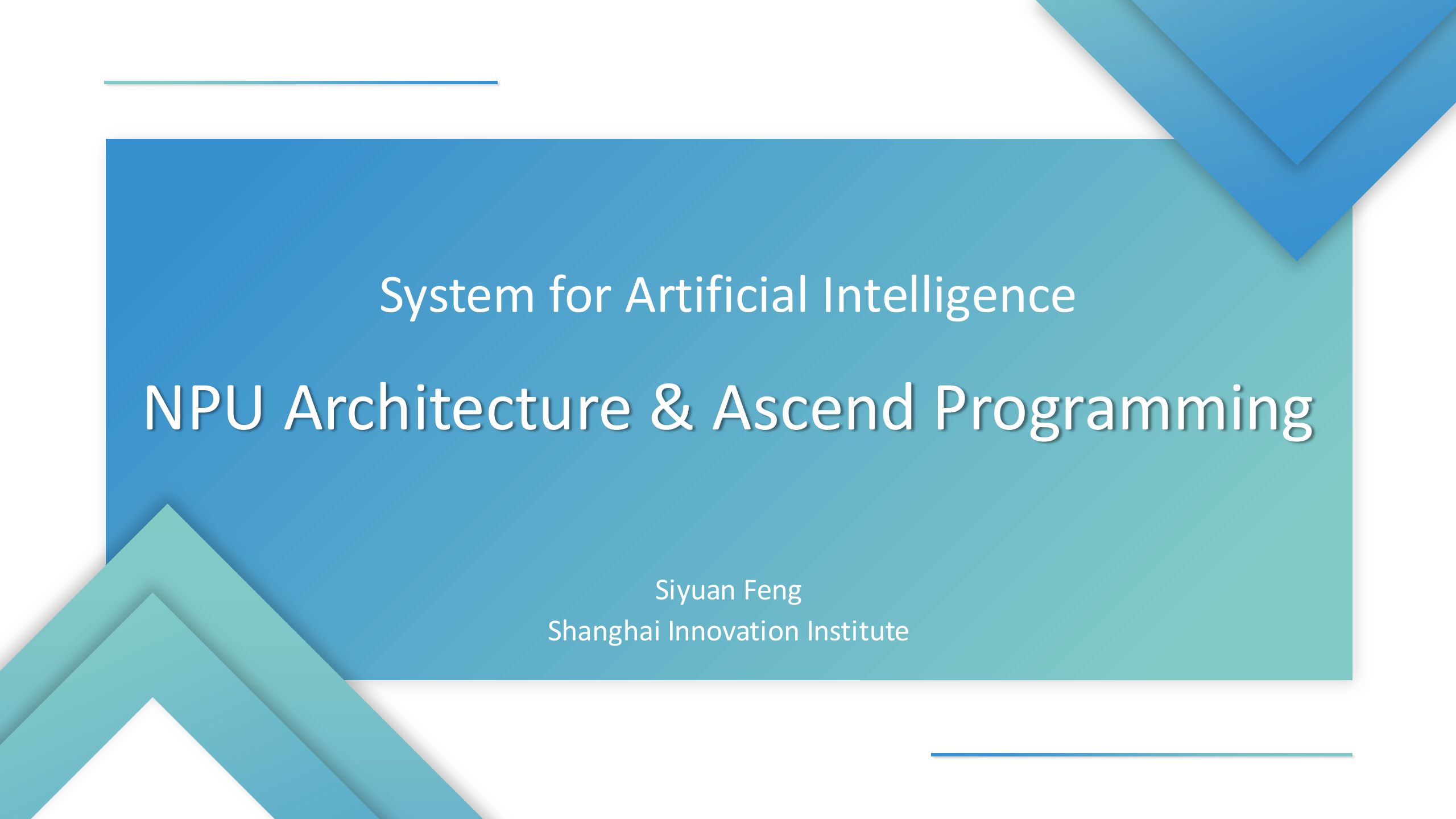




System for Artificial Intelligence

NPU Architecture & Ascend Programming

Siyuan Feng
Shanghai Innovation Institute

Recap: Overview of Machine Learning Systems



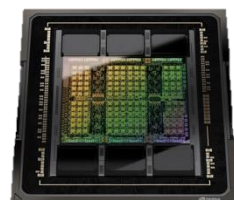
ML Models

Automatic Differentiation

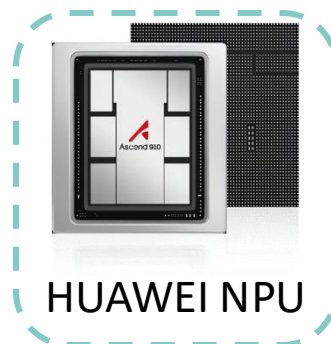
Graph Optimization

Parallelism / Distributed

Hardware Acceleration



NVIDIA GPU



HUAWEI NPU

This Lecture



Mobile devices

OUTLINE

- 01 ▶ HUAWEI NPU Architectures
- 02 ▶ Ascend C Programming

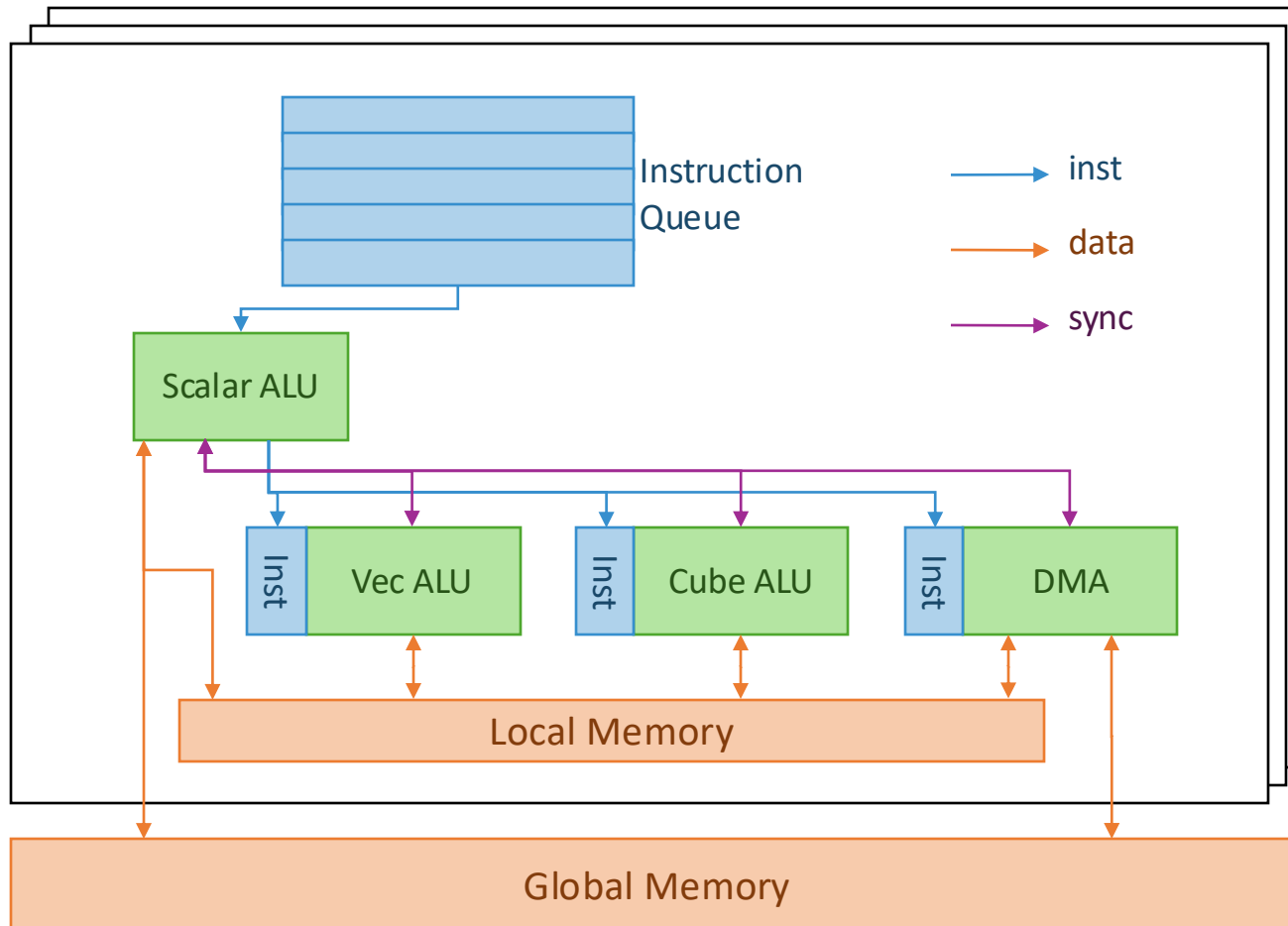


01



HUWEI NPU Architectures

NPU AI Core Architecture



Scalar Unit:

$$c = a + b$$

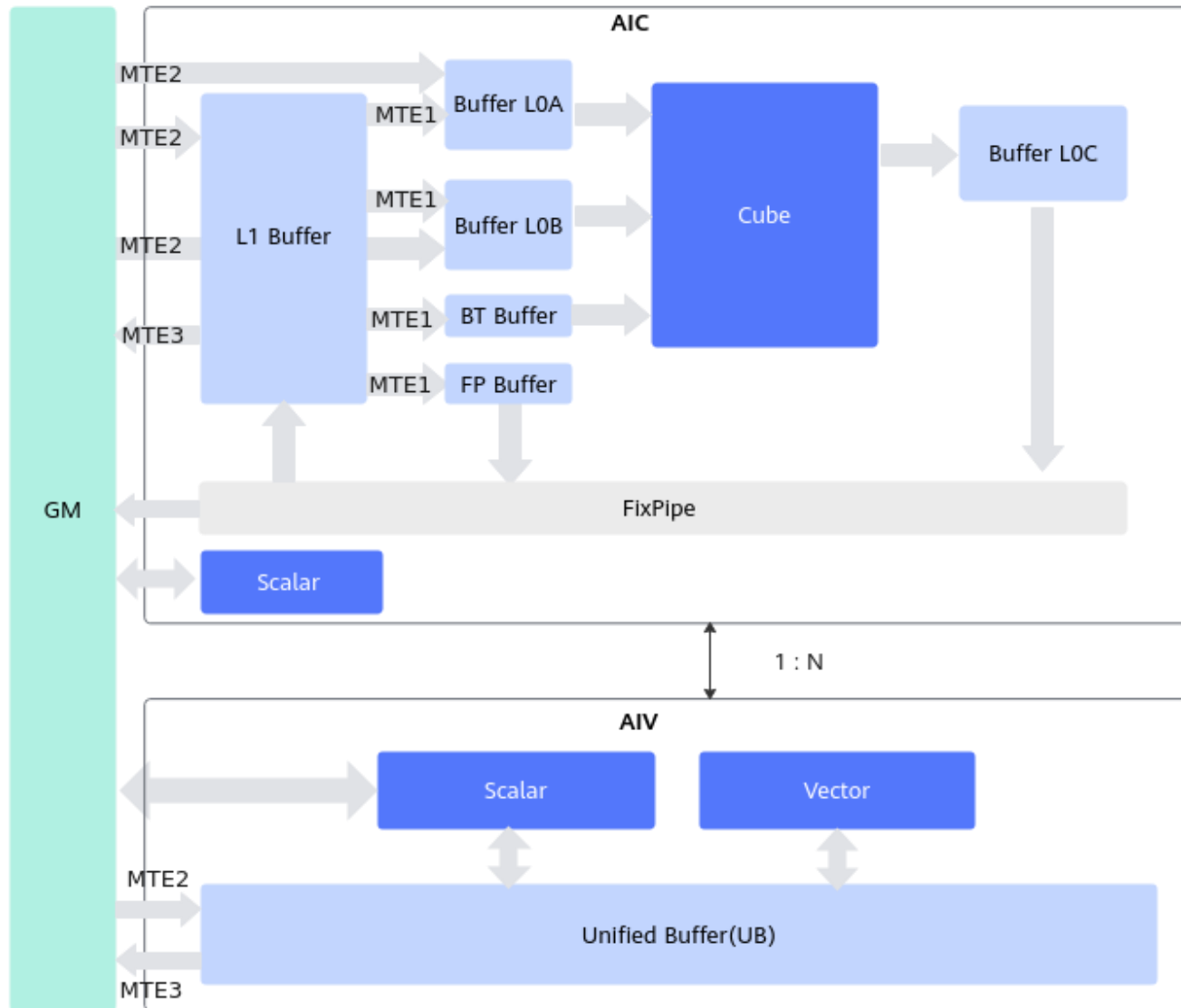
Vector Unit:

$$C[0:1024] = A[0:1024] + B[0:1024]$$

Cube Unit:

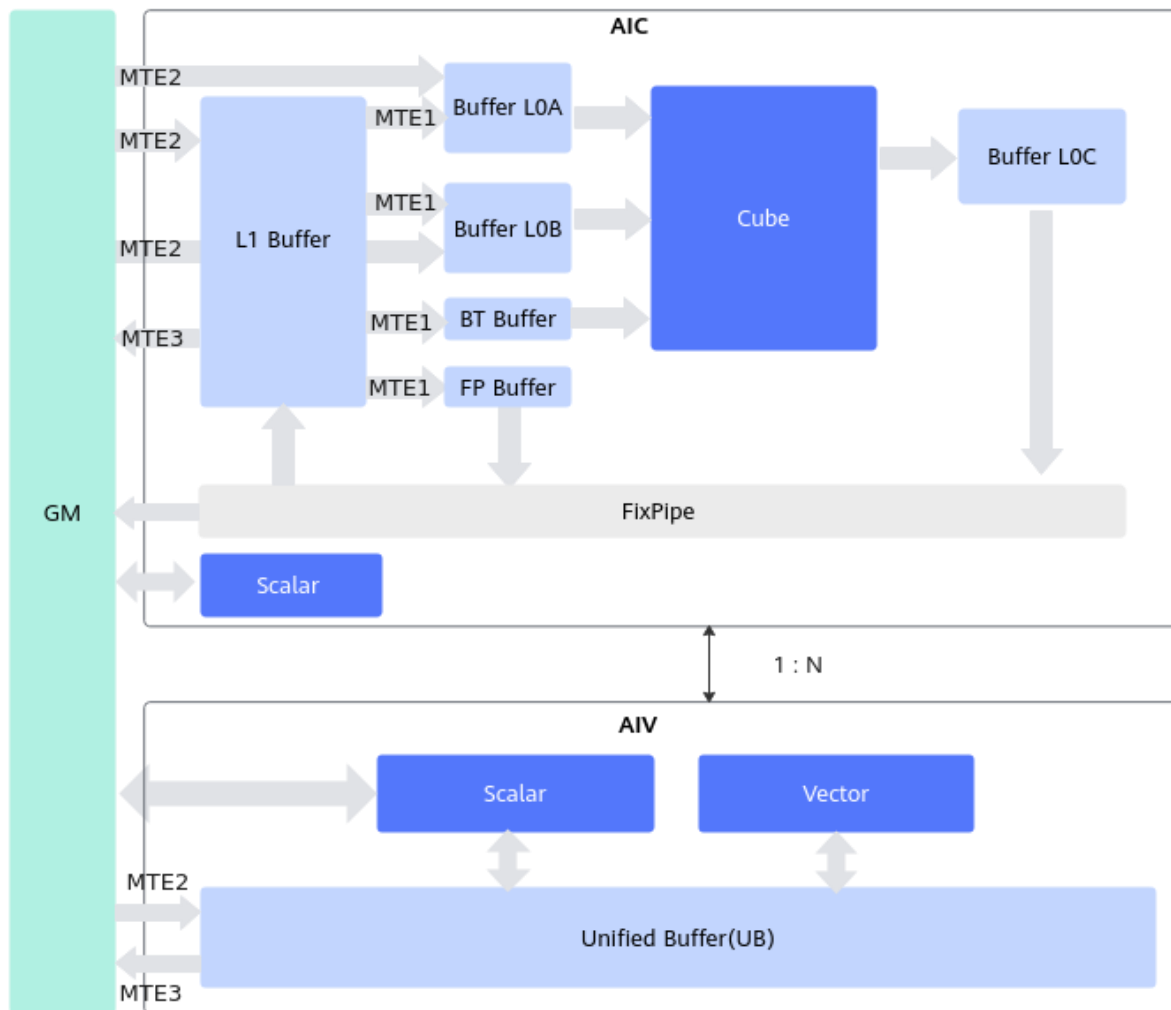
$$C[0:16, 0:16] += A[0:16, 0:16] * B[0:16, 0:16]$$

Disaggregated AI Core Architecture (910 B / C)



- Multi-level Memory Scope
- **Not "full-mesh"** access across ALU and memory

Discussion: comparison between NPU and GPU



HUAWEI NPU Architecture



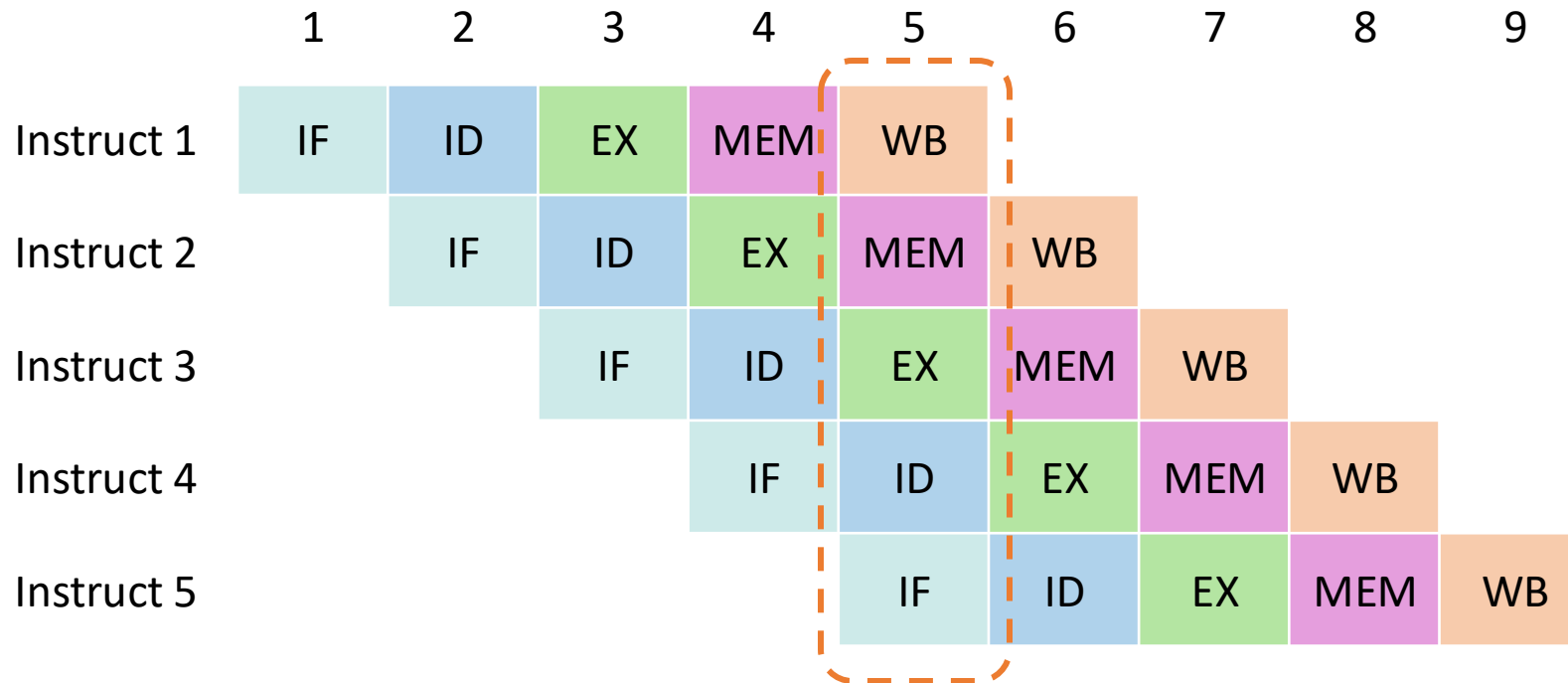
NVIDIA GPU Architecture



02

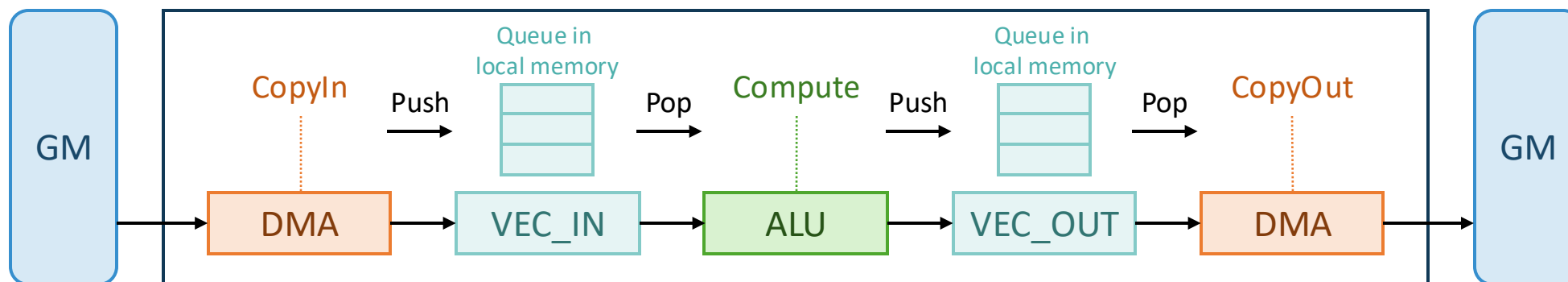
Ascend C Programming

Recap: Five-Stage Pipeline in CPU



Discussion: Why? Utilize different hardware components at the same time

Ascend C Programming for AIV



Three tasks:

1. **CopyIn:** load data from global memory to local memory and push to queue VEC_IN
2. **Compute:** load from queue VEC_IN, do computation, and then store to VEC_OUT
3. **CopyOut:** load data from queue VEC_IN, and store data to VEC_OUT

Pseudocode for AIV Programming

```
TQue<VecIn, 1> queIn;
TQue<VecOut, 1> queOut;

for-loop {
    // CopyIn Stage
    auto tensor = queIn.AllocTensor<half>();
    DataCopy(tensor, gm, len);
    queIn.Enqueue(tensor);

    // Compute Stage
    auto tensor = queIn.DeQueue<half>();
    auto tensorOut = queOut.AllocTensor<half>();
    OP(tensorOut, tensor, 1024);
    queIn.FreeTensor(tensor);
    queOut.Enqueue(tensorOut);

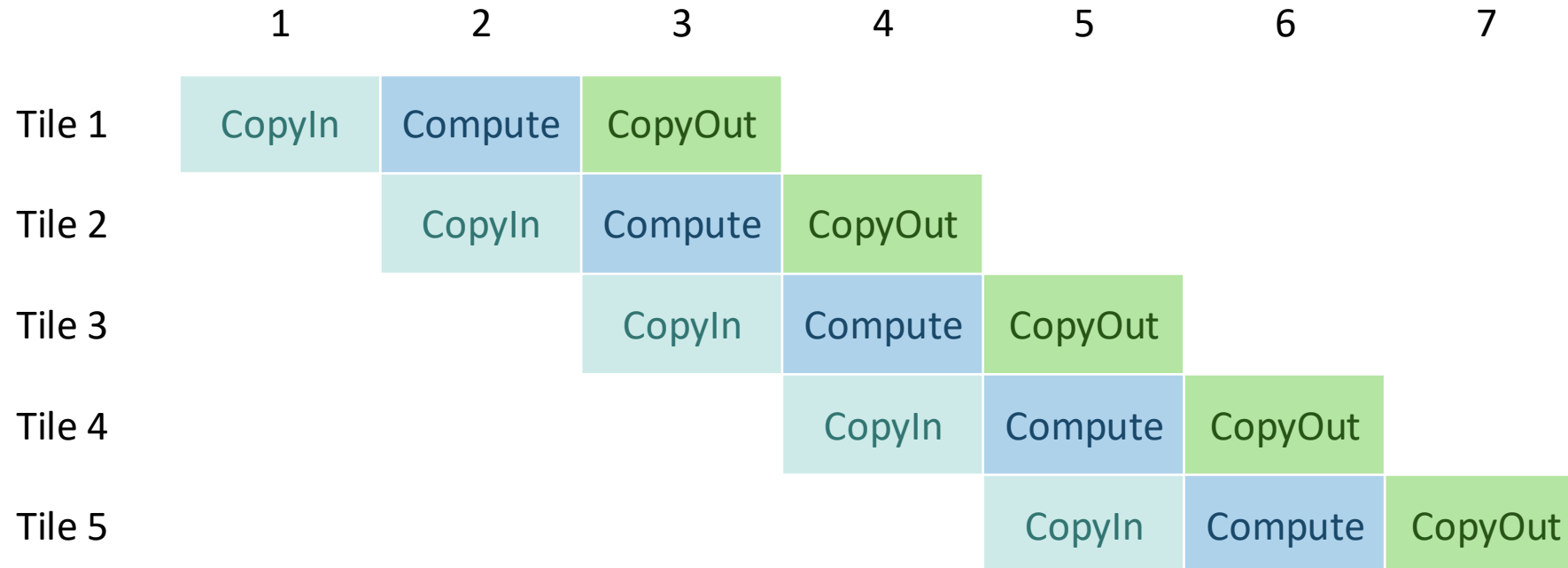
    // CopyOut Stage
    auto tensor = queOut.DeQueue<half>();
    DataCopy(gmOut, tensor, 1024);
    queOut.FreeTensor(tensor);
}
```

```
// Allocate tensor at local for input
// Copy data from gm to local
// push tensor to queIn after load

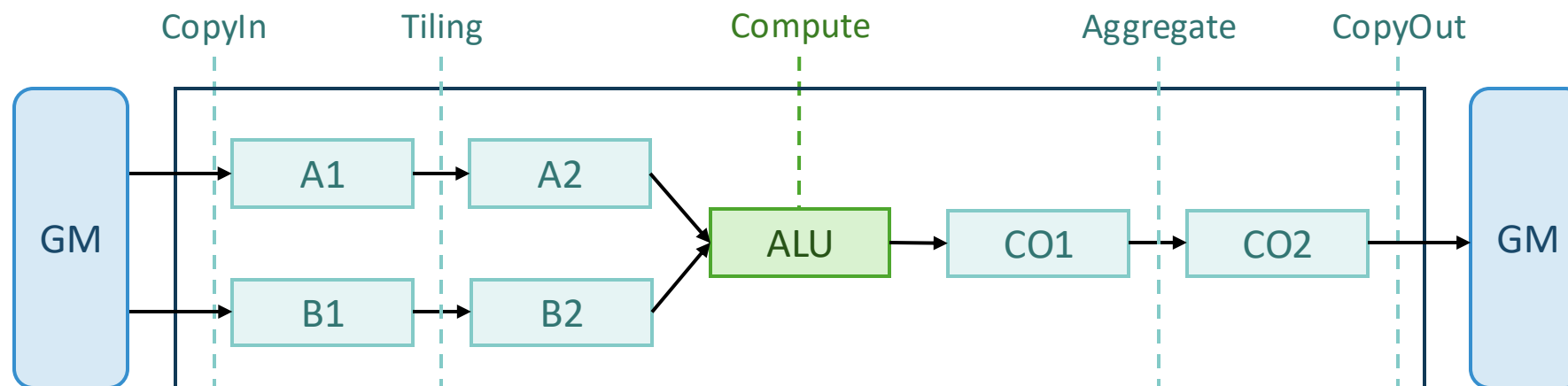
// Get tensor from queIn when ready
// Allocate tensor at local for output
// Main computation
// Free input tensor at local
// Push output tensor to queOut

// Get tensor from queIn when ready
// Copy data from local to gm
// Free output tensor at local
```

Pipeline of AIV



Ascend C Programming for AIC



A1: L2 cache for matrix A
A2: L1 cache for matrix A
B1: L2 cache for matrix B
B2: L1 cache for matrix B

CO1: Output buffer for tiled C
CO2: Output buffer for C

Pseudocode for AIC Programming

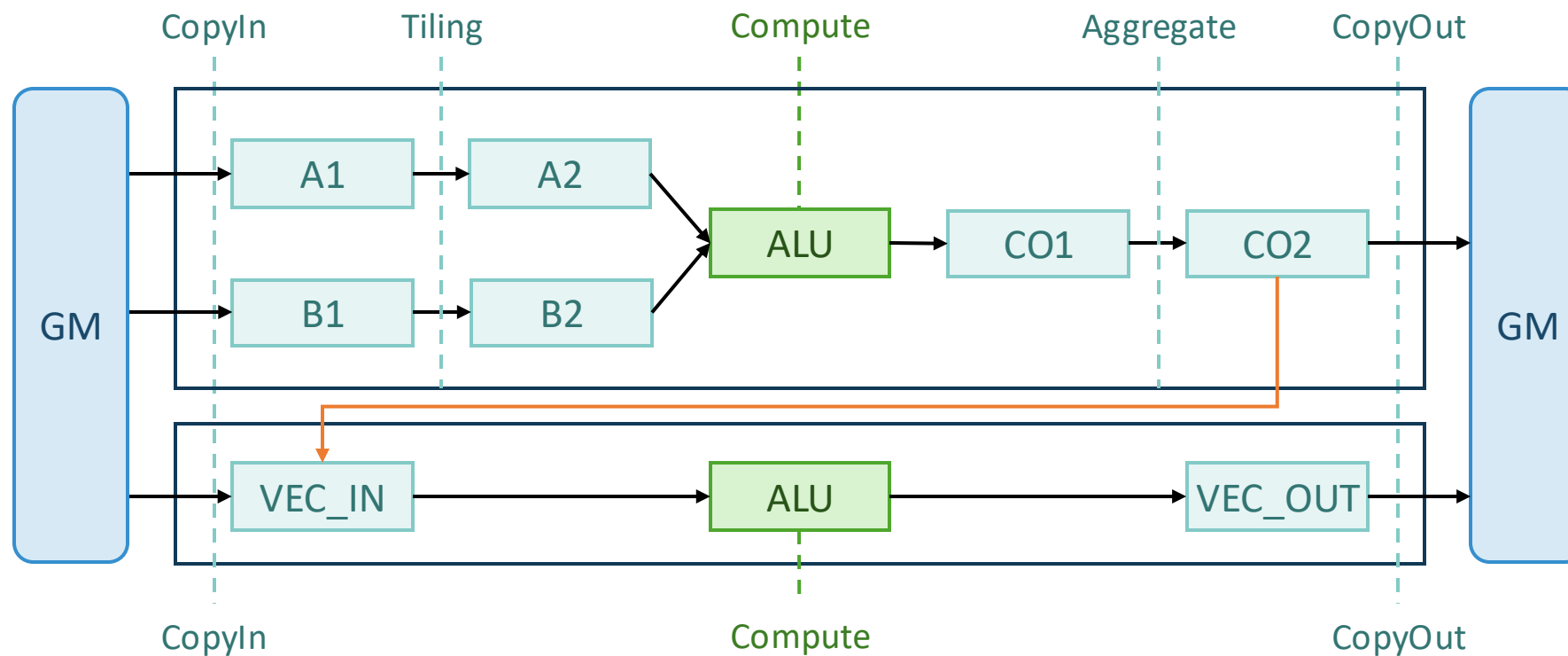
```
// Init
typedef MatmulType<TPosition::GM, CubeFormat::ND, half> aType;
typedef MatmulType<TPosition::GM, CubeFormat::ND, half> bType;
typedef MatmulType<TPosition::GM, CubeFormat::ND, float> cType;
typedef MatmulType<TPosition::GM, CubeFormat::ND, float> biasType;
Matmul<aType, bType, cType, biasType> mm;

REGIST_MATMUL_OBJ(&pipe, GetSysWorkSpacePtr(), mm, &tiling);

// CopyIn stage
mm.SetTensorA(gm_a);
mm.SetTensorB(gm_b);
mm.SetBias(gm_bias);

// Compute stage
while (mm.Iterate()) {
    // CopyOut Stage
    mm.GetTensorC(gm_c);
}
mm.End();
```

Ascend C Programming for Fused Operator



A1: L2 cache for matrix A
A2: L1 cache for matrix A
B1: L2 cache for matrix B
B2: L1 cache for matrix B

CO1: Output buffer for tiled C
CO2: Output buffer for C
VEC_IN: input buffer for AIV
VEC_OUT: output buffer for AIC

Pseudocode for Fused Operator

```
template<typename aType, typename bType, typename cType, typename biasType>
__aicore__ inline void MatmulLeakyKernel<aType, bType, cType, biasType>::Process() {
    // Step 1. Init
    REGIST_MATMUL_OBJ(&pipe, GetSysWorkSpacePtr(), matmulObj);
    matmulObj.Init(&tiling);
    matmulObj.SetTensorA(aGlobal);
    matmulObj.SetTensorB(bGlobal);
    matmulObj.SetBias(biasGlobal);

    // Step 2. iterate the tiling
    while (matmulObj.template Iterate<true>()) {
        // Step 3. Move the output to AIV
        reluOutLocal = reluOutQueue_.AllocTensor<cType>();
        matmulObj.template GetTensorC<true>(reluOutLocal, false, true);
        // Step 4. Compute on AIV
        AscendC::LeakyRelu(reluOutLocal, reluOutLocal, (cType)alpha, tiling.baseM * tiling.baseN);
        reluOutQueue_.EnQueue(reluOutLocal);
        // Step 5. Write result back to GM
        reluOutQueue_.DeQueue<cType>();
        ...
        AscendC::DataCopy(cGlobal[startOffset], reluOutLocal, copyParam);
        reluOutQueue_.FreeTensor(reluOutLocal);
    }
    matmulObj.End();
}
```

Programming Paradigm

- Utilize hardware by applying instruction pipeline
- Each stage of tasks:
 - Allocate local memory or get local memory from upstream queue
 - Compute or copy data
 - Push processed data to downstream queue
 - Free useless local memory

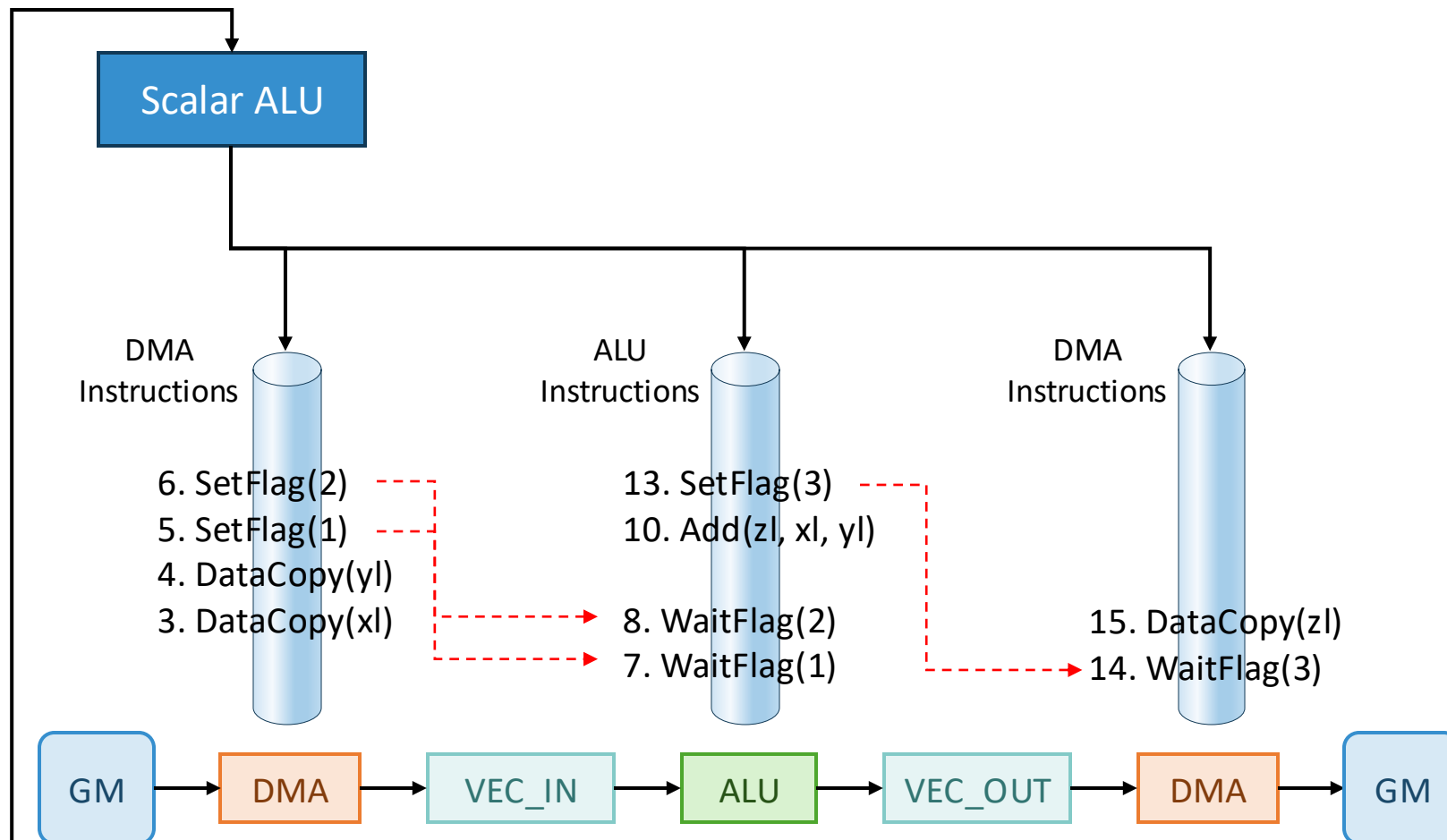
Appendix A: The Processing Flow of Enqueue / Dequeue

1. Scalar unit read the instruction
2. The instructions are dispatched to the corresponding unit's queues
3. The various execution units execute these instruction in parallel.
4. Enqueue/Dequeue resolves the **Write-After-Read (WAR) hazard** for memory:
 1. An **Enqueue** call issues a synchronization instruction (**set**), sending a signal to activate the wait.
 2. A **Dequeue** call issues a synchronization instruction (**wait**), waiting for the data write to complete.
 3. The wait must wait for the set signal before it can execute; otherwise, it is blocked.

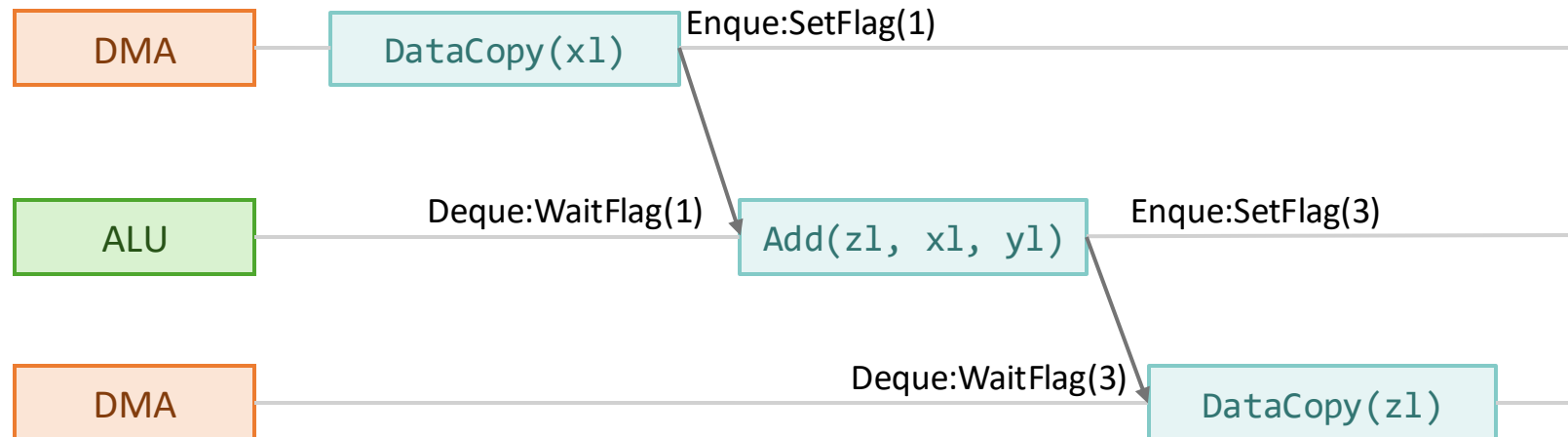
Appendix A: The Processing Flow of Enque / Deque

Operator Instruction Sequence

16. z.FreeTensor(zl)
15. DataCopy(zl)
14. zl = z.Deque()
13. z.Enqueue(zl)
12. y.FreeTensor(y1)
11. x.FreeTensor(x1)
10. Add (zl, x1, y1)
9. zl = z.AllocTensor()
8. y1 = y.Deque()
7. x1 = x.Deque()
6. y.Enqueue(y1)
5. x.Enqueue(x1)
4. DataCopy(y1)
3. DataCopy(x1)
2. y1 = y.AllocTensor()
1. x1 = x.AllocTensor()



Appendix A: The Processing Flow of Enqueue / Dequeue



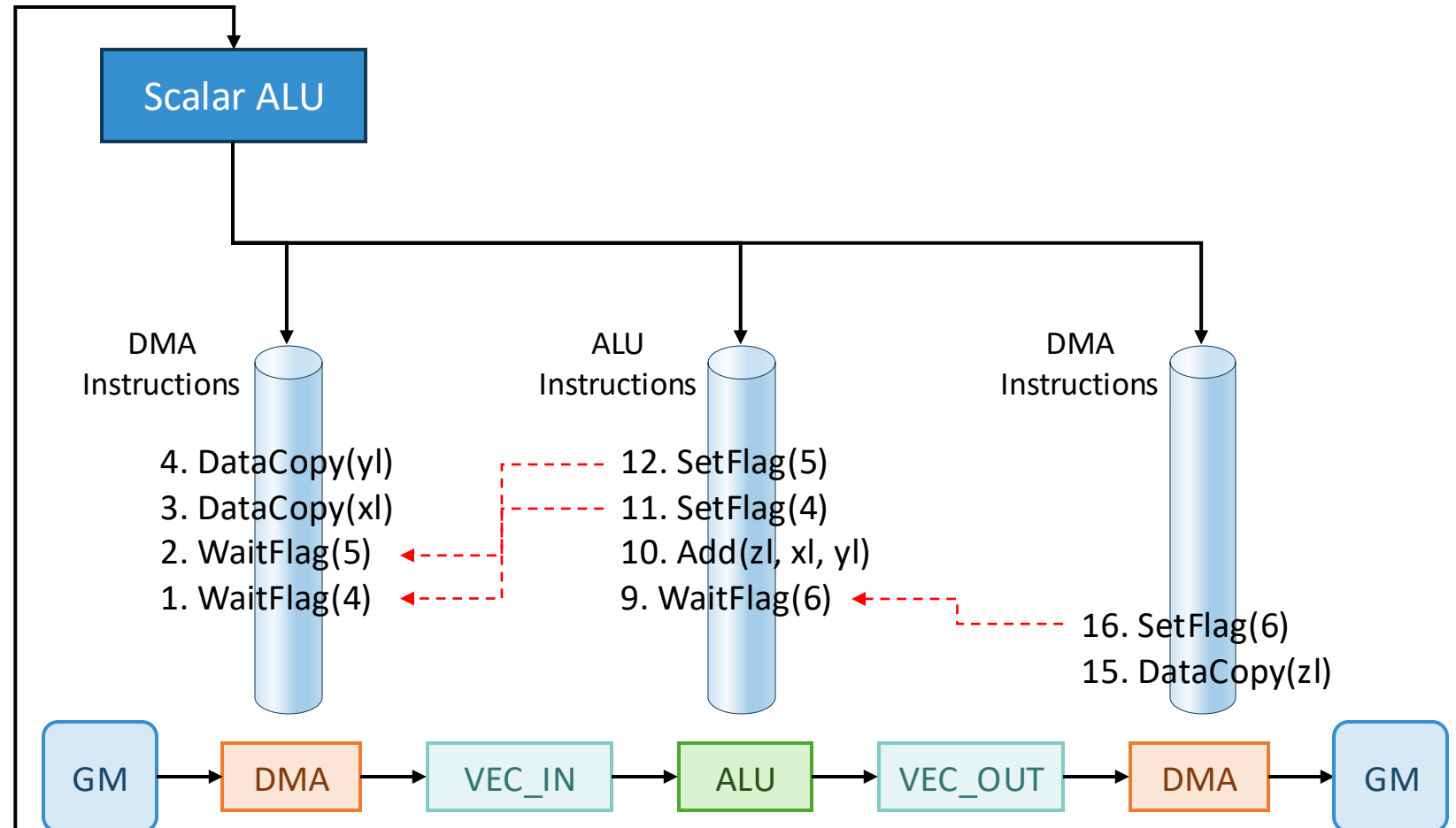
Appendix B: The Processing Flow of AllocTensor / FreeTensor

1. Scalar unit read the instruction
2. The instructions are dispatched to the corresponding unit's queues
3. The various execution units execute these instruction in parallel.
4. AllocTensor/FreeTensor resolves the **Read-After-Write (RAW)** hazard:
 1. An **AllocTensor** call issues a synchronization instruction (**wait**), waiting for the memory to finish being read.
 2. A **FreeTensor** call issues a synchronization instruction (**set**), signaling that the memory is released and can be repeatedly written to.
 3. The wait must wait for the set signal before it can execute; otherwise, it is blocked.

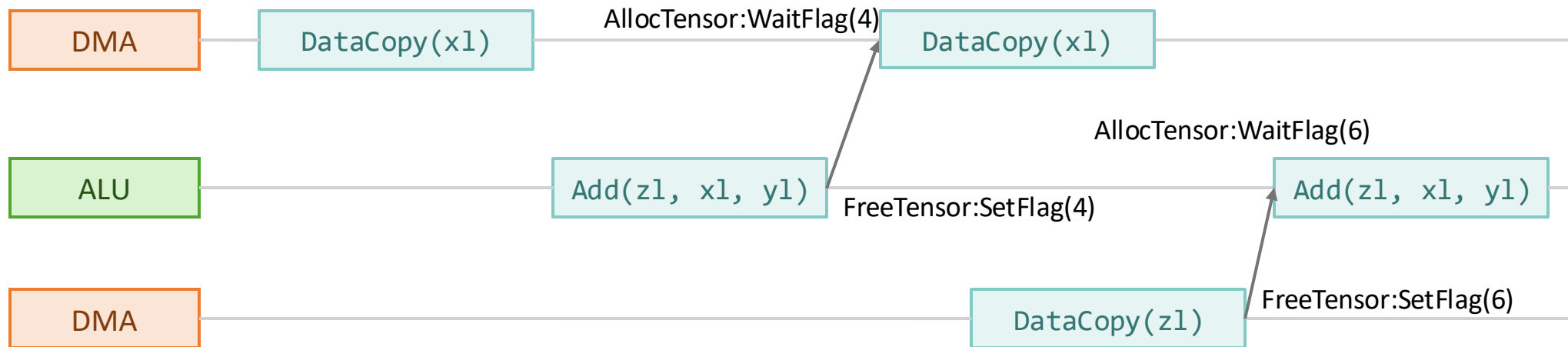
Appendix B: The Processing Flow of AllocTensor / FreeTensor

Operator Instruction Sequence

16. z.FreeTensor(zl)
15. DataCopy(zl)
14. zl = z.Dequeue()
13. z.Enqueue(zl)
12. y.FreeTensor(y1)
11. x.FreeTensor(x1)
10. Add (zl, x1, y1)
9. zl = z.AllocTensor()
8. y1 = y.Dequeue()
7. x1 = x.Dequeue()
6. y.Enqueue(y1)
5. x.Enqueue(x1)
4. DataCopy(y1)
3. DataCopy(x1)
2. y1 = y.AllocTensor()
1. x1 = x.AllocTensor()



Appendix B: The Processing Flow of AllocTensor / FreeTensor



Acknowledgement

The development of this course, including its structure, content, and accompanying presentation slides, has been significantly influenced and inspired by the excellent work of instructors and institutions who have shared their materials openly. We wish to extend our sincere acknowledgement and gratitude to the following courses, which served as invaluable references and a source of pedagogical inspiration:

- Machine Learning Systems[15-442/15-642], by **Tianqi Chen** and **Zhihao Jia** at **CMU**.
- Advanced Topics in Machine Learning (Systems)[CS6216], by **Yao Lu** at **NUS**

While these materials provided a foundational blueprint and a wealth of insightful examples, all content herein has been adapted, modified, and curated to meet the specific learning objectives of our curriculum. Any errors, omissions, or shortcomings found in these course materials are entirely our own responsibility. We are profoundly grateful for the contributions of the educators listed above, whose dedication to teaching and knowledge-sharing has made the creation of this course possible.



System for Artificial Intelligence

Thanks

Siyuan Feng
Shanghai Innovation Institute