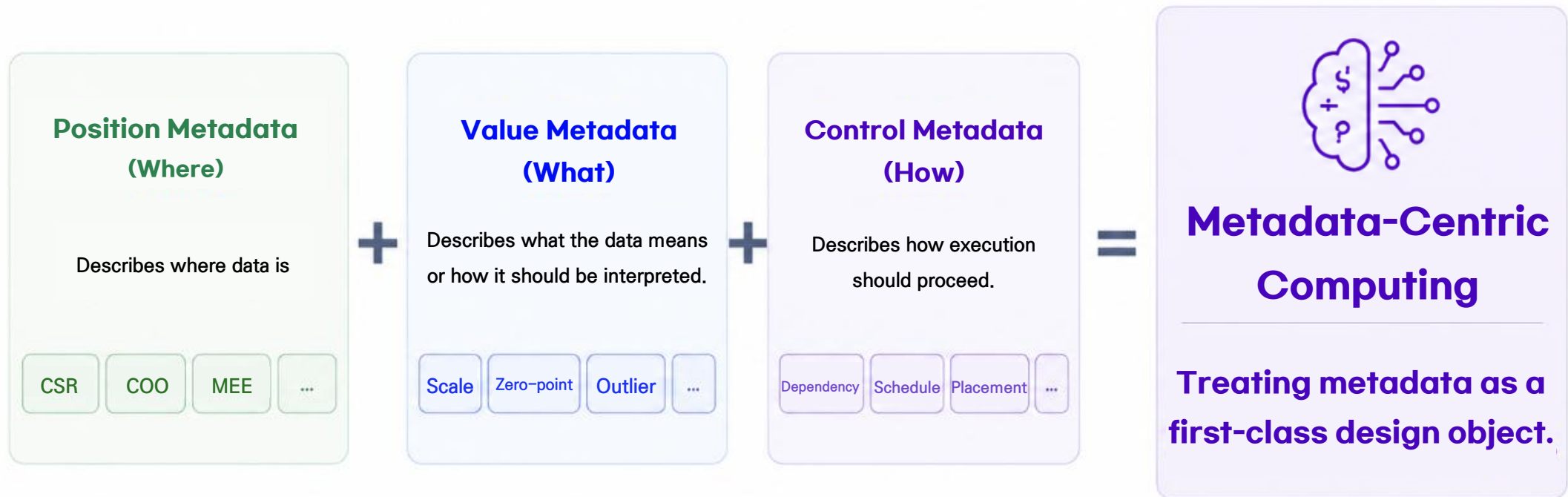
Metadata is used in many area as essential information for describing



WHERE data resides, **WHAT** it represents, and **HOW** computation proceeds.

Metadata for Efficient Execution in AI Systems

From Overhead to First-Class Design Object



Metadata describes how data is accessed, interpreted, and executed.

This perspective motivates us to think about Metadata-Centric Computing.

What makes metadata effective?

3 simple design principles learned from MEE design



Metadata is stored together with the data

Compressed & Quantized Data



Metadata (Index)



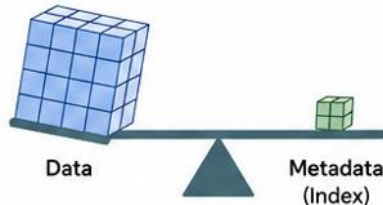
□ Data
■ Metadata (Index)

Good metadata enables efficient, predictable, and hardware-friendly execution.

1

Metadata should be much smaller than the data itself.

Metadata overhead can outweigh its benefits.



Size Constraint

$$|\text{Metadata}| \ll |\text{Data}|$$

Example

Quantization (INT4)



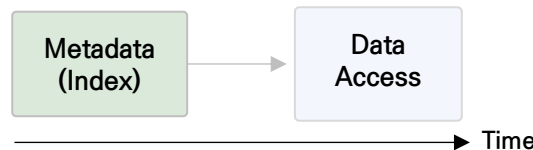
Pruning (sparse)



2

Accessed before actual data, so, it should be predictable and regular

Irregular metadata structures make prefetching difficult and introduce additional overhead.



- ✓ Must be accessed before data
- ✓ Size and number should be predictable
- ✓ Regularity enables efficient prefetching, scheduling and data access

Good (Predictable)

Fixed-size
Fixed-count per unit



- Easy to schedule
- Easy to prefetch

Bad (Irregular)

Varying-size
Varying-count per unit

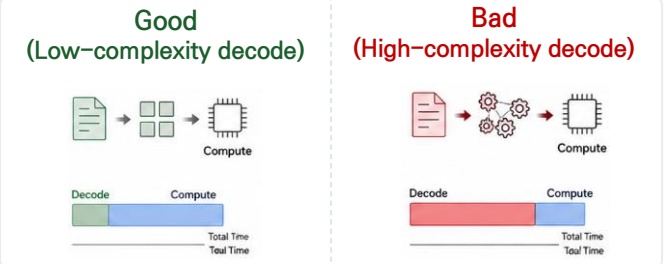


- Hard to schedule
- Unpredictable access

3

Metadata decoding should be simple

If the metadata requires excessive processing, the decoding overhead may become larger than the computation itself.



Decode overhead $\ll$ Compute

Decode overhead $\geq$ Compute

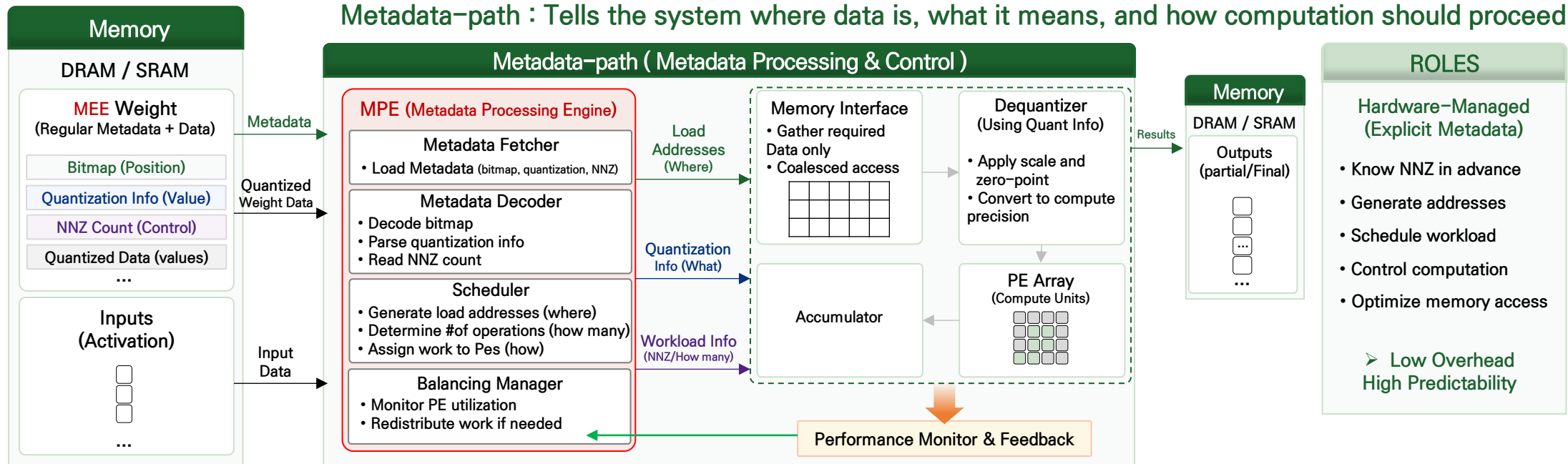
Rule of Thumb

- Simple and fast to decode
- Low memory access for decoding
- Limited branching and irregularity

MEE (with bitmap structure) is compact, fixed-size, predictable, and simple to decode

⇒ These principles extend beyond MEE, may serve as useful guidelines for designing metadata structures in many AI systems.

One Possible Metadata-Centric Accelerator Architecture



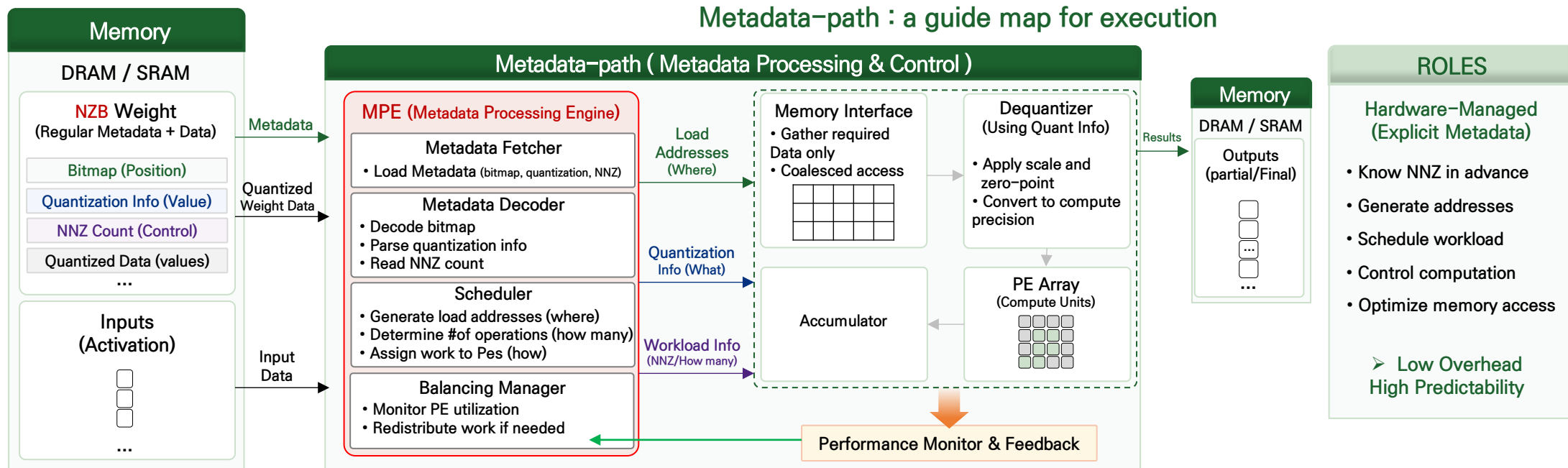
key idea is to separate metadata processing from data processing.

Instead of treating metadata as a passive byproduct, explicitly process metadata and use it to guide execution.

Metadata Processing Engine (MPE) reads metadata first, extracts useful information, and helps orchestrate execution

One Possible Metadata-Centric Accelerator Architecture

Metadata-Guided Execution Flow



Key Difference

Conventional
(CSR/COO, etc.)

Data-Driven, CPU-Controlled
Metadata is implicit or absent
→ CPU computes & controls everything

MEE
(Bitmap+Quant+NNZ)

Metadata-Driven, Hardware-Managed
Metadata is implicit or regular
→ Metadata guides & orchestrates execution

How Metadata Enables Predictable Execution

Predictable Memory
Access Patterns (Where)

- Bitmap → Know locations
- Only necessary data is loaded
- Coalesced access

Predictable Computation
(What / How many)

- Quant Info → Know precision
- NNZ Count → Know # of ops
- Workload known in advance

Predictable Scheduling
(How)

- Scheduler assigns work to PEs
- Balanced execution
- High utilization

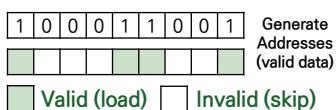
MEE Metadata-path Realizes

Performs CPU-like control-path
functions within the accelerator
hardware

Enables predictable, metadata-
driven high-efficiency acceleration

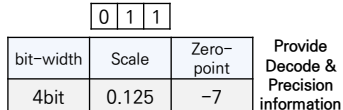
Metadata as a guide map for execution to help determine where to load, what to compute, and how many operations should be performed

1. Position (Bitmap) → Where



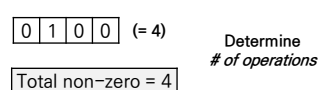
Reduce memory traffic & energy

2. Quantization Info → What



Correctly decode values
& select proper precision

3. Count (NNZ) → How many



Schedule workload & balance PEs

4. Execution Flow

- [1] MPE reads metadata
- [2] Generate load addresses & schedule
- [3] Load required data from memory
- [4] Decode, compute, accumulate in PEs
- [5] Write back results

Predictable, efficient, and high-utilization execution

Outcome

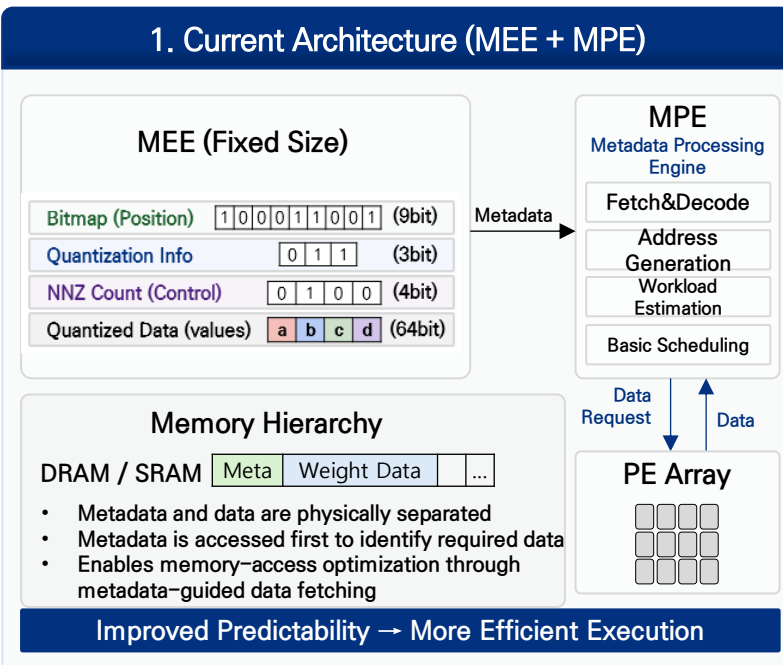
- High PE Utilization
- Low Memory Traffic
- Low Energy
- High Accuracy

Metadata helps determine where to load, what to compute, and how many operations should be performed,
→ **Predictable Execution with High Efficiency and Accuracy**

Baseline Metadata-Centric Execution Architecture

From MEE to MPU : Metadata-Centric Computing for Predictable and Efficient AI Acceleration

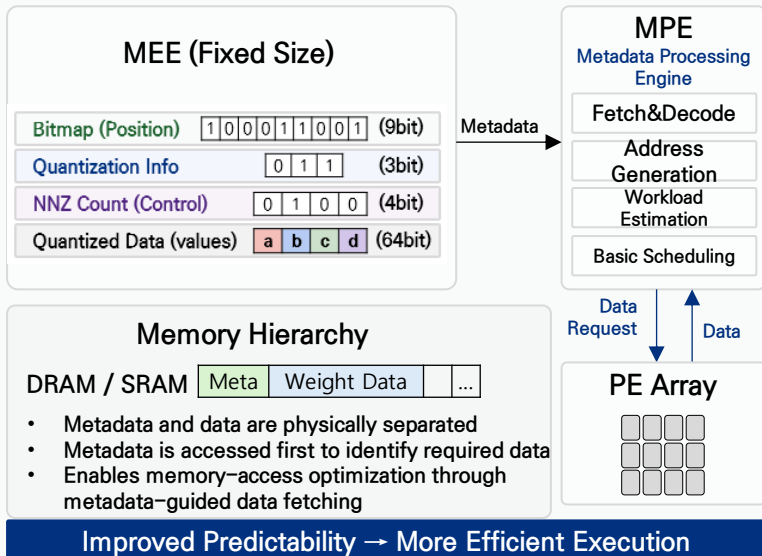
1. Current Architecture (MEE + MPE)



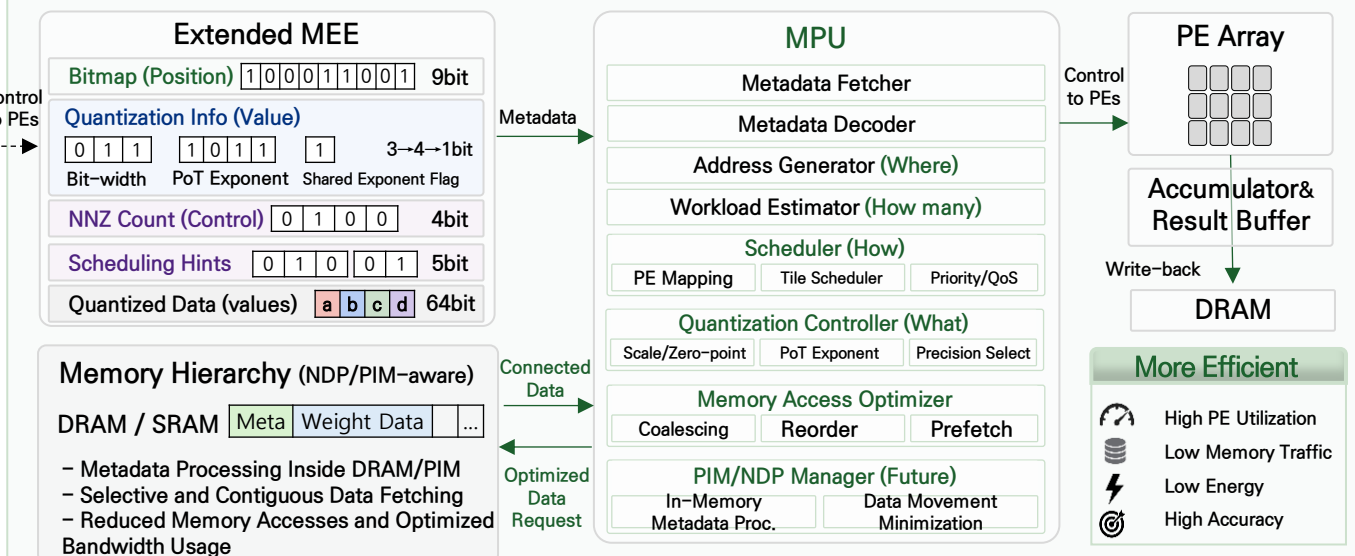
Toward Metadata-Centric Computing with MPU

From MEE to MPU : Metadata-Centric Computing for Predictable and Efficient AI Acceleration

1. Current Architecture (MEE + MPE)



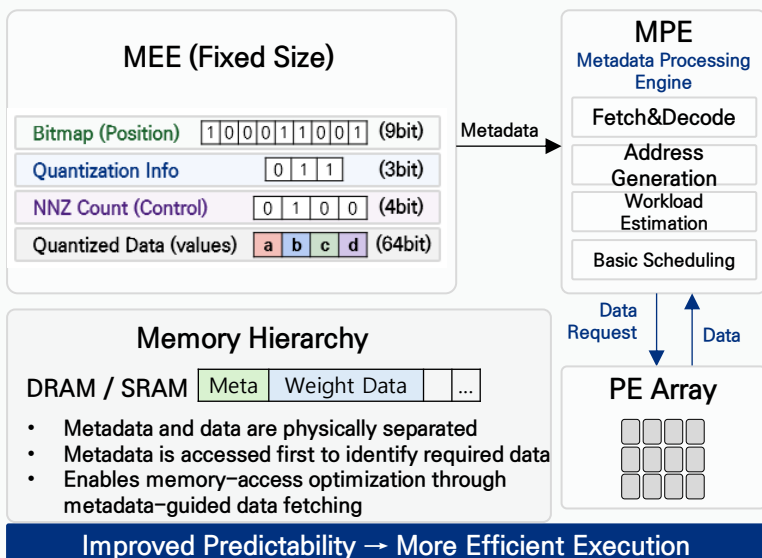
2. Future Direction : MPU (Metadata Processing Unit)



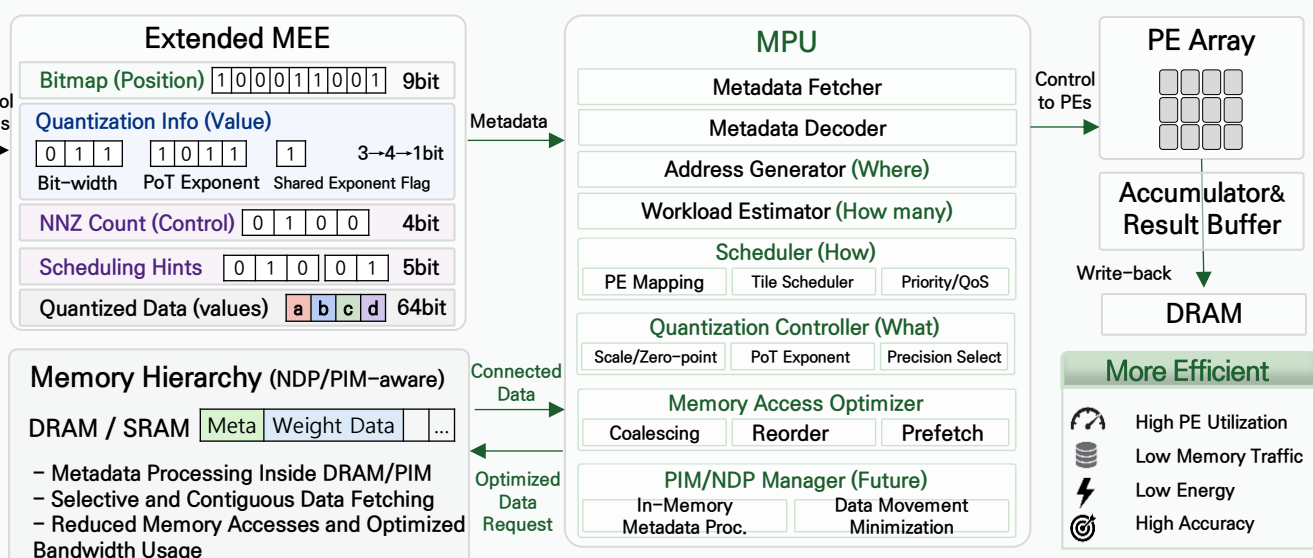
From Metadata Processing to Metadata-Centric Computing

From MEE → MPU → Metadata-Centric Computing: A new computing paradigm? – an open question

1. Current Architecture (MEE + MPE)



2. Future Direction : MPU (Metadata Processing Unit)



Why CSR is inefficient (vs. MEE)

Row Pointer	1	3		4	...	(M+1)
Column Index	0	1	2	2	...	(NNZ)
Values	a	b	c	d	...	(NNZ)

Problems

- No Prior Knowledge of Workload Size
- Poor Memory Access Locality
- Complex Metadata Processing
- Variable-Length Metadata Structure

MEE Advantages

- ✓ Fixed-size metadata enables efficient metadata access
- ✓ Explicitly stores request-related information such as NNZ
- ✓ Provides operation-related information, including quantization parameters
- ✓ Extensible Metadata with Technique-Specific Scheduling Hints

PoT-aware Quantization Support

Example (Pot Quantization)

Scale = 2^{-3}
(PoT Element = -3)

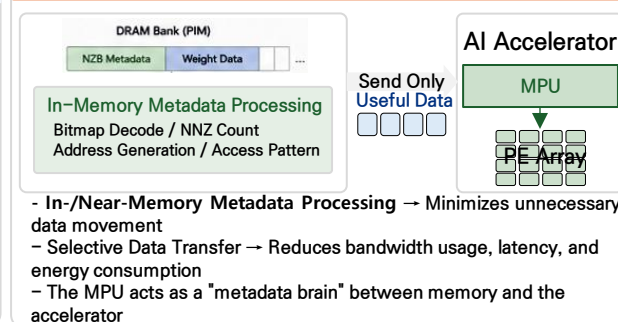
Multiply by 2^{-3}
= Shift Right by 3

NZB Value Field

Bit-width	0 1 1	(3 bit)
PoT Exponent	1 0 1 1	(-3)
Shared Exp Flag	1	(1 bit)

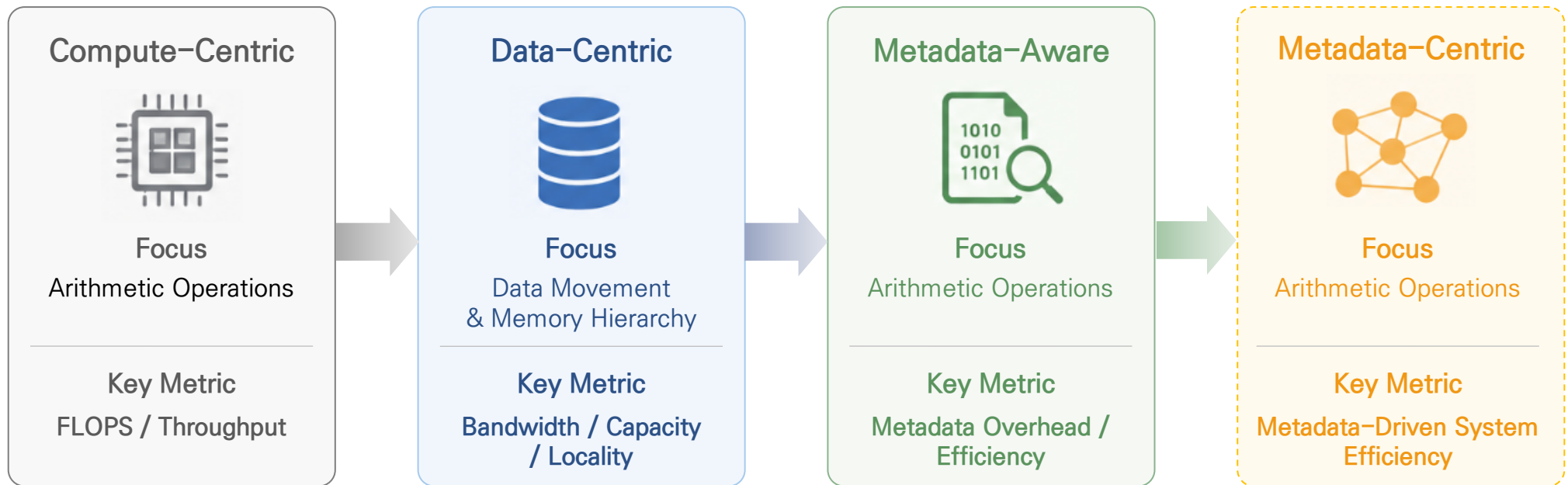
Easily extensible by adding new metadata fields for emerging techniques

Toward PIM / NDP with MPU



Metadata-Centric Computing as Next Opportunity for AI Systems?

Evolution of Computing Paradigms for AI Systems



As AI systems grow in scale and complexity, metadata increasingly determines
Where data is, what it means, and how computation should proceed.

Is Metadata-Centric Computing The Next Opportunity for AI Systems?

I plan to continue thinking about this direction and would appreciate your feedback.

From **optimizing computation and data movement** to unlocking the potential of **metadata**

Acknowledgements

- This work was supported by the grant of the National Research Foundation of Korea (NRF) (No. NRF-2021R1A2C2014557 and No. 2021R1A4A1031864) and Institute of Information & communications Technology Planning & Evaluation (IITP) (No. RS-2022-00167143)
- This research was supported by the Commercialization Promotion Agency for R&D Outcomes(COMPA) funded by the Ministry of Science and ICT(MSIT). (RS-2025-02311436, IP Enhancement and Commercialization for AI Model Lightweighting Technology to Promote the Business Expansion of Low-Cost AI Accelerators)

THANK YOU

