

Towards **Compute-Aware In-Switch Computing** for LLMs Tensor-Parallelism on Multi-GPU Systems

Chen Zhang¹, Qijun Zhang^{1*}, Zhuoshan Zhou¹, Yijia Diao¹,
Haibo Wang², Zhe Zhou², Zhipeng Tu², Zhiyao Li²,
Guangyu Sun³, Zhuoran Song¹, Zhigang Ji¹, Jingwen Leng^{1*}, Minyi Guo

¹Shanghai Jiao Tong University, ²Huawei Technologies Co. Ltd., ³Peking University

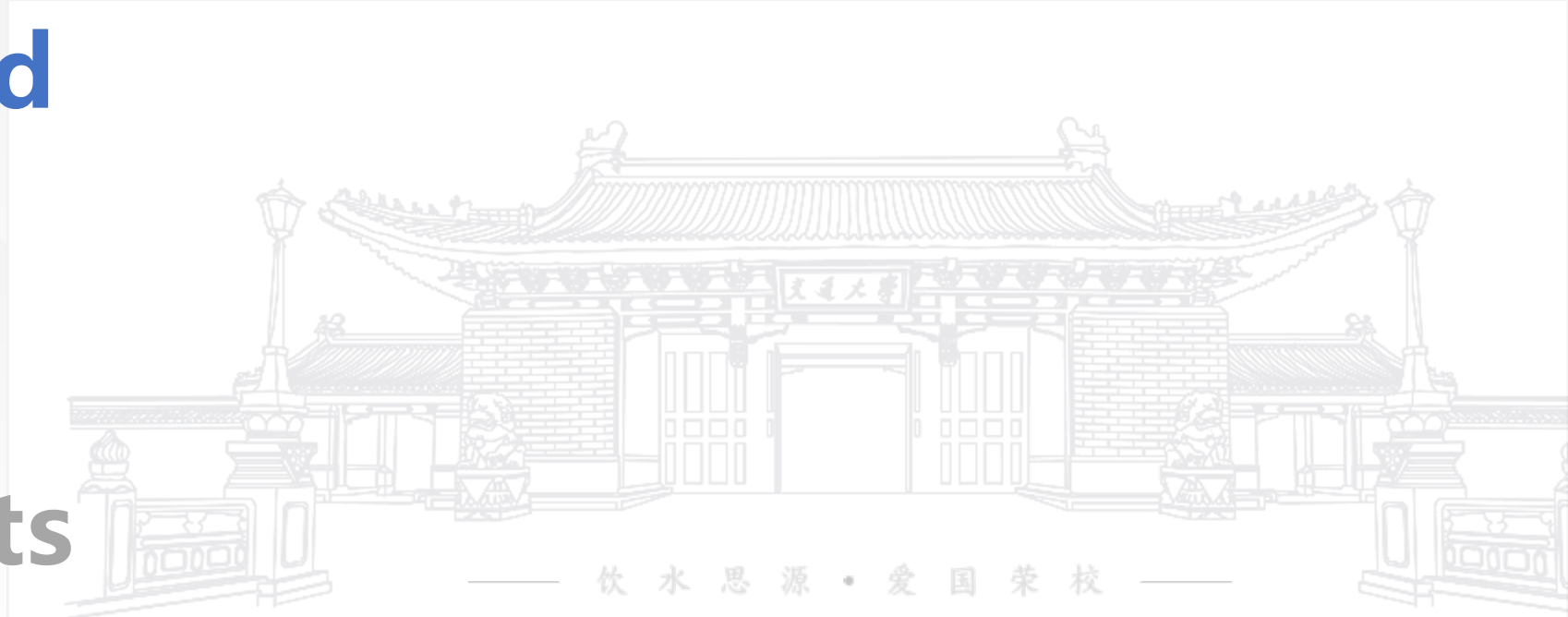
1. Background

2. Motivation

3. Solutions

4. Experiments

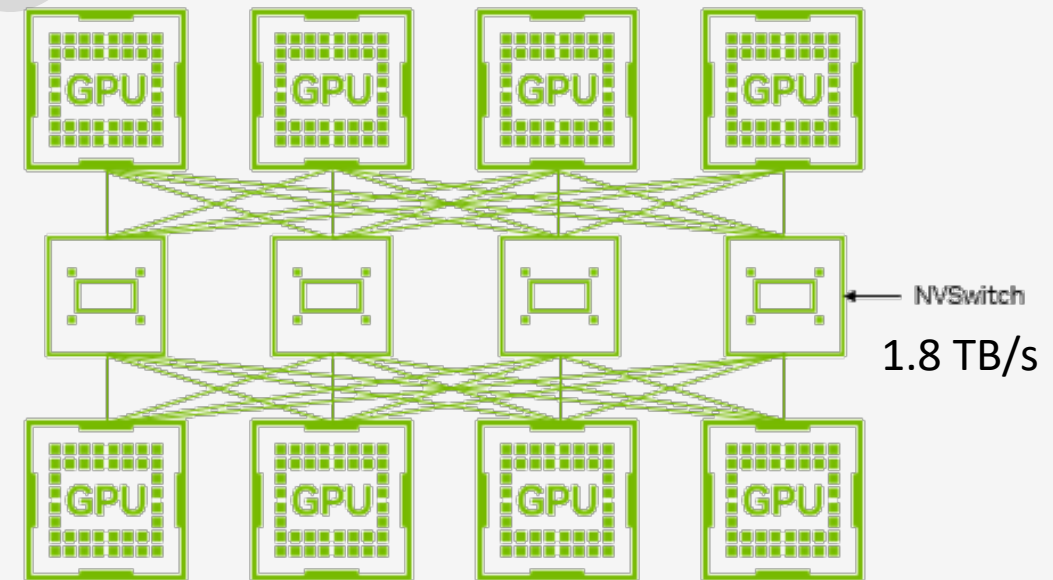
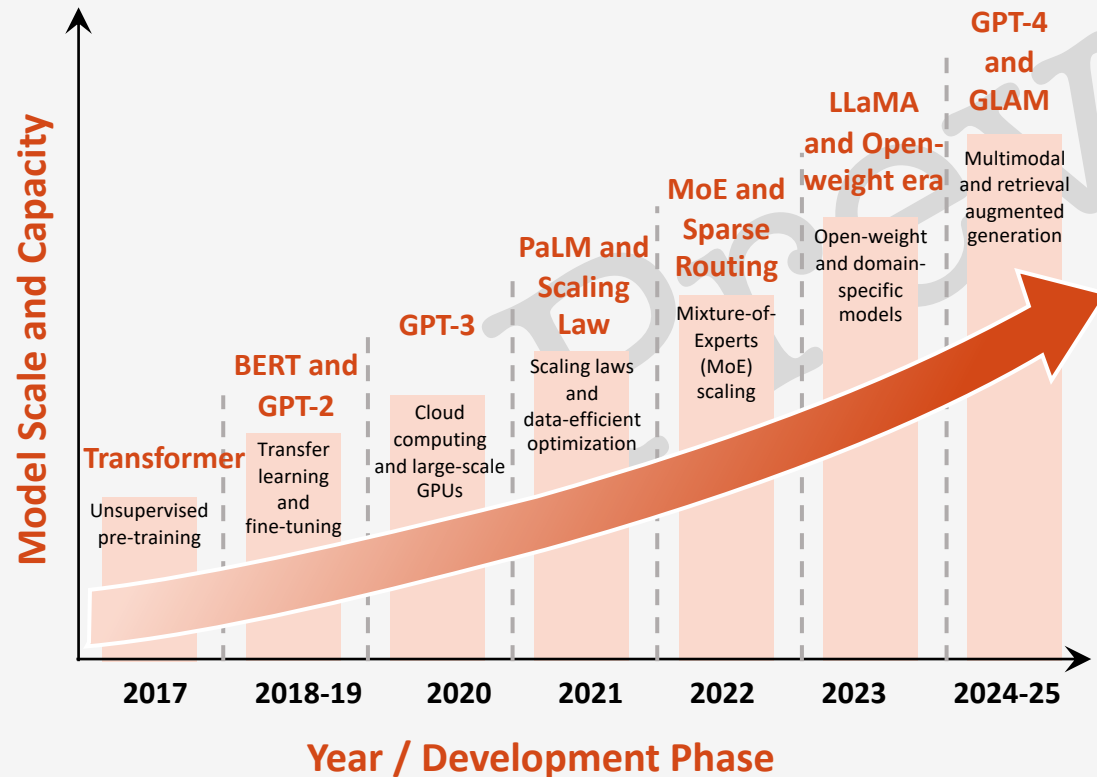
5. Conclusion



LLM Scaling Laws Drive Multi-GPU System Scale-Up

3

- LLM scaling laws incentivize ever-larger models, pushing parameter growth beyond single-GPU capacity and motivating tightly coupled multi-GPU nodes interconnected by NVLink/NVSwitch.



Tensor Parallelism Contributes to 97% Traffic

4

- Unlike DP and PP, TP incurs collective communication at every forward and backward pass, e.g., All-Reduce (basic TP) or Reduce-Scatter / All-Gather (TP+SP), making its communication overhead the dominant bottleneck as models and GPU counts scale.

Communication Volume and Bandwidth Requirements

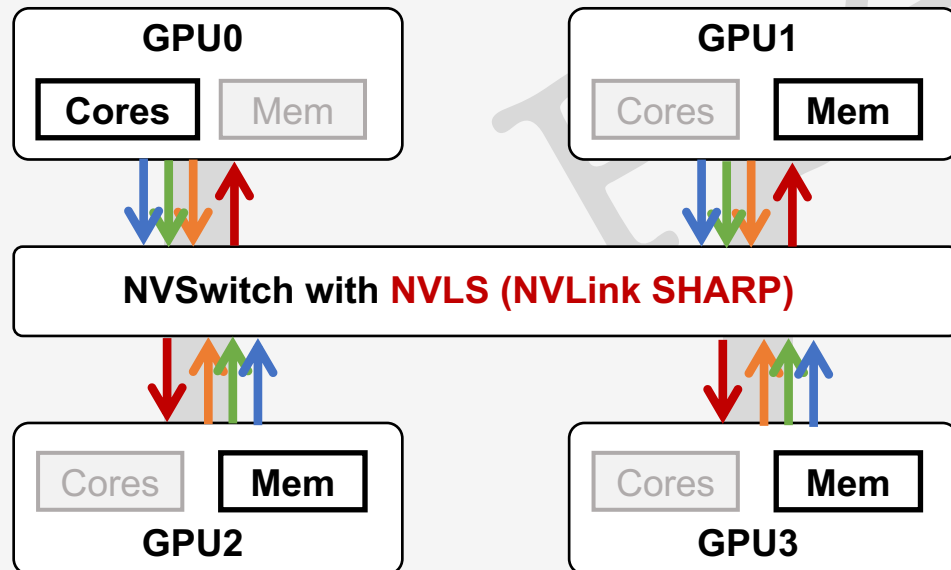
Parallelism	Traffic per Round (MB)	# Rounds	Total Traffic (GB)
PP	96	120	11
DP	1277	32	40
EP	6	7680	42
TP	180	7680	1350
SP	762	11520	8573

97%

NVLink SHARP (a.k.a. NVLS)

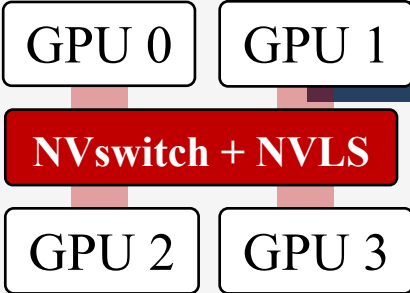
5

- **Dedicated to NVSwitch-based Multi-GPU Systems**
 - Introduced since Hopper Architecture
- **Native, instruction-level support for in-switch operations**
 - **multimem.*** instructions via **inline PTX** to implement the corresponding collectives
- **2–8× speedups for collectives compared to GPU-driven implementations**



Instructions	Accelerated Collectives
multimem.st	AllGather
multimem.ld_reduce	Reduce-Scatter
multimem.red	AllReduce

CAIS: Computation-aware In-Switch Computing

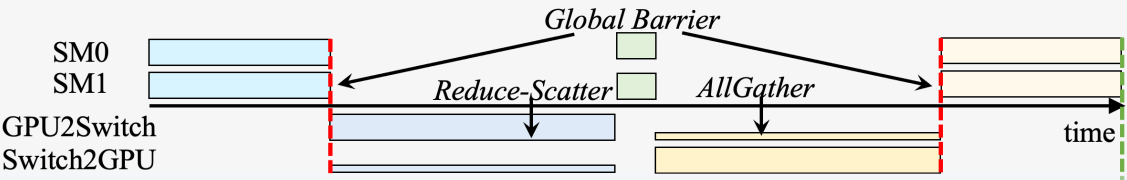


Existing NVLS design is
Communication Centric

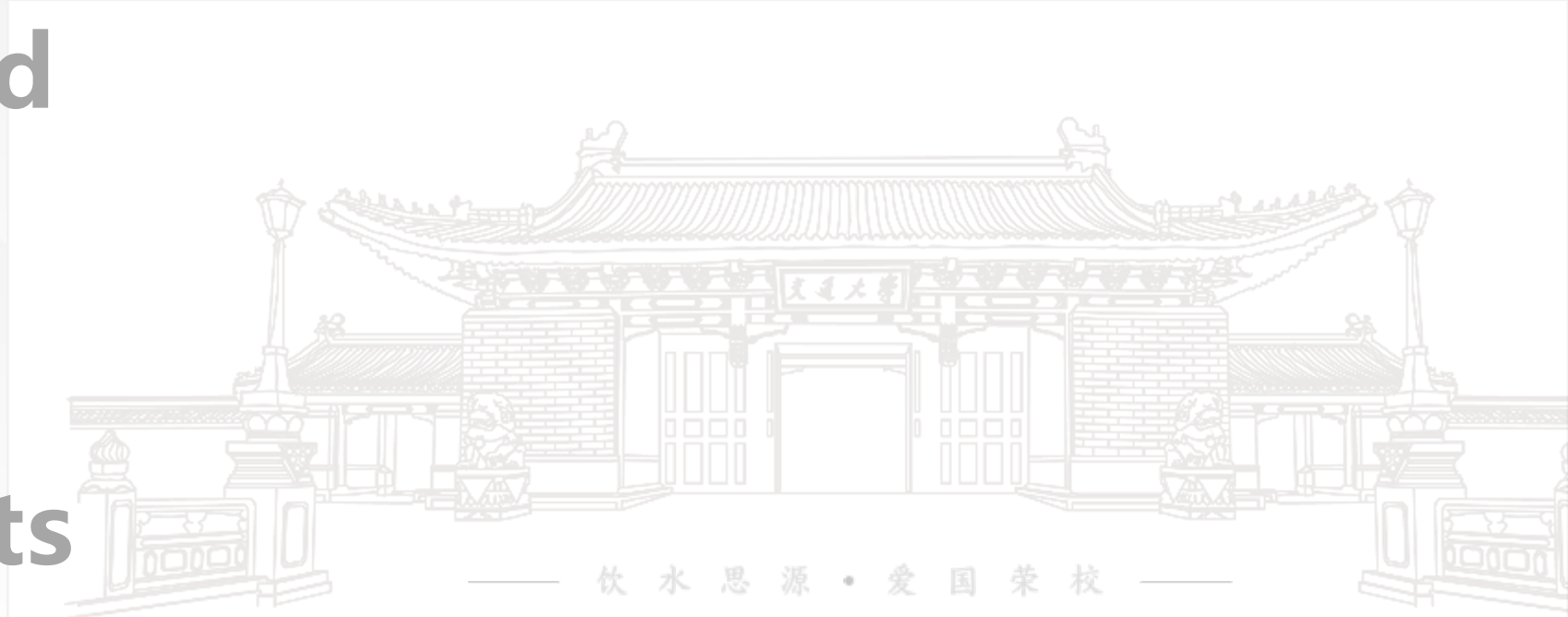
Our proposal:
Computation Aware

Instructions	Accelerated Collectives
multimem.st	AllGather
multimem.ld_reduce	Reduce-Scatter
multimem.red	AllReduce

Tensor Parallel Operations	
TP with SP	AllGather + GEMM
	GEMM + ReduceScatter
Basic TP	AllReduce + GEMM
	GEMM + AllReduce



1. Background
- 2. Motivation**
3. Solutions
4. Experiments
5. Conclusion



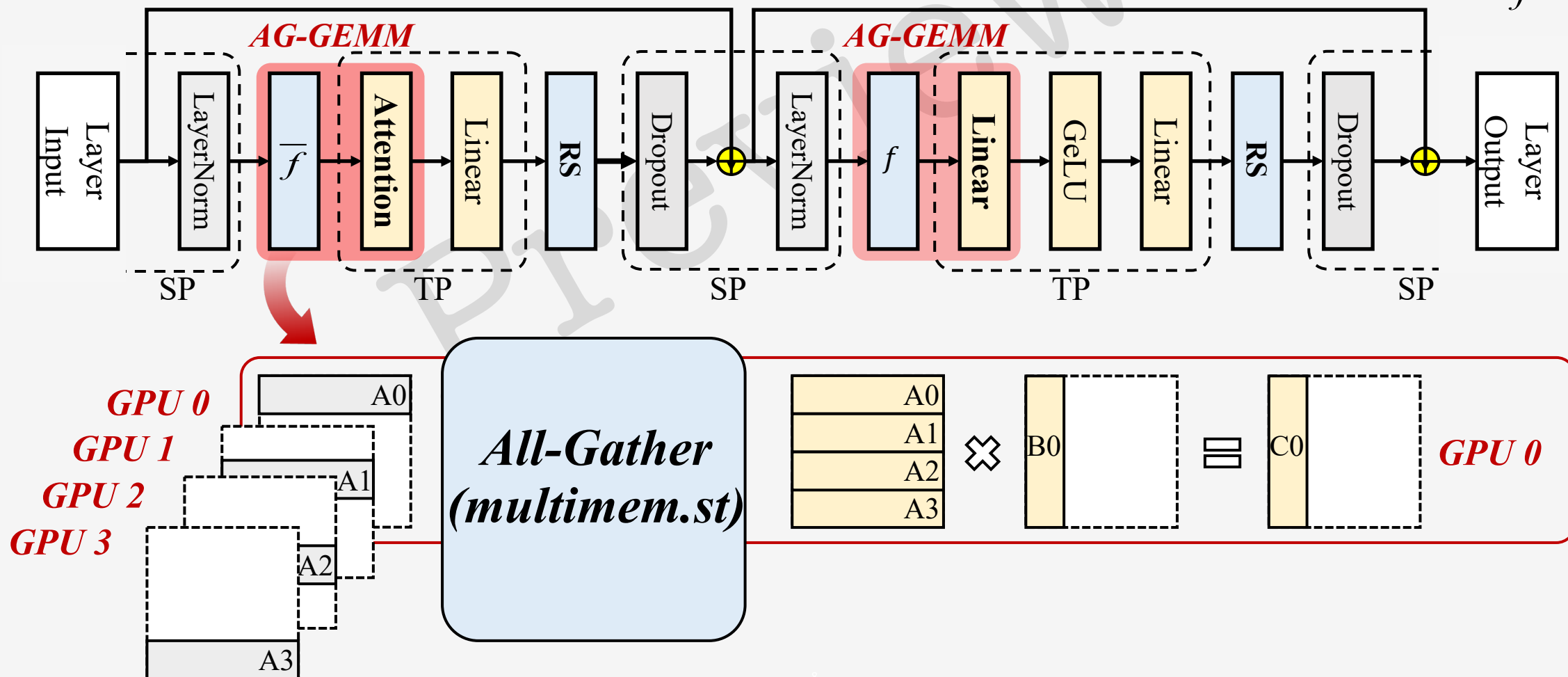
AllGather + GEMM Example

8

- All-Gather is often followed by a GEMM (e.g., attention or FFN), which requires reading data from both local and remote devices.

f :forward

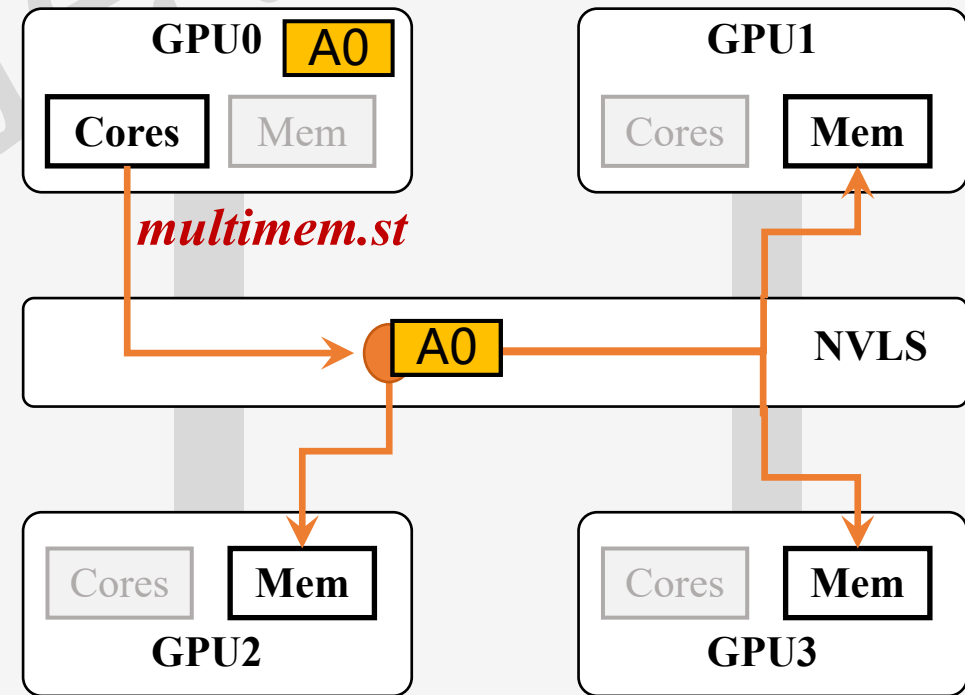
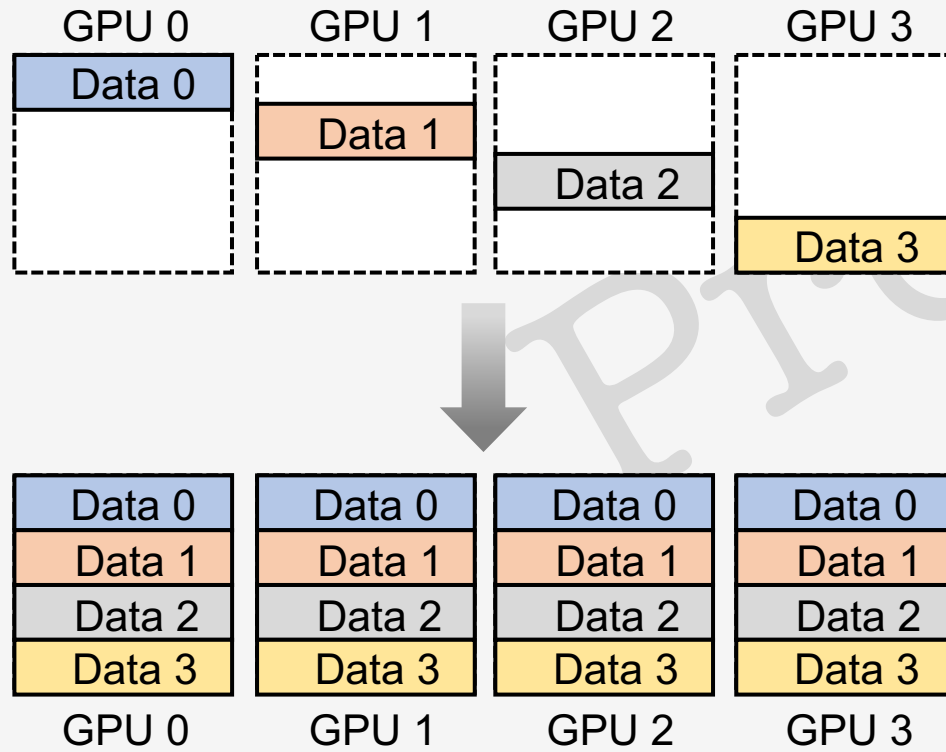
$\overline{f}$:backward



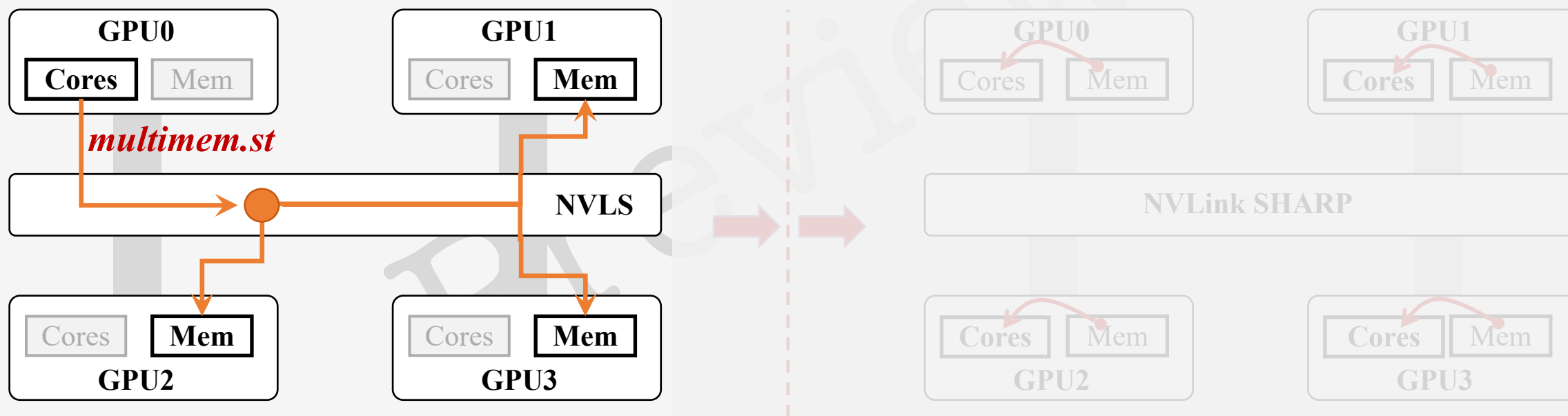
AllGather works in **PUSH** mode

9

- Each GPU pushes its local data to other GPUs via the *multimem.st* instruction
 - Inline PTX for in-switch computing function
 - Automatic Data Duplication/Reduction in the Switch



All-Gather + *GEMM*: remote *store* $\rightarrow$ *syn.* $\rightarrow$ local *read*

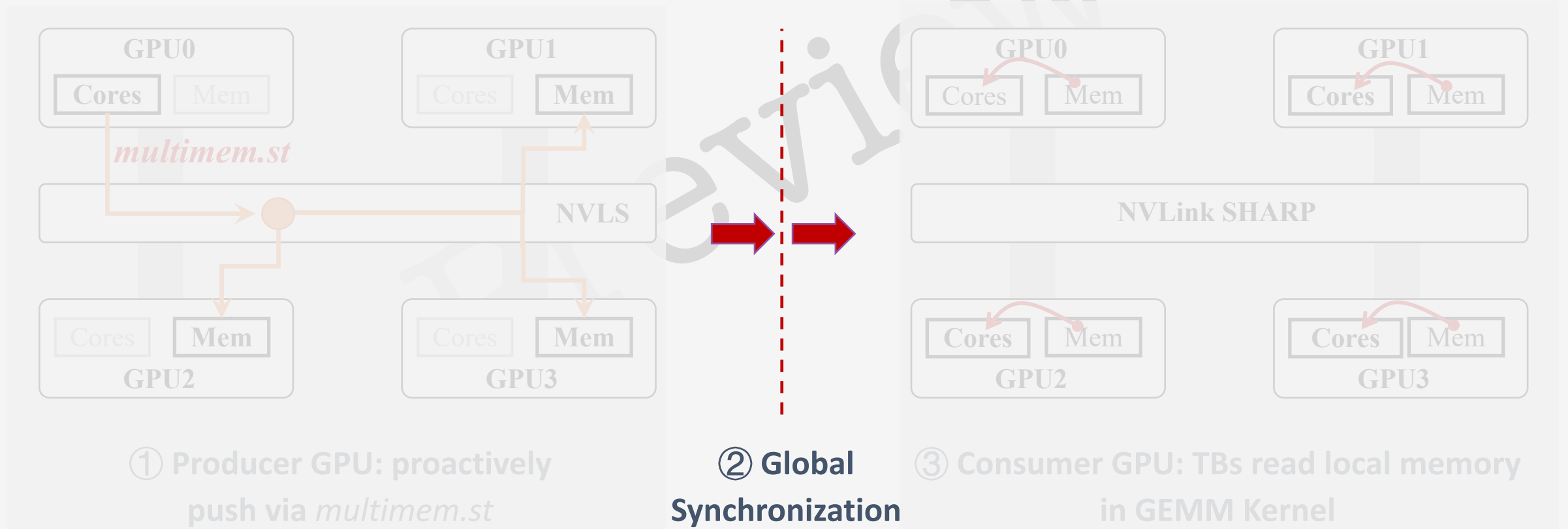


① **Producer GPU: proactively push via *multimem.st***

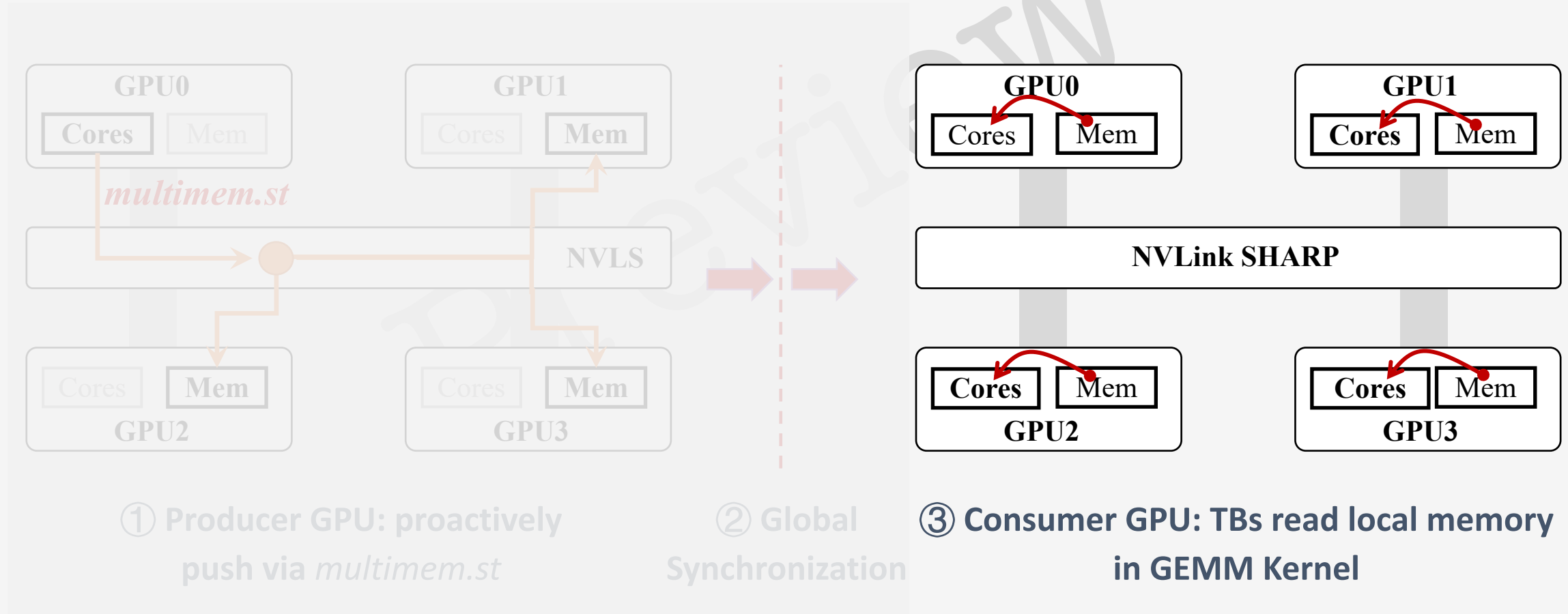
② **Global Synchronization**

③ **Consumer GPU: TBs read local memory in GEMM Kernel**

All-Gather + GEMM: remote store $\rightarrow$ *syn.* $\rightarrow$ *local read*



*All-Gather + **GEMM**: remote store $\rightarrow$ syn. $\rightarrow$ local read*

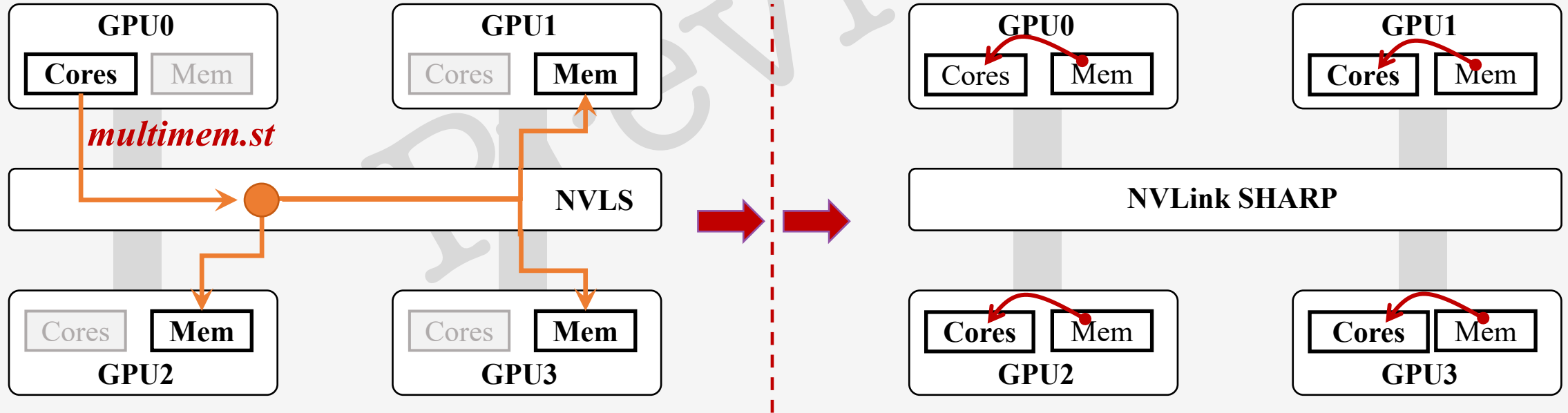


In summary: Producer **PUSH**, then Consumer **READ**

14

- The GEMM computation requires memory reads from both local and remote devices, yet *multimem.st*, the in-switch instruction used for All-Gather, operates in push mode.

*All-Gather + GEMM: remote **store** → syn. → local **read***



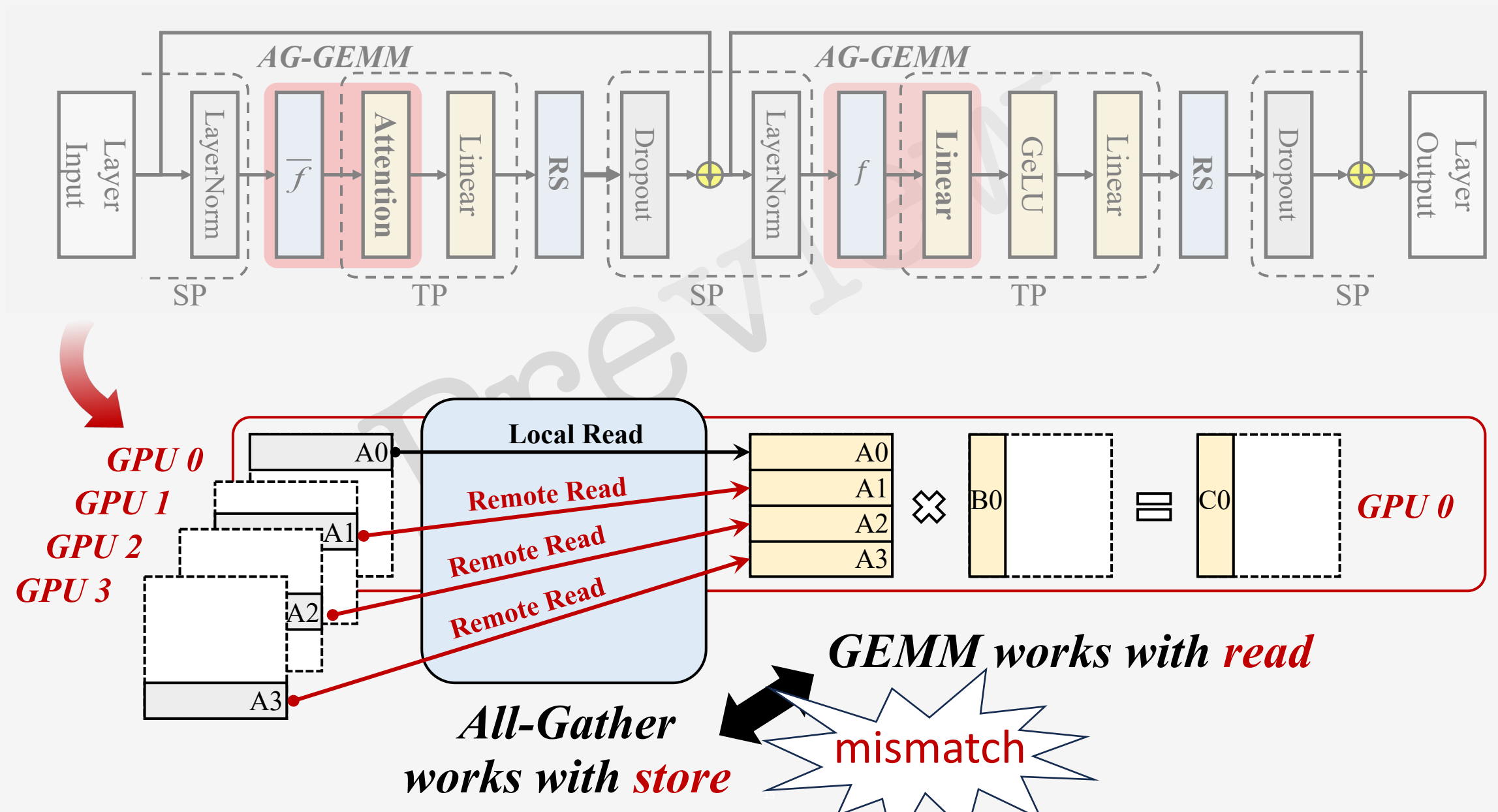
① Producer GPU: proactively push via *multimem.st*

② Global Synchronization

③ Consumer GPU: TBs read local memory in GEMM Kernel

From GEMM side: **Rethinking** All-Gather+GEMM

15

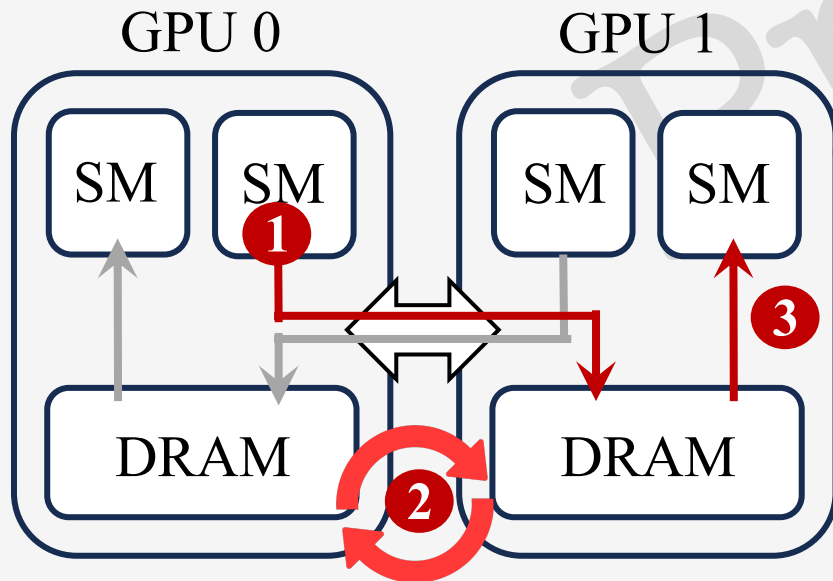


From GEMM side: **Rethinking** All-Gather+GEMM

16

- Issue 1: isolated data-transfer & compute in separate phases (GPU 0 vs. 1)
- Issue 2: frequent global synchronization
- Issue 3: more programming effort & software cost

→ in-efficiency in fine-grained compute-communicate overlap

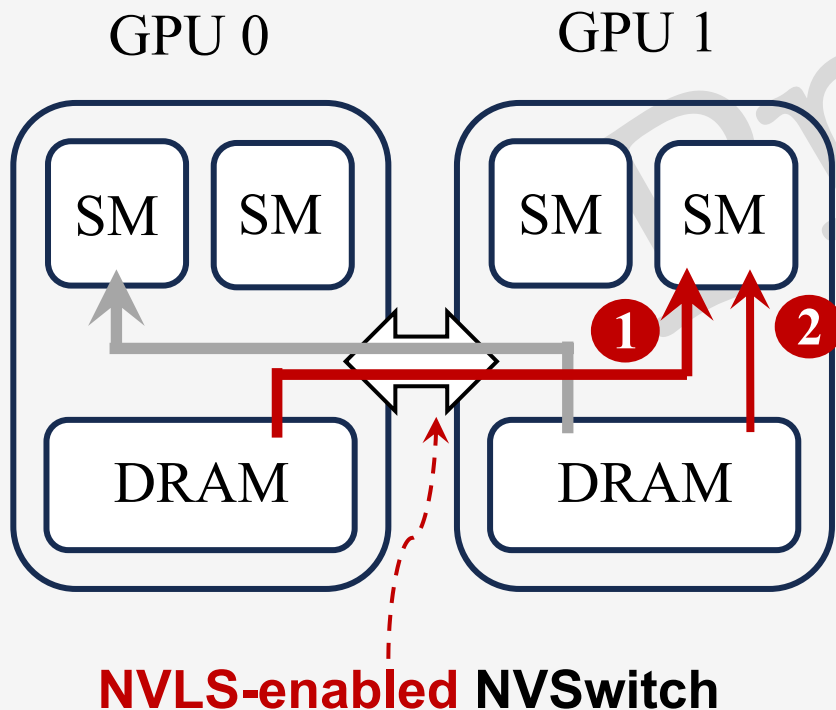


```
__global__ AG_GEMM_MultimemSt(.....) {  
    // Perform AllGather with multimem.st  
    1 AllGather_MultimemSt();  
    // Guarantee remote GPU has pushed data to local DRAM  
    2 Global_Barrier();  
    // Load data to shared memory from local DRAM  
    3 Local_Load();  
    // Guarantee data has been in shared memory  
    Local_Barrier();  
    // Perform computation with data in shared memory  
    Computation();  
}
```

What if AllGather in “Pull Mode” ?

17

- Benefit 1: ~~disjoint~~ data-move & compute **unified** in one single end
 - Benefit 2: frequent **less** global synchronization
 - Benefit 3: ~~more~~ **less** programming effort & software cost
- **smoother data-flow & compute-communicate overlap**

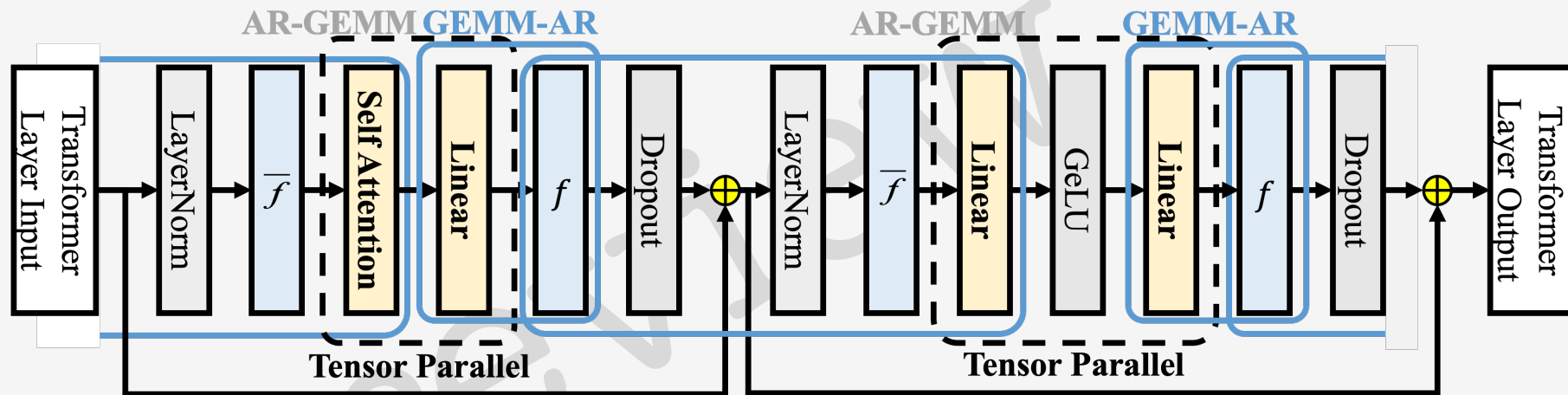


```
__global__ AG_GEMM_Load(.....) {  
    // Directly read remote data to shared memory with load  
    1 Remote_Load();  
    // Load data to shared memory from local DRAM  
    2 Local_Load();  
    // Guarantee data has been in shared memory  
    Local_Barrier();  
    // Perform computation with data in shared memory  
    Computation();  
}
```

(TP | SP) & (Forward | Backward)

19

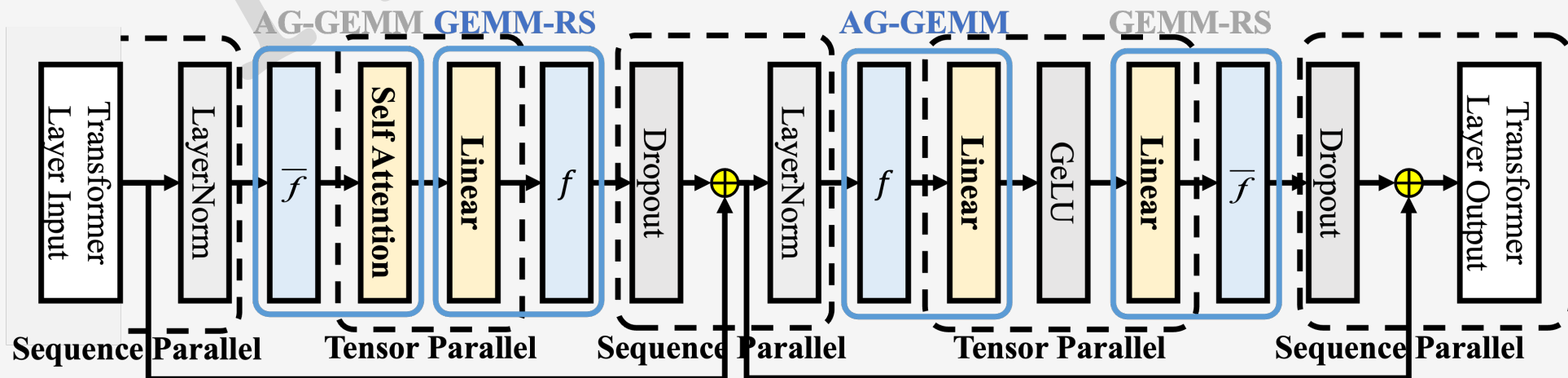
Basic Tensor Parallelism (Basic TP)



f :forward

$\bar{f}$:backward

Tensor Parallelism with Sequence Parallelism (TP with SP)



Semantic **Mismatch** between Comp. & Comm.

Existing NVLink SHARP Primitives

Existing Primitives	Collective Op	Comm. Mode
multimem.st	AllGather	Push
multimem.ld_reduce	Reduce-Scatter	Pull
multimem.red	AllReduce	Push

Tensor Parallelism Requirement

TP Operation		Required Memory Semantics
TP with SP	AG-GEMM	Read Memory
	GEMM-RS	Write Memory
Basic TP	AR-GEMM	Read Memory
	GEMM-AR	Write Memory

Mis-Matches

✗

“Read-Push”

✗

“Write-Pull”

✗

“Read-Push”

CAIS: Re-design NVLS in a **Compute-aware** Fashion

Existing NVLink SHARP Primitives

Existing Primitives	Collective Op	Comm. Mode
multimem.st	AllGather	Push
multimem.ld_reduce	Reduce-Scatter	Pull
multimem.red	AllReduce	Push

Our Proposal

Primitives	Comm. Mode
ld.cais	Pull
red.cais	Push

Tensor Parallelism Requirement

TP Operation		Required Memory Semantics
TP with SP	AG-GEMM	Read Memory
	GEMM-RS	Write Memory
Basic TP	AR-GEMM	Read Memory
	GEMM-AR	Write Memory

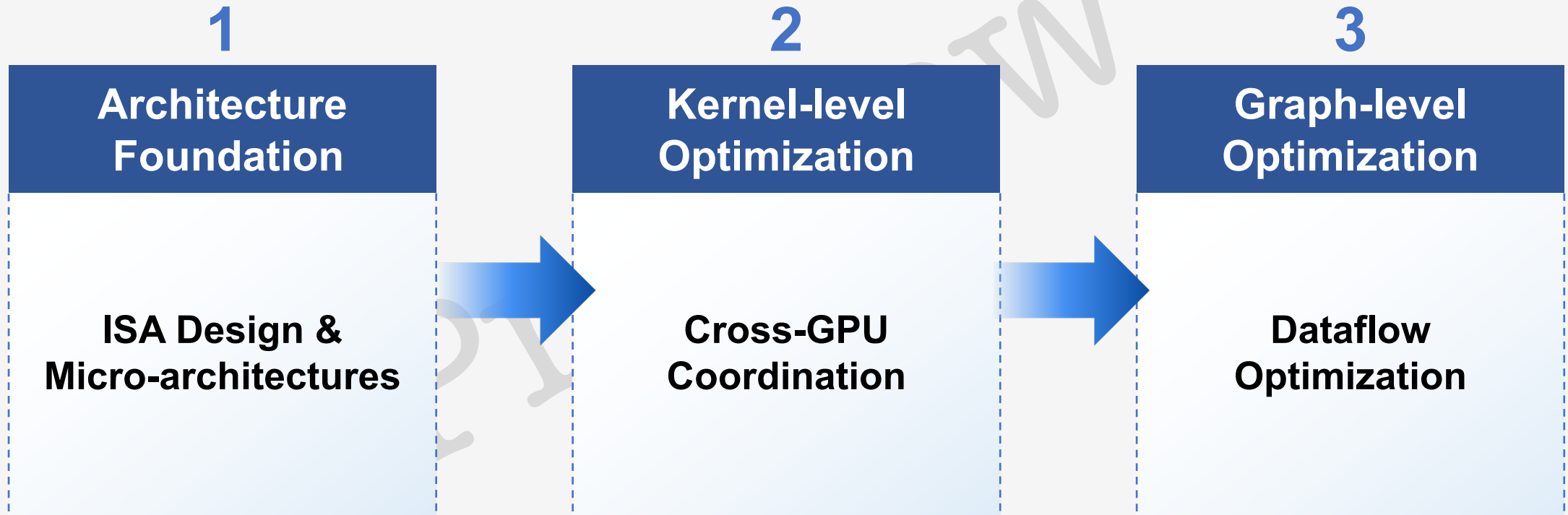
✓ **Match!**

1. Background
2. Motivation
- 3. Solutions**
4. Experiments
5. Conclusion



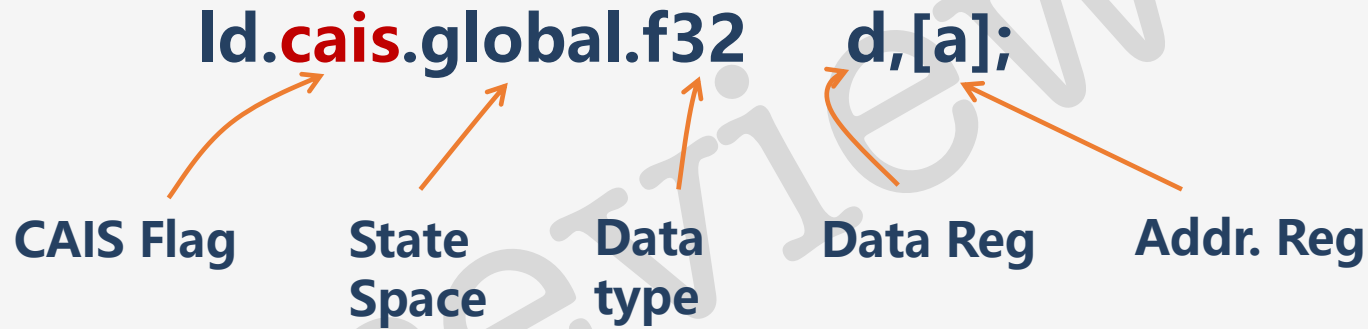
Three Key Pieces in CAIS

24

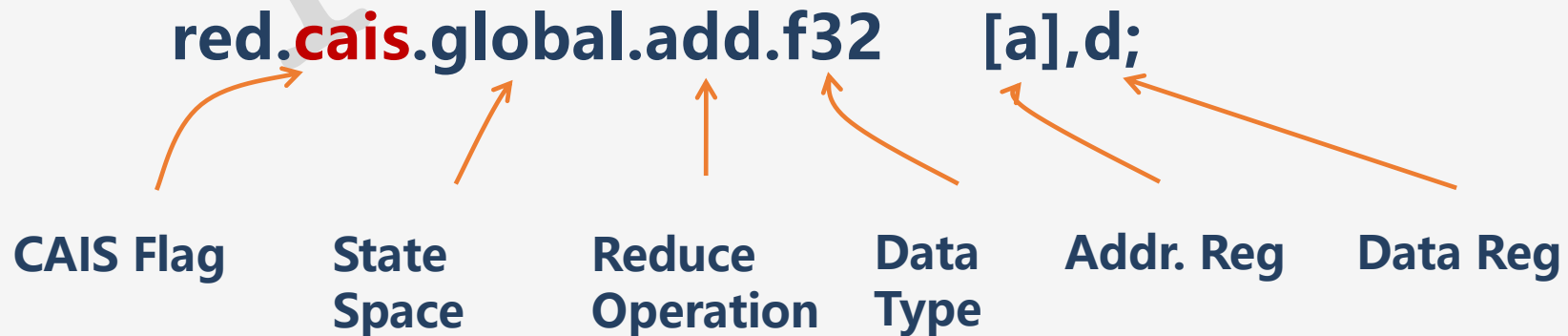


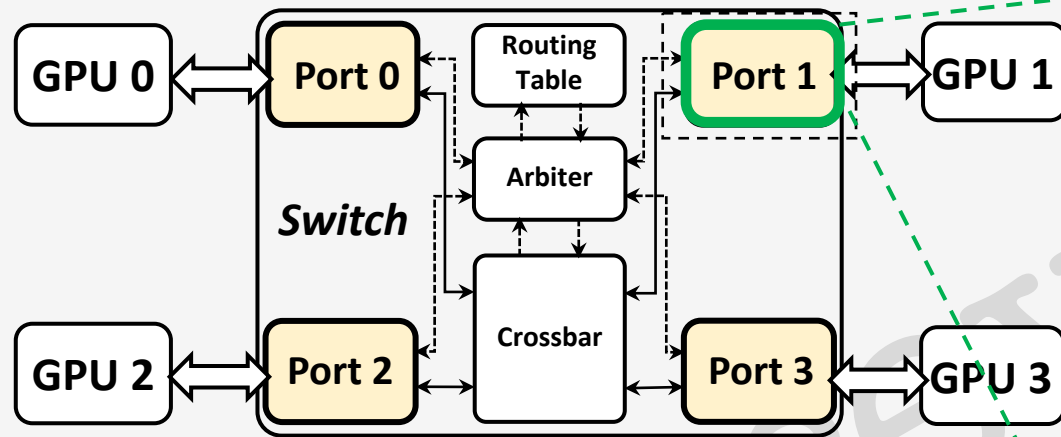
	Original <i>Inst.</i> in Existing Arch	The Proposed <i>Inst.</i> in CAIS
Instruction 1: <i>load</i>	ld.global.f32	<u>ld.cais.global.f32</u>
Instruction 2: reduce	red.global.add.f32	red.cais.global.add.f32

■ Instruction 1: *load*

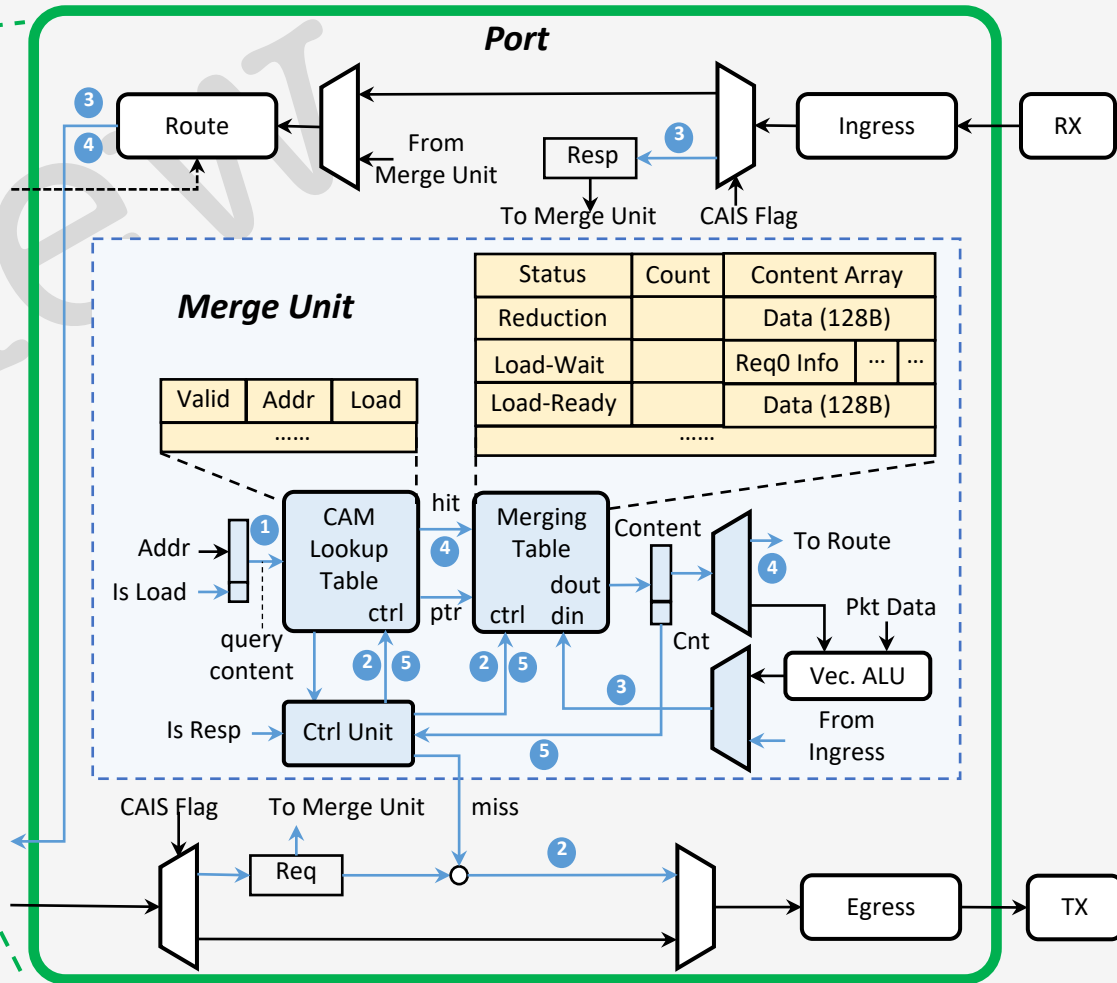


■ Instruction 2: *reduce*



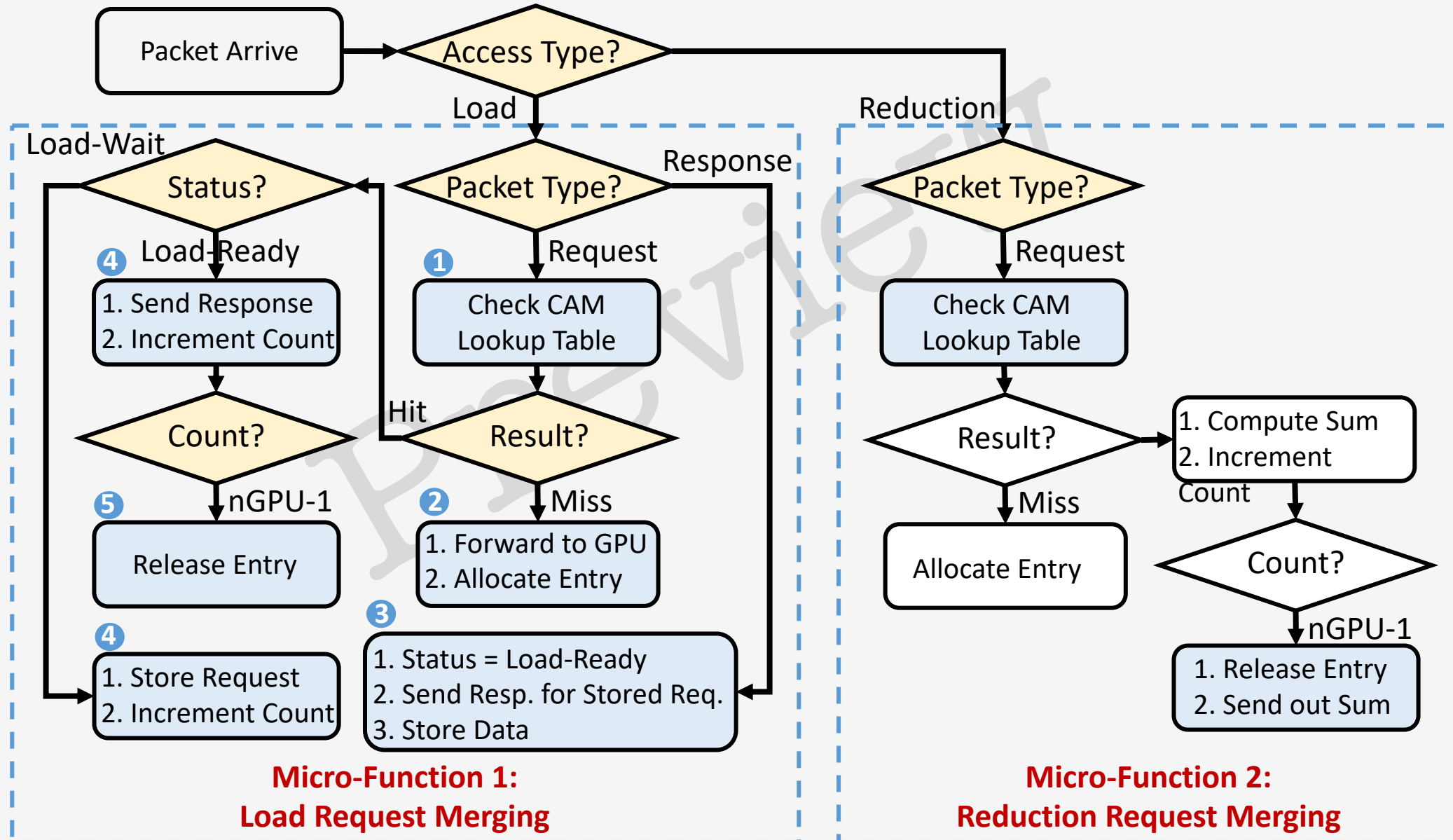


- **CAM Lookup Table**
 - Matching CAIS requests
- **Merge Table**
 - Maintain Partial Results



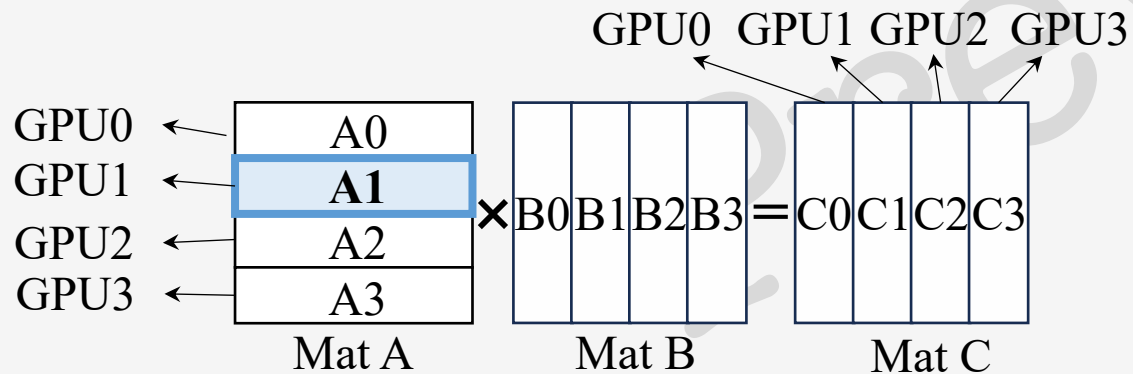
Tech 1: In-Switch Micro-Functions

28

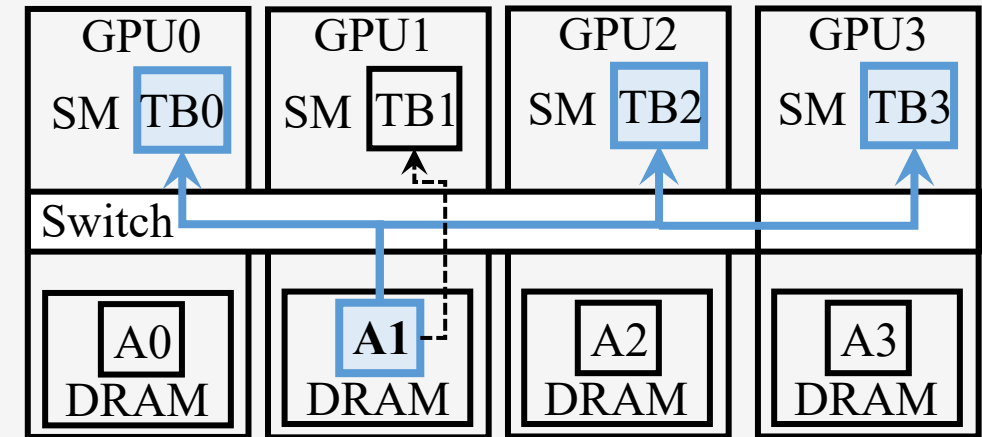


■ AllGather-GEMM Example:

“A1” data is requested by multiple GPUs (0/2/3)

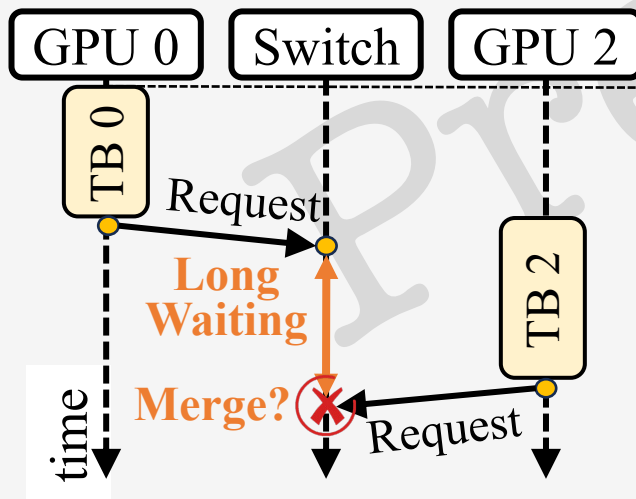


**AllGather-GEMM
Example**

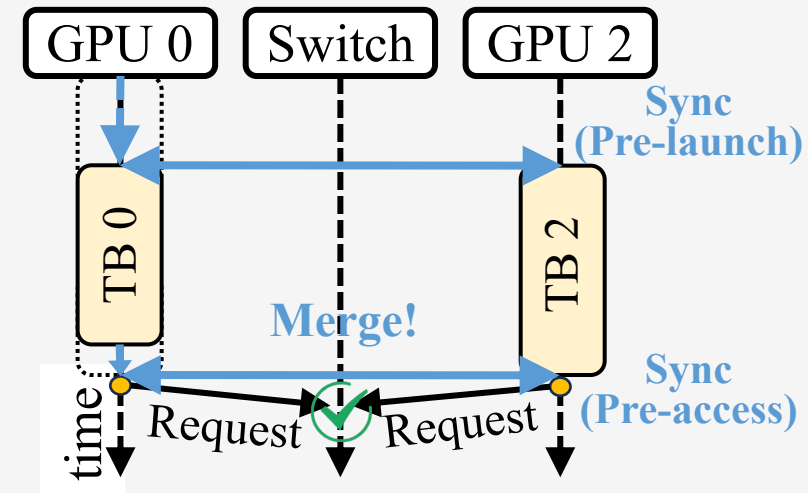


**In-switch Request
Broadcasting**

- Without coordination, mergeable load or reduction requests from different GPUs may **arrive** at the switch **at different times**, resulting in missed merging opportunities or buffer pressure due to delayed aggregation.



Timeline without
Coordination



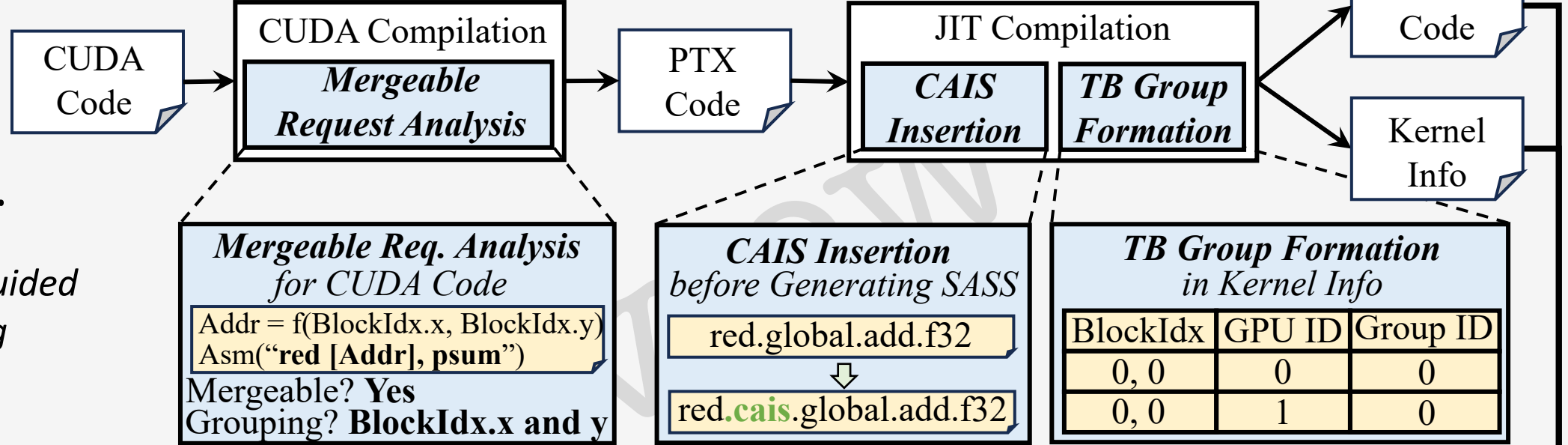
Timeline with
Coordination

Tech 2: Cross-GPU TB Coordination

31

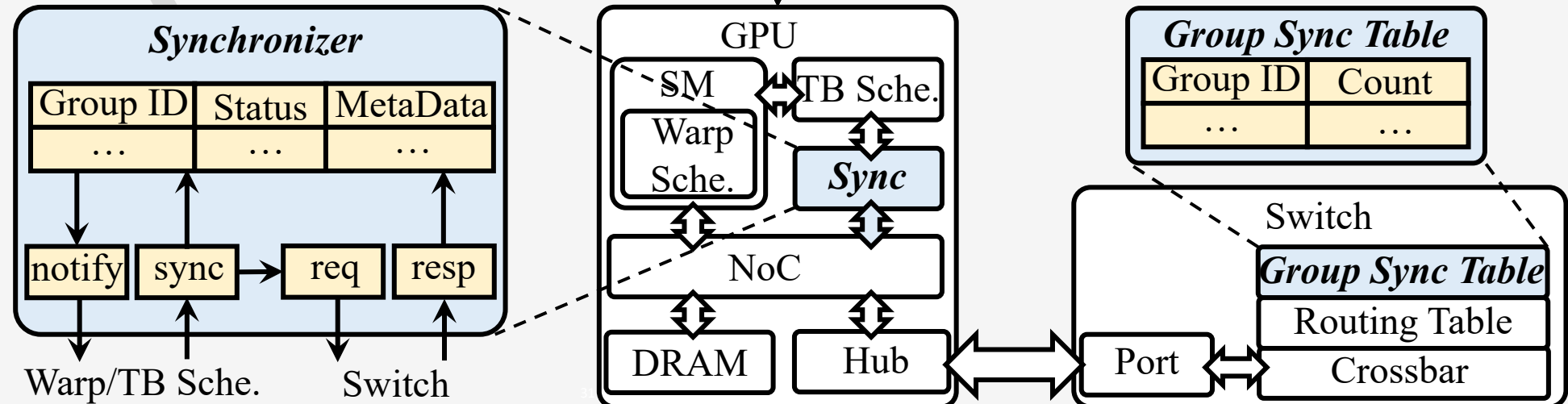
Compiler

Compiler-Guided
TB Grouping

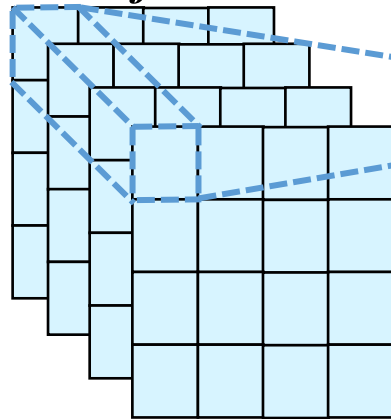


Architecture

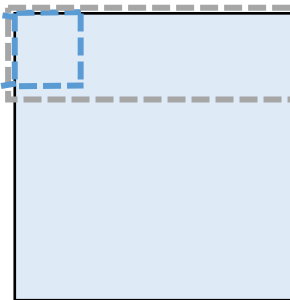
Group-Aware
Synchronization
Mechanisms



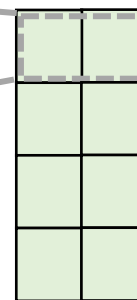
*Output Projection Layer
of Attention*



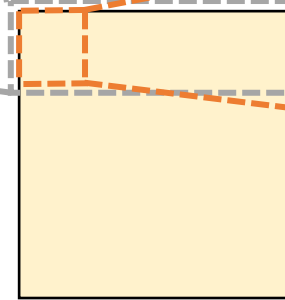
GEMM-1



Mat A

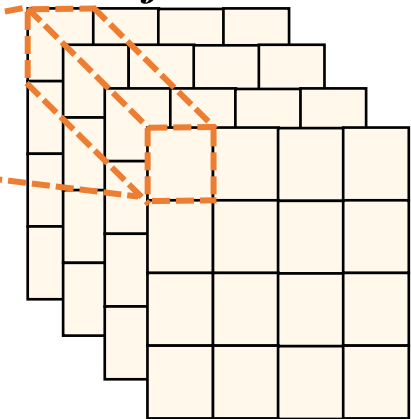


LN



Mat B

*First Layer
of FFN*

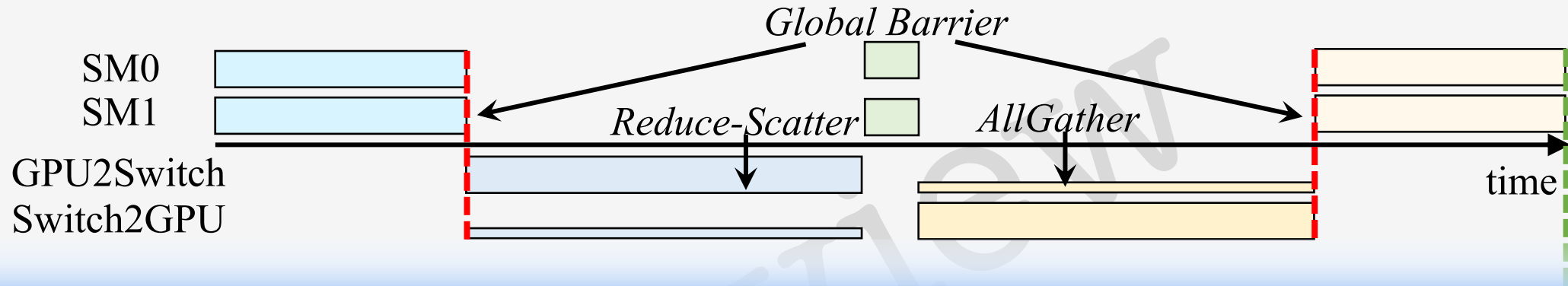


GEMM-2

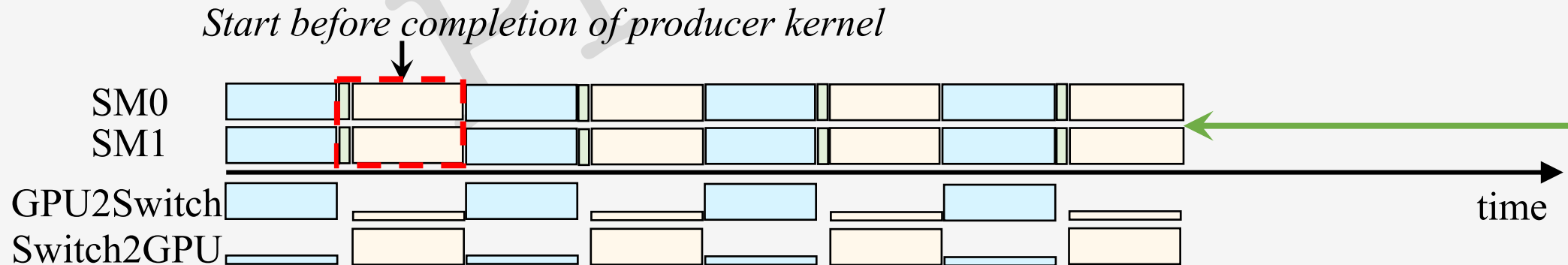
Sub-layer of Transformer layer and its fine-grained dependency

Tech 3: CAIS-enabled Kernel Fusion

33



Timeline of communication-centric in-switch computing

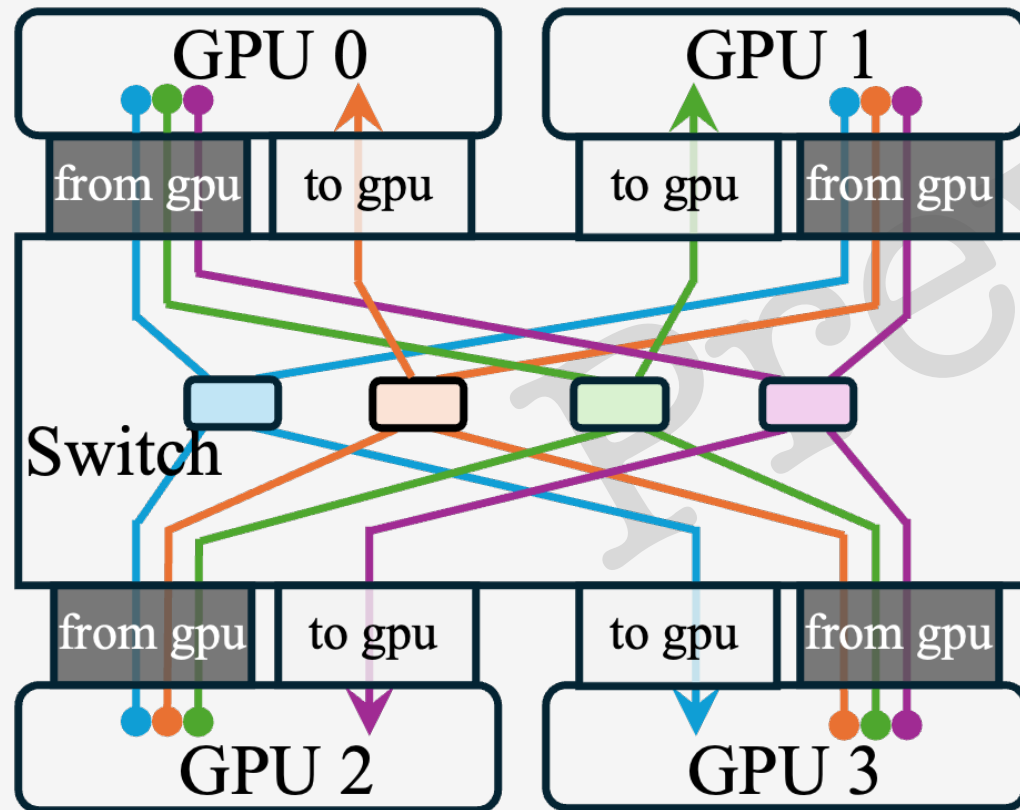


Timeline of compute-aware in-switch computing + fine-grained dependency

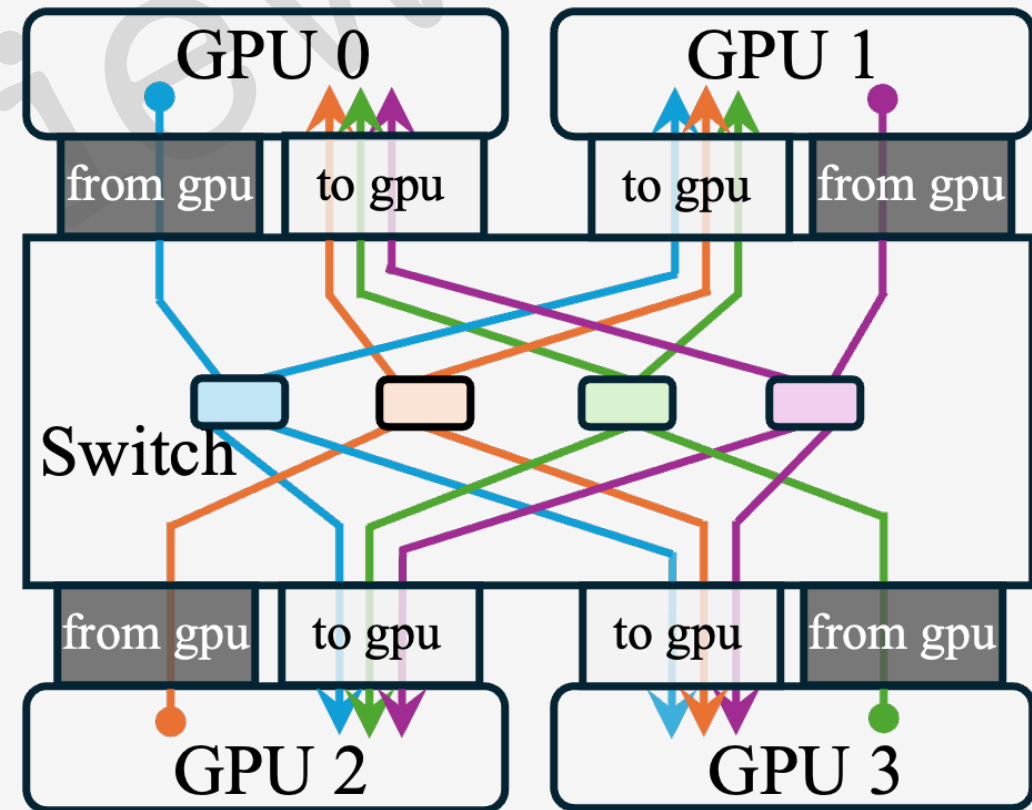
Tech 3: Opportunities of **Complementary** Traffic

34

Reduction in GEMM-RS

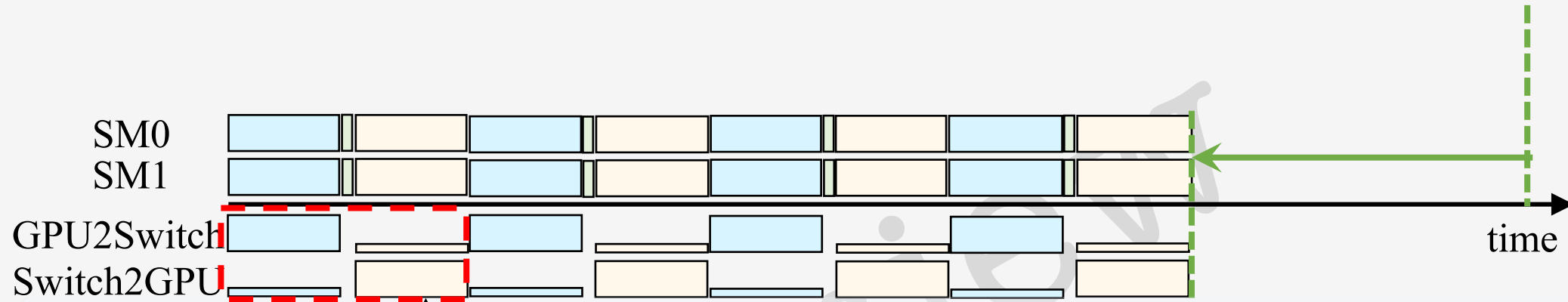


Load in AG-GEMM



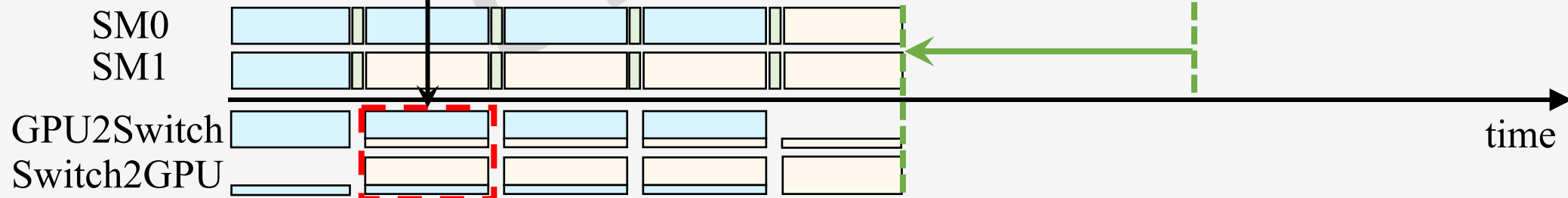
Tech 3: Kernel Fusion for Complementary Traffic

35



Timeline of compute-aware in-switch computing + fine-grained dependency

Utilize complementary Communication pattern



Timeline of compute-aware in-switch computing + fine-grained dependency + asymmetric kernel overlapping

1. Background
2. Motivation
3. Solutions
- 4. Experiments**
5. Conclusion



■ Machine Setup:

- 8 GPUs + 4 NVSwitches, similar to NVIDIA DGX-H100
- Modified Accel-Sim + BookSim2

■ Workload:

- Mega-GPT-4B, Mega-GPT-8B, LLaMA-7B
- Training and Inference

■ Previous work:

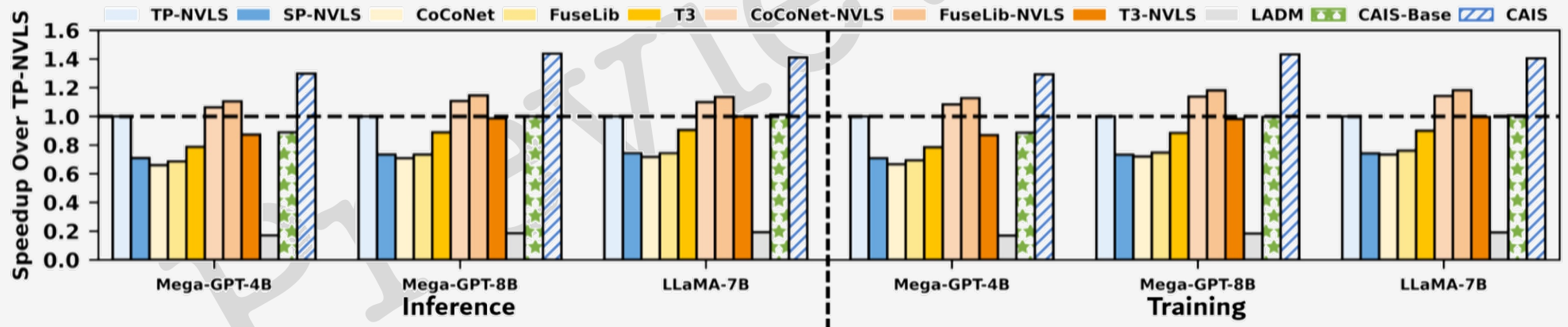
	Baseline	MICRO'20	ASPLOS'22	SC'24	ASPLOS'24
Original	-	LADM	CoCoNet	FuseLib	T3
+NVLS	TP/SP+NVLS	-	CoCoNet+NVLS	FuseLib+NVLS	T3+NVLS

Result 1: Model Speedup (Training/Inference)

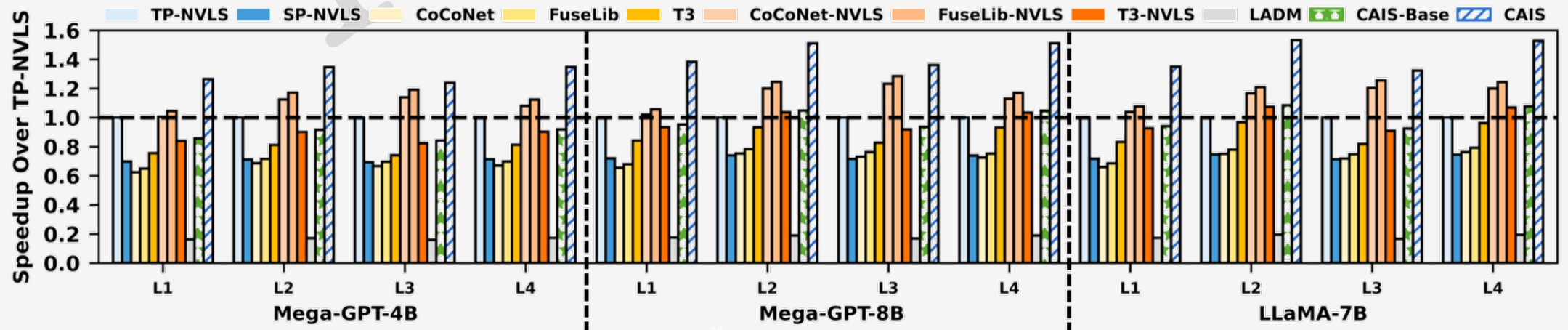
38

- Inference: $1.21\times \sim 1.99\times$
- Training: $1.25\times \sim 2.03\times$

End-to-End



Layer wise



Result 2: Benefits of TB Coordination

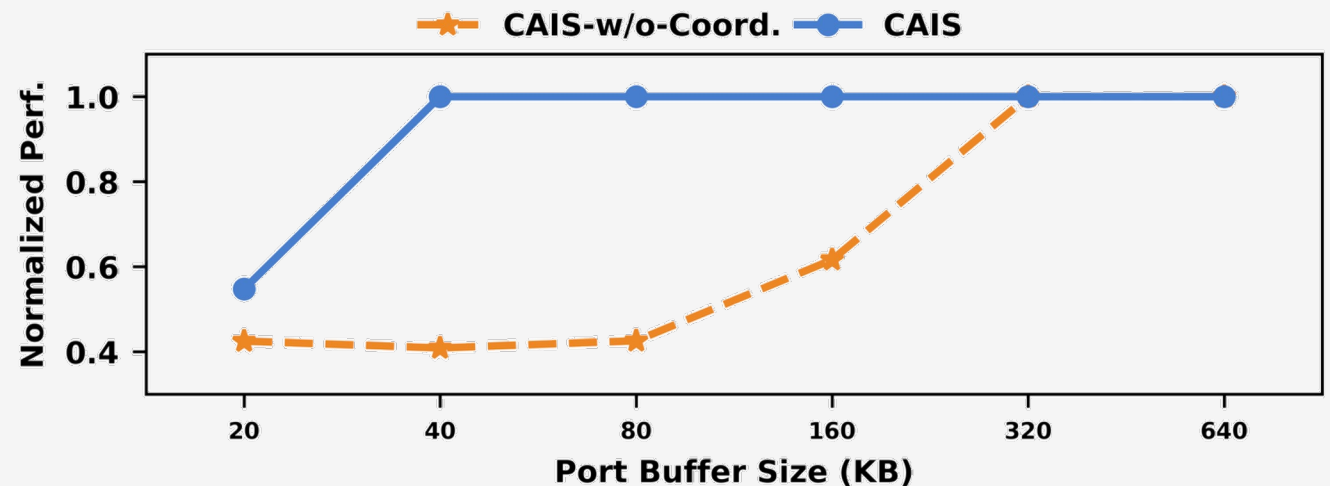
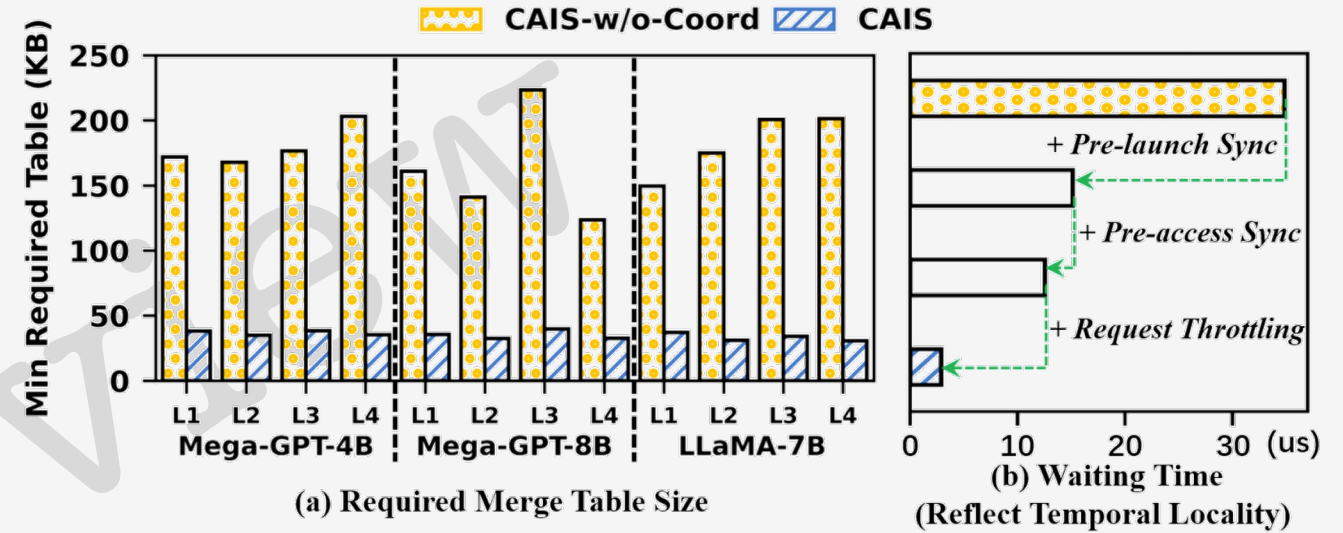
39

■ Higher Performance:

- The waiting time was reduced from 35 to 3 microseconds for merge requests.

■ Lower Hardware Cost:

- This further decreases the required capacity of the lookup table, leading to higher performance with fewer in-switch caches.



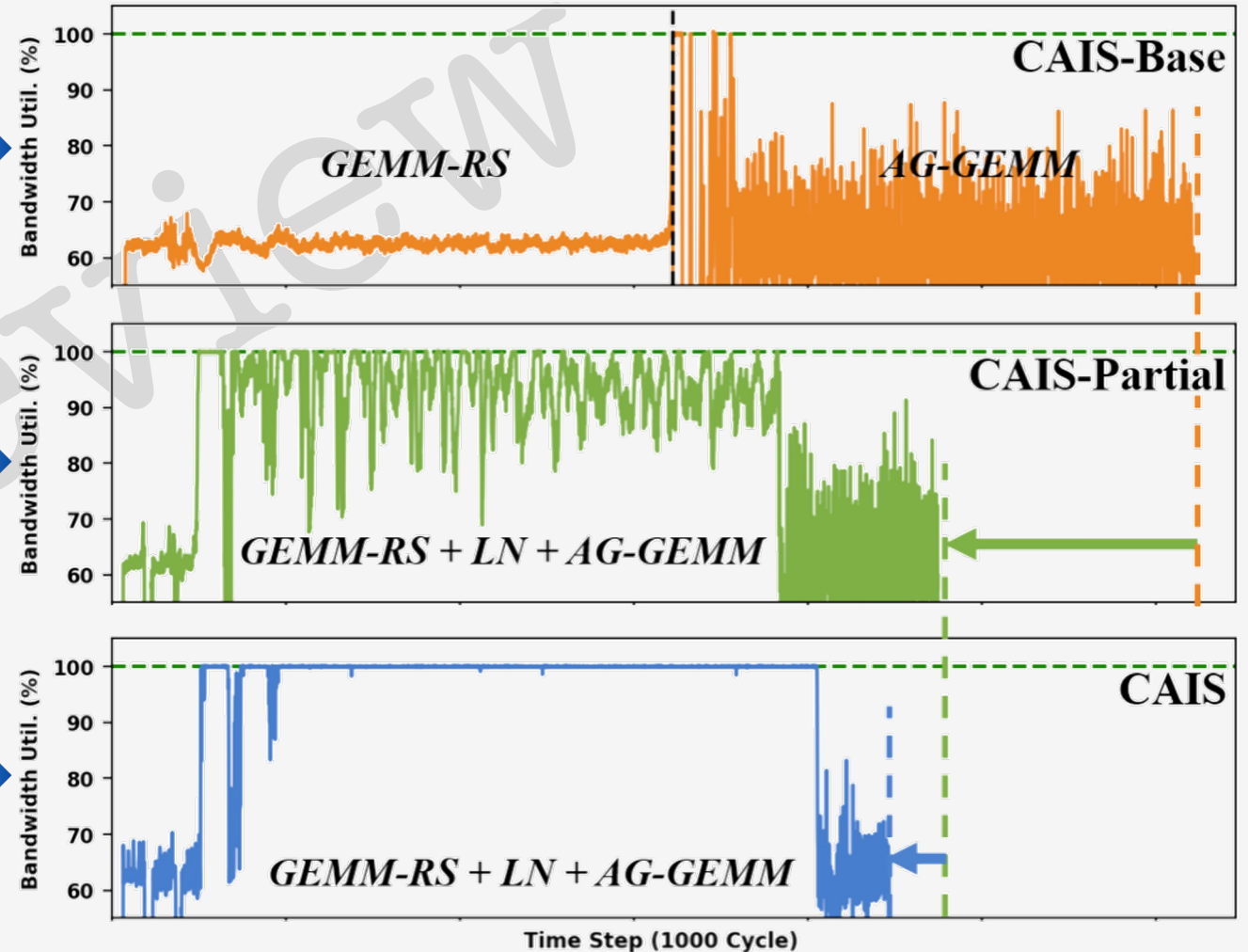
Result 3: Bandwidth Utilization Improvements

40

CAIS base: 62.4%
Only ISA & Arch

CAIS Partial: 84.7%
ISA & Arch + Graph

CAIS Full: 90.2%
ISA & Arch + Graph + Trfc Ctl



Result 4: Scalability and HW Cost

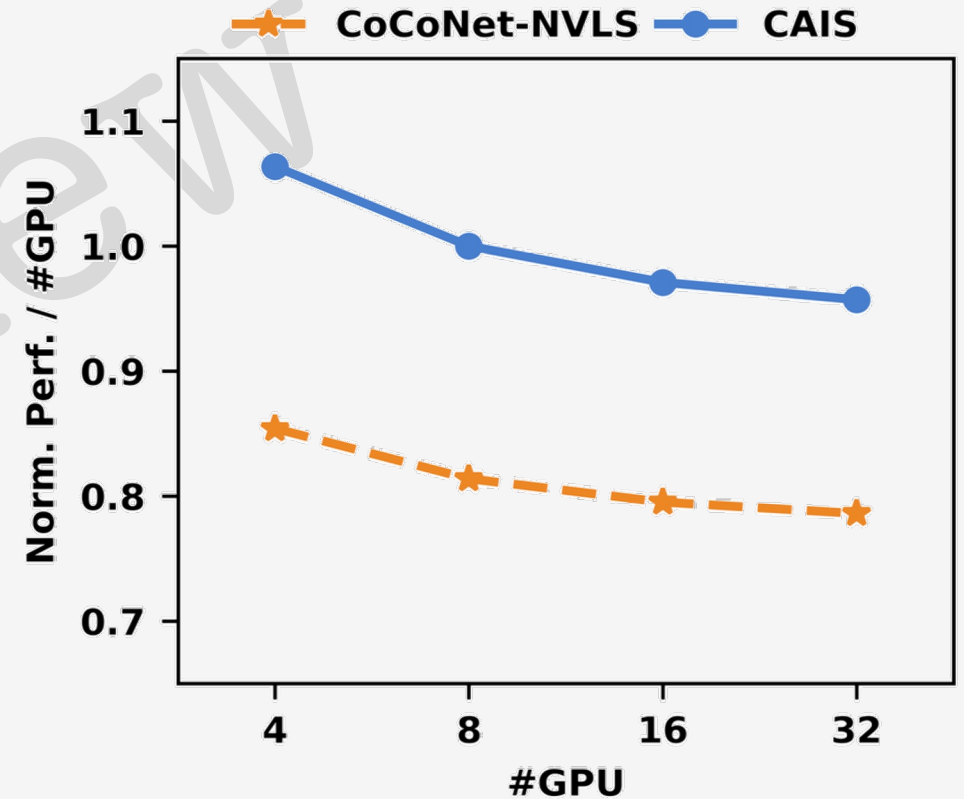
41

■ Scalability:

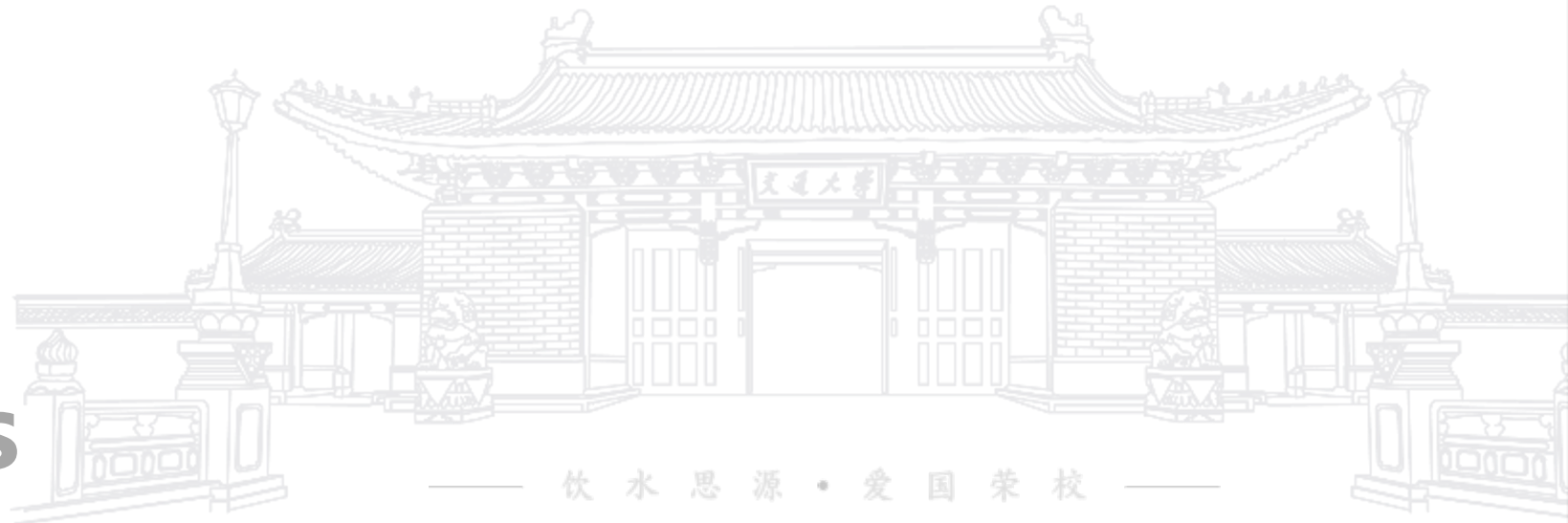
- The performance with 32 GPUs is less than 5% lower than with 8 GPUs.
- The hardware overhead does not increase with the number of GPUs

■ HW Cost:

- Switch: 0.5 mm², < 1% of NVswitch.
- GPU: 0.019 mm², <0.01% of H100.
- 12 nm process



1. Background
2. Motivation
3. Solutions
4. Experiments
- 5. Conclusion**



- The existing NVSwitch is designed with a **communication-centric** approach, **lacking consideration for up/down-stream operators** (e.g., GEMM), limiting the performance of in-switch computing in Scale-up GPU System.
- This work provides a detailed analysis of the TP dataflow during both LLM training and inference, and subsequently redesigns a **Compute-aware In-switch Computing** communication mechanism.
- Building upon this mechanism, we propose corresponding designs at the **instruction set, architecture, and software** stacks. These designs substantially enhance the collective performance of multi-GPU systems, thereby demonstrating the advantages of the proposed scheme.



上海交通大学
SHANGHAI JIAO TONG UNIVERSITY

Thank You !

Chen Zhang

Assistant Professor

Shanghai Jiao Tong University

chenzhang.sjtu@sjtu.edu.cn

