

PystachIO: Efficient Distributed GPU Query Processing with PyTorch over Fast Networks & Fast Storage

VLDB 2026

Jigao Luo^{*}, Nils Boeschen^{*}, Muhammad El-Hindi, Carsten Binnig

^{*}: Equal contribution
Systems Group, Technische Universität Darmstadt

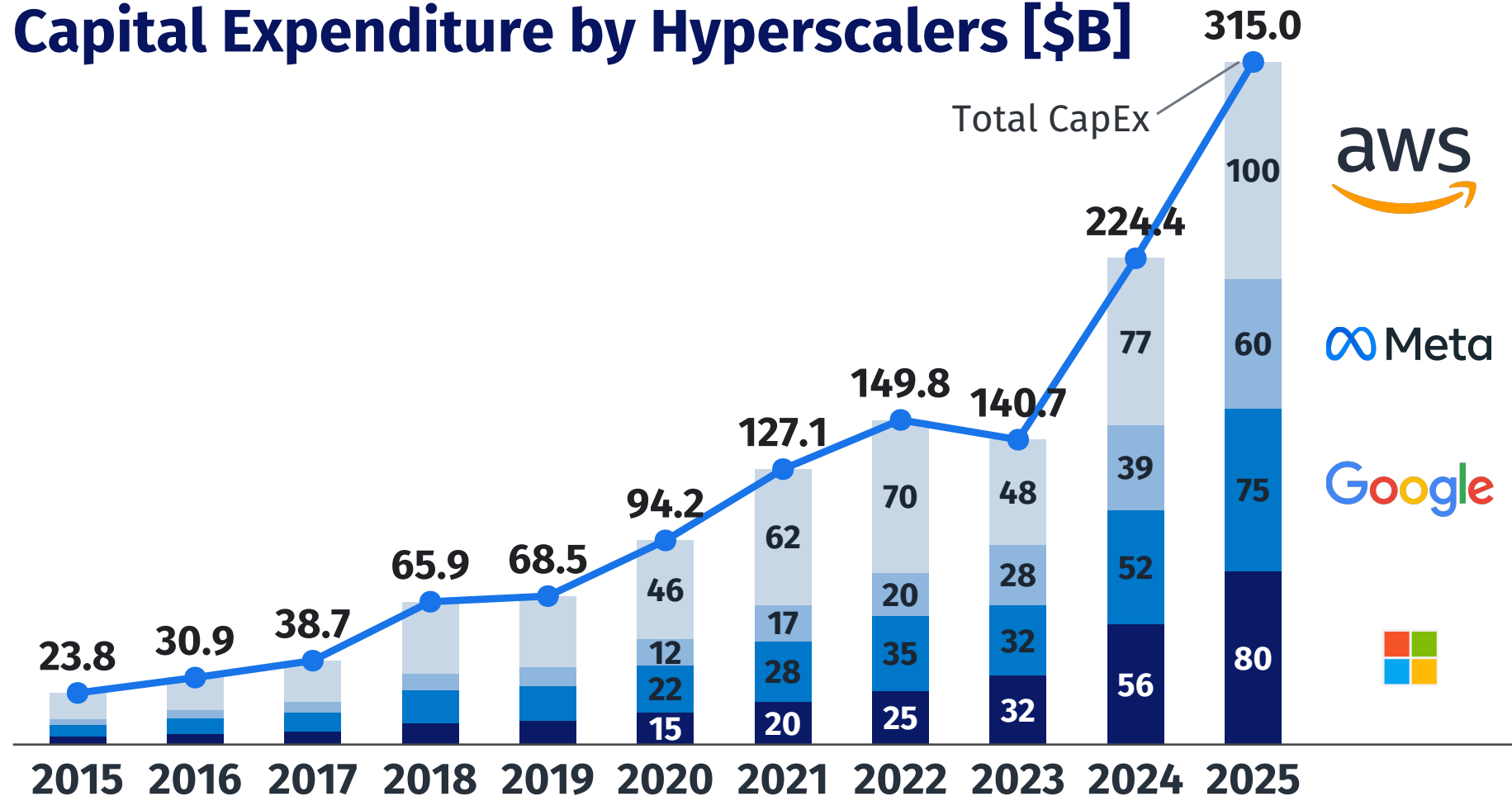


TECHNISCHE
UNIVERSITÄT
DARMSTADT

SYSTEMS

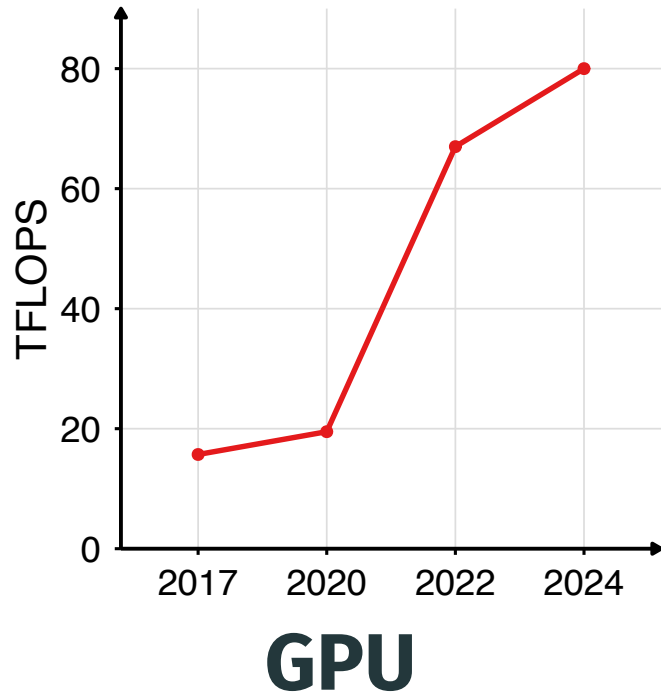
AI Investment Is Shaping the Future

Capital Expenditure by Hyperscalers [\$B]

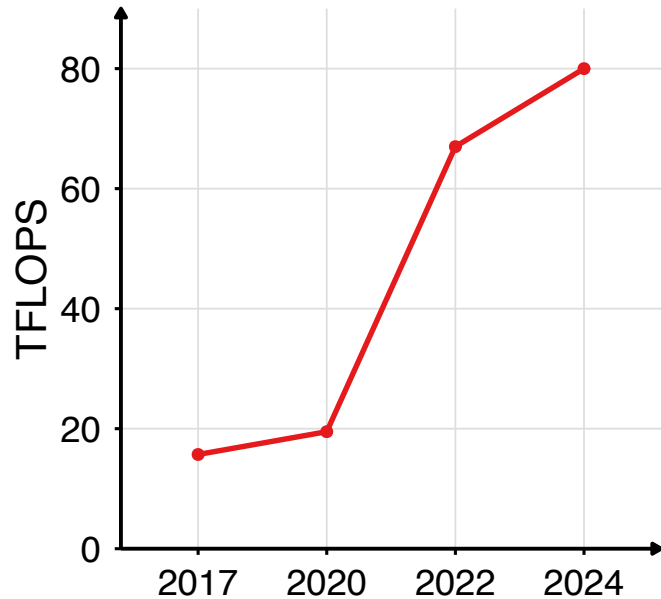


Source: DC Pulse

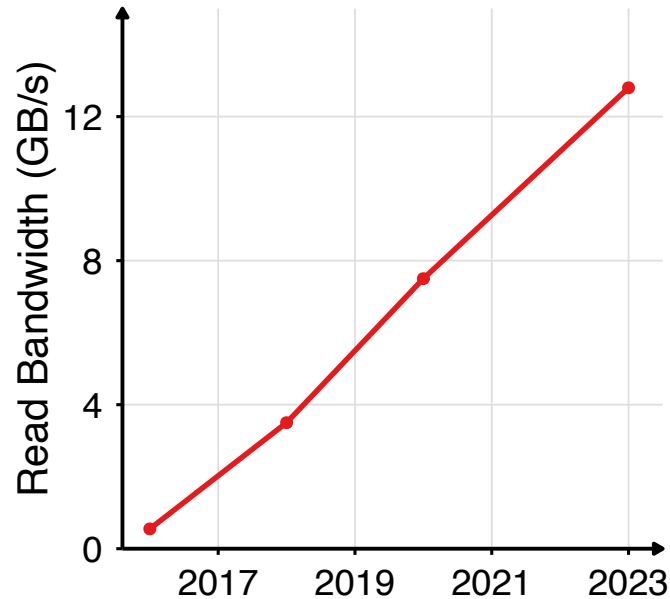
AI Investment Is Shaping the Future HW & DC



AI Investment Is Shaping the Future HW & DC

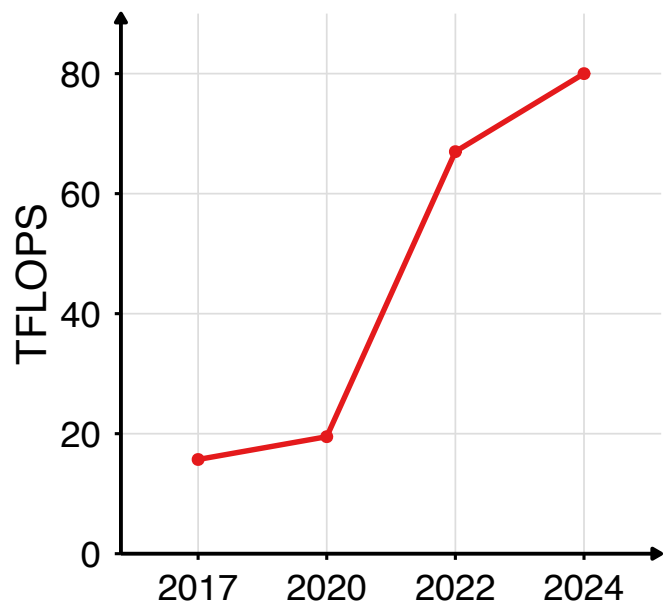


GPU

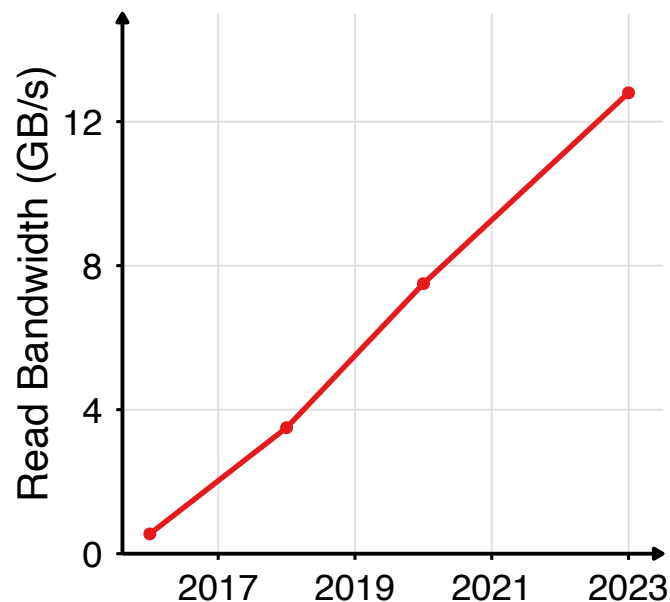


**Storage:
NVMe SSD**

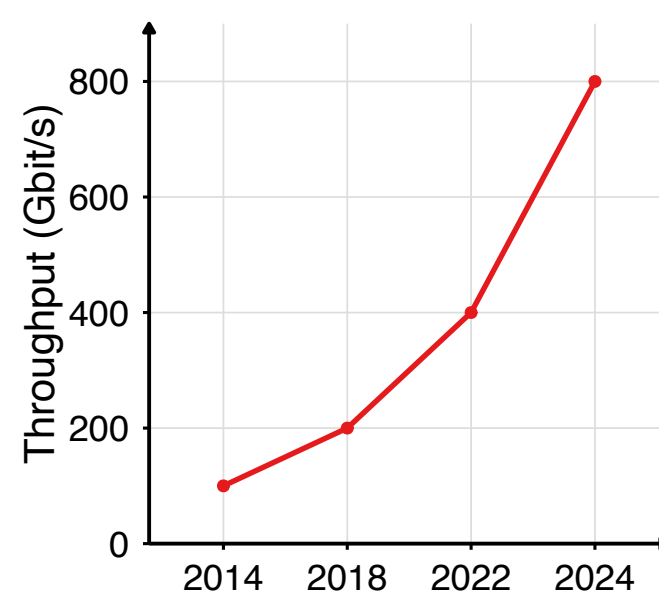
AI Investment Is Shaping the Future HW & DC



GPU

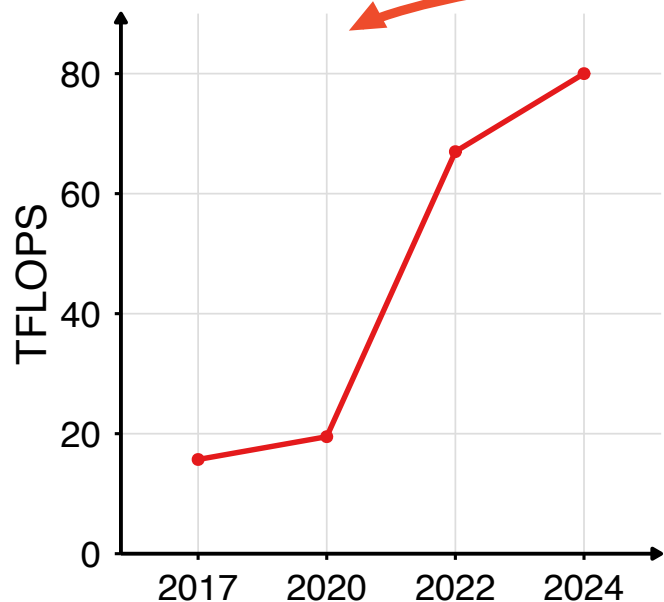


**Storage:
NVMe SSD**

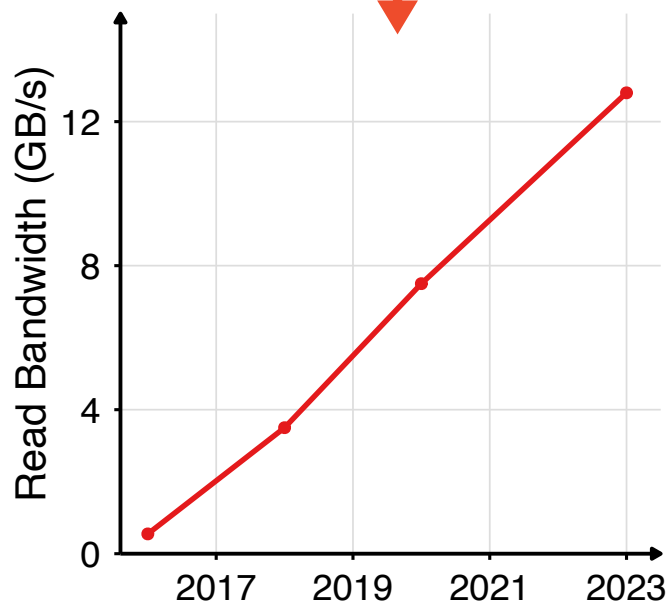


**Network:
RDMA NIC**

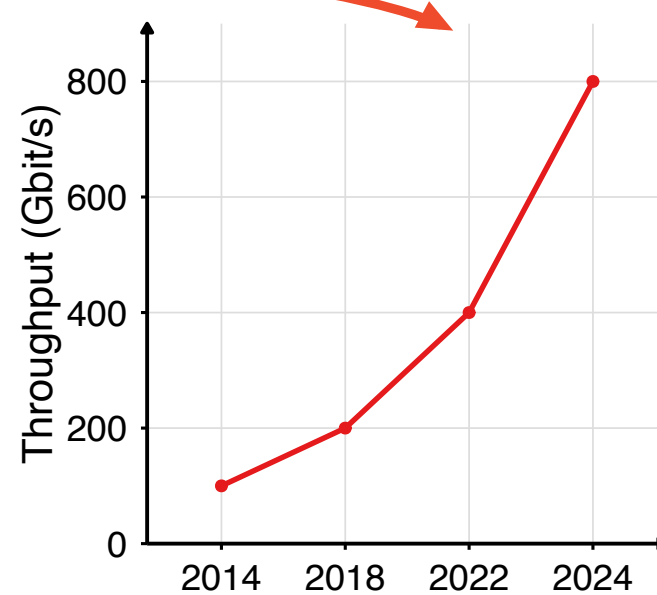
AI Investment Is Shaping the Future HW & DC



GPU



**Storage:
NVMe SSD**



**Network:
RDMA NIC**

Build a DB on PyTorch?

Query Processing on Tensor Computation Runtimes



Dong He¹, Supun Nakandala², Dalitso Banda³, Rathijit Sen³, Karla Saur³, Kwanghyun Park³,
Carlo Curino³, Jesús Camacho-Rodríguez³, Konstantinos Karanasos⁴, Matteo Interlandi³

¹University of Washington, ²University of California, San Diego, ³Microsoft, ⁴Meta

¹donghe@cs.washington.edu, ²snakanda@eng.ucsd.edu, ³firstname.lastname@microsoft.com, ⁴kkaranasos@fb.com

What is missing?

PyTorch

GPU



Storage



Network



What is missing?

	PyTorch	TQP  
GPU	✓	✓
Storage	✓	✗
Network	✓	✗

What is missing?

	PyTorch	TQP  	Cloud OLAP  
GPU	✓	✓	✗
Storage	✓	✗	✓
Network	✓	✗	✓

What is missing?

PyTorch

TQP  

Cloud OLAP  

GP

How to build a distributed

Stor

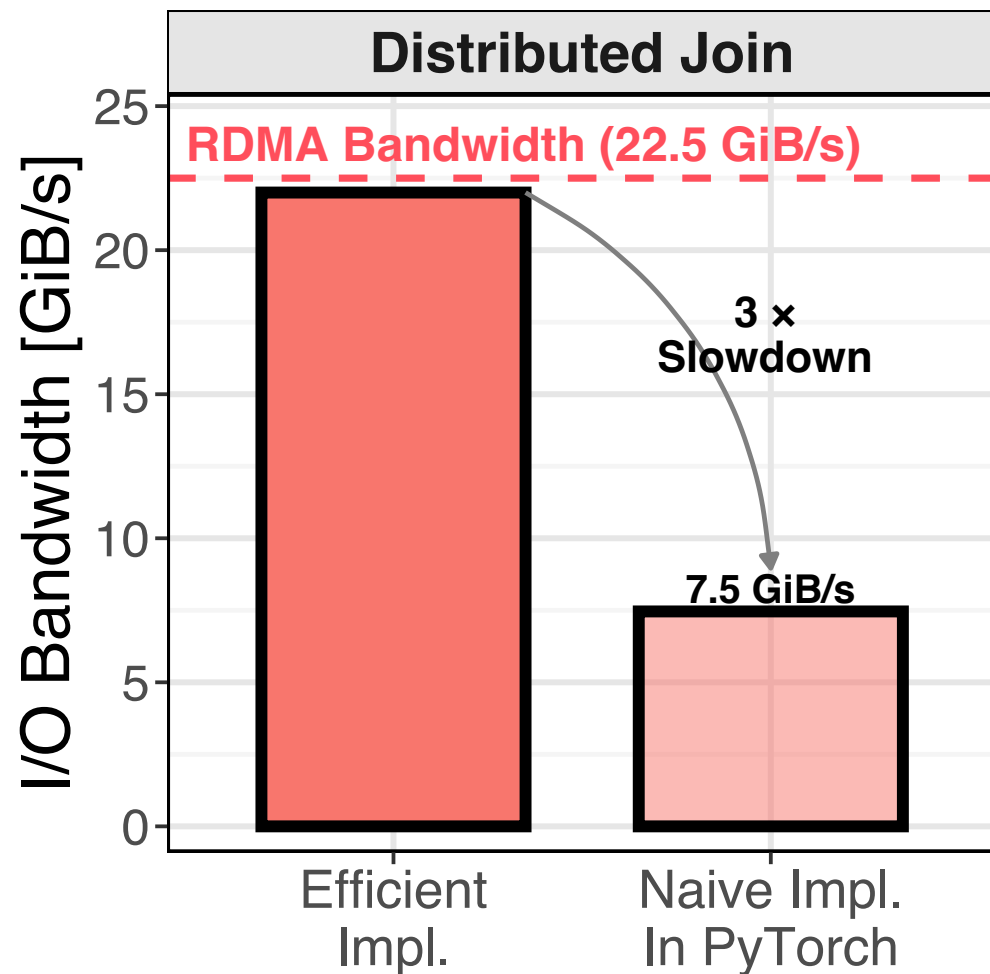
GPU DB on PyTorch?

Network

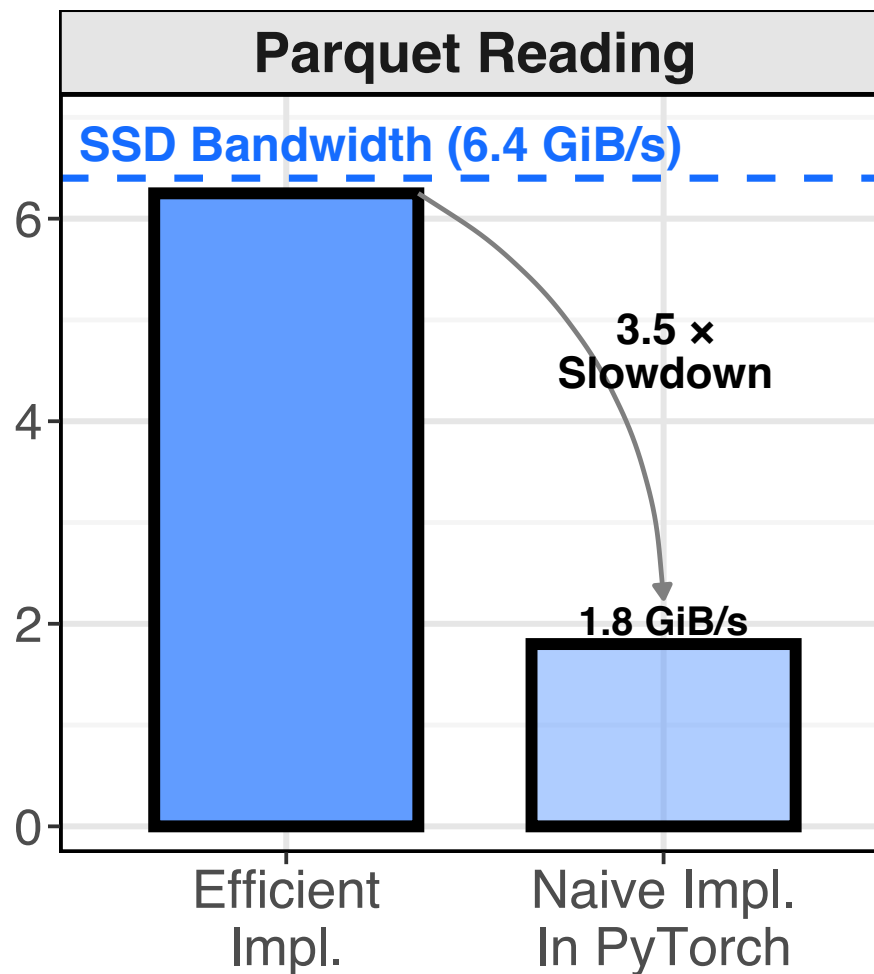
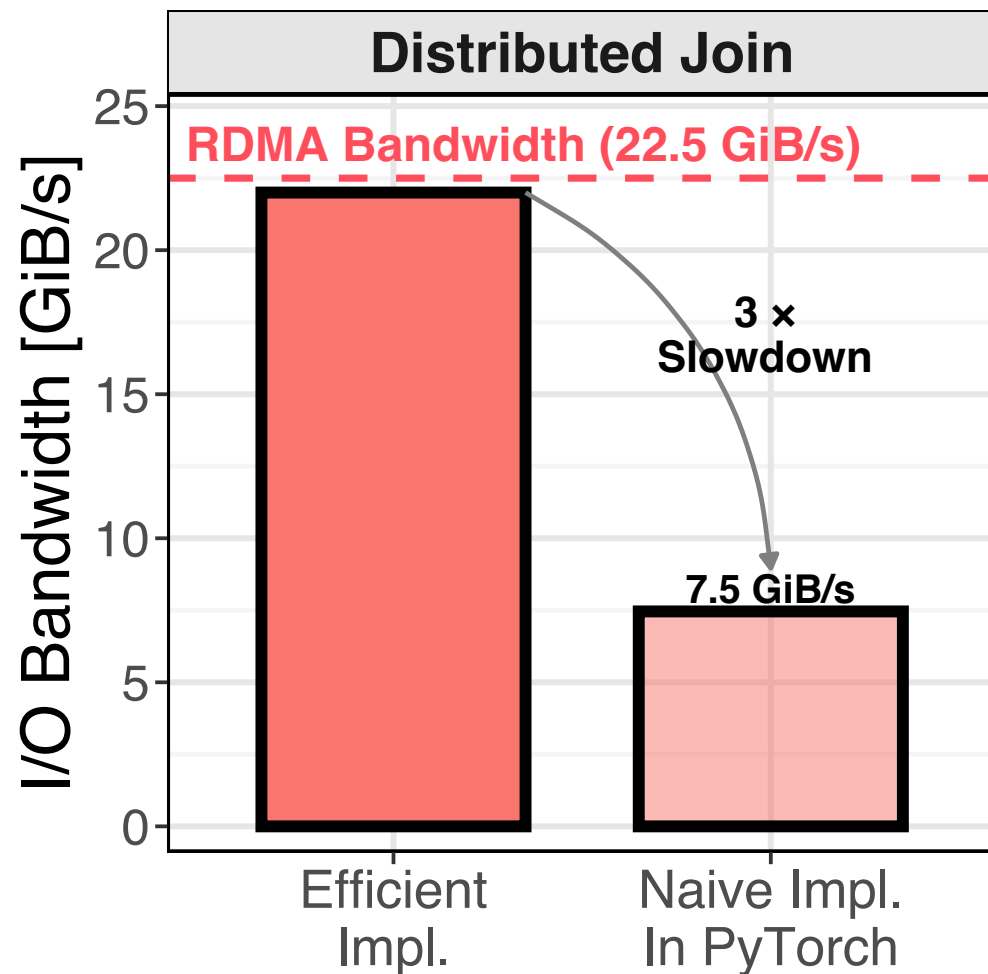


Challenge: Network and Storage I/O

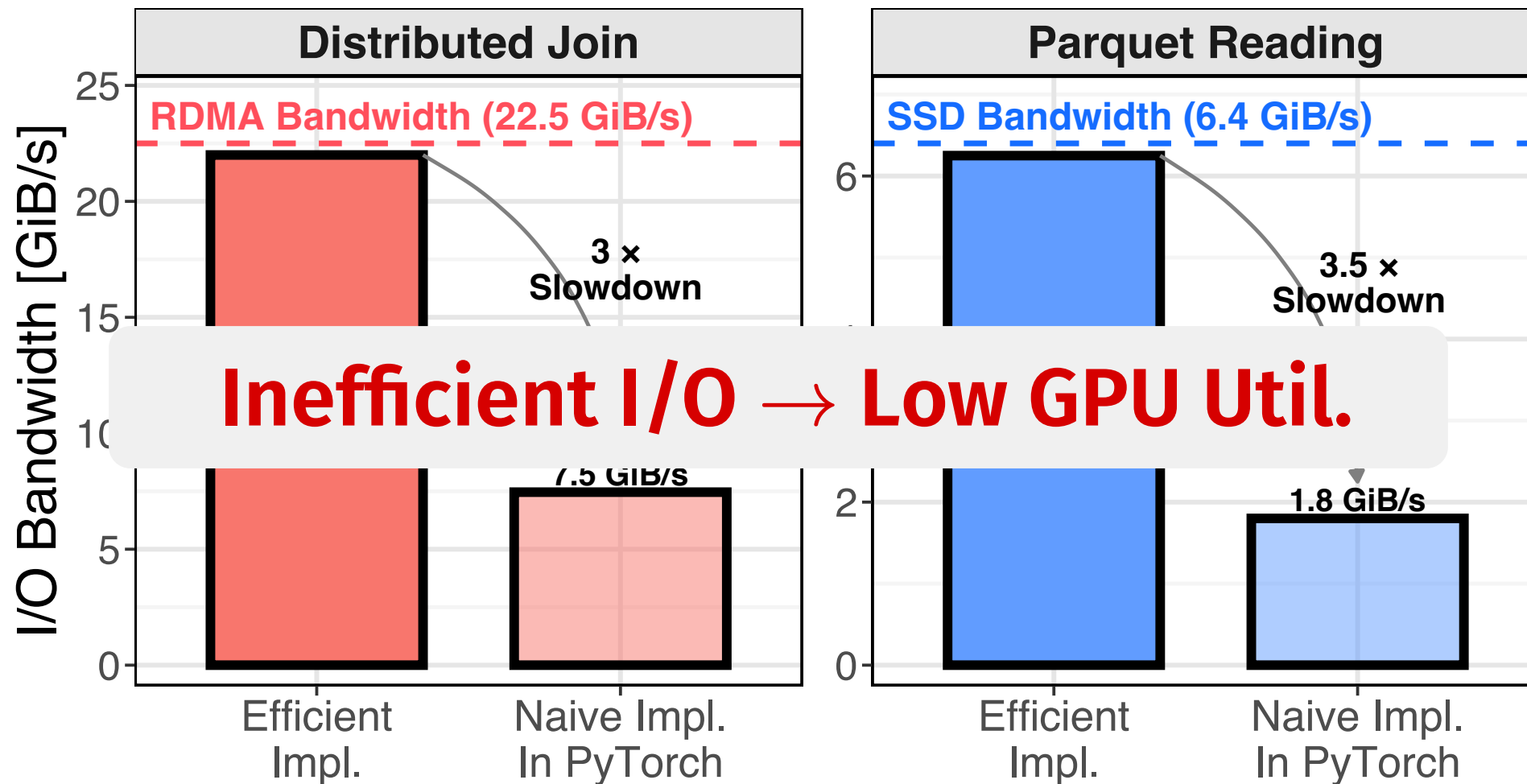
Challenge: Network and Storage I/O



Challenge: Network and Storage I/O

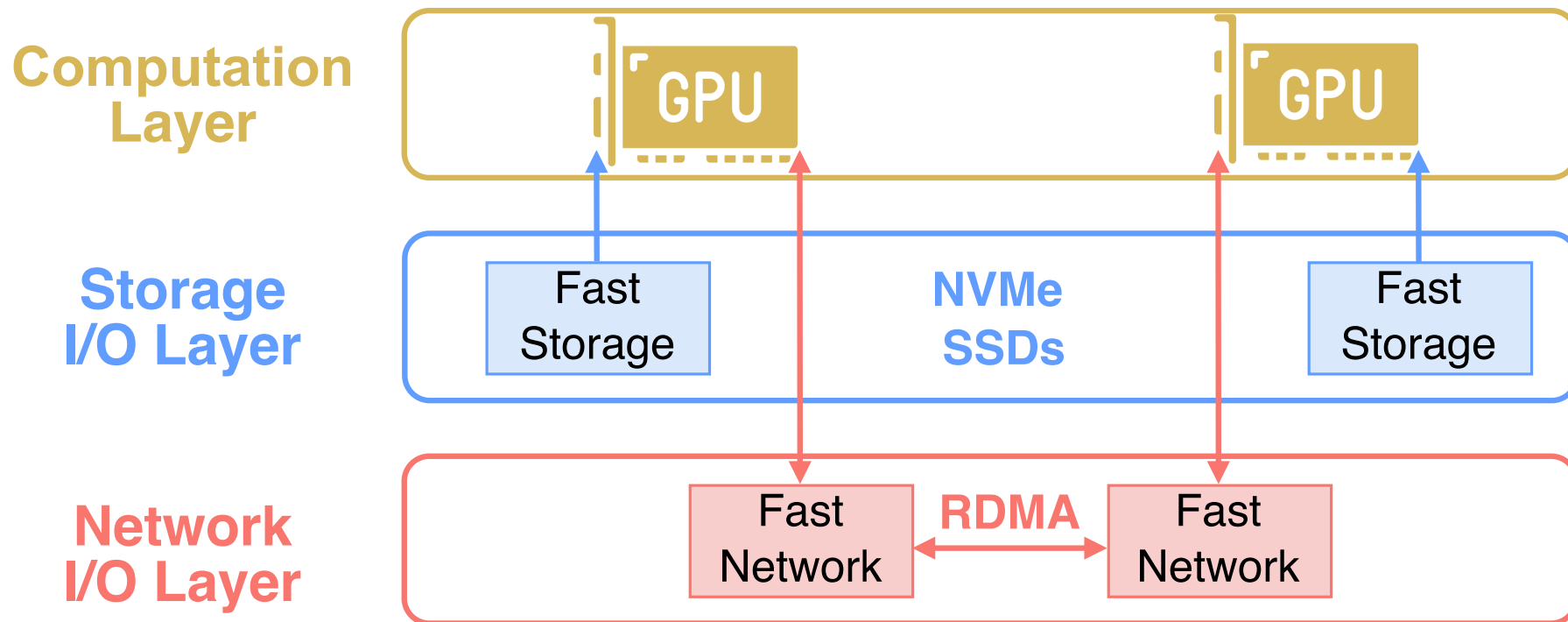


Challenge: Network and Storage I/O



PystachIO

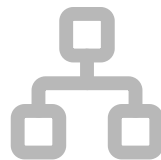
Distributed GPU Query Engine PystachIO



PystachIO: Storage



**Efficient
Storage
Access**



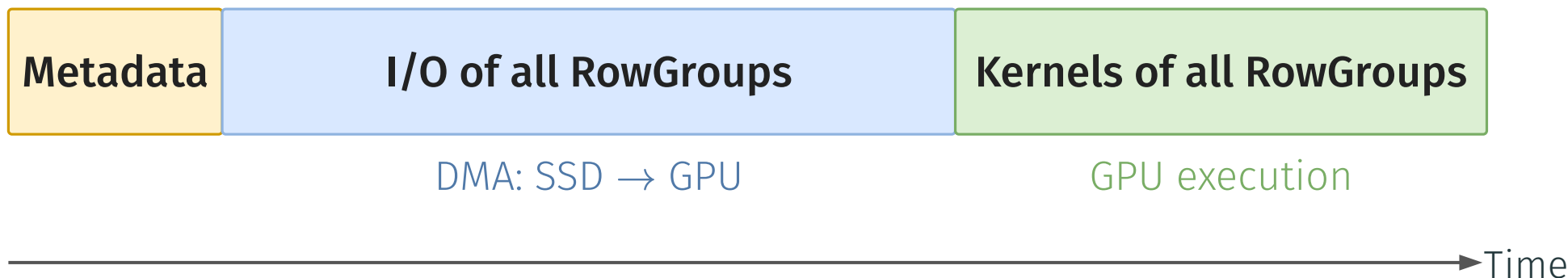
**Distributed
Join Over
Networks**



**Coordinating
Storage and
Network I/O**

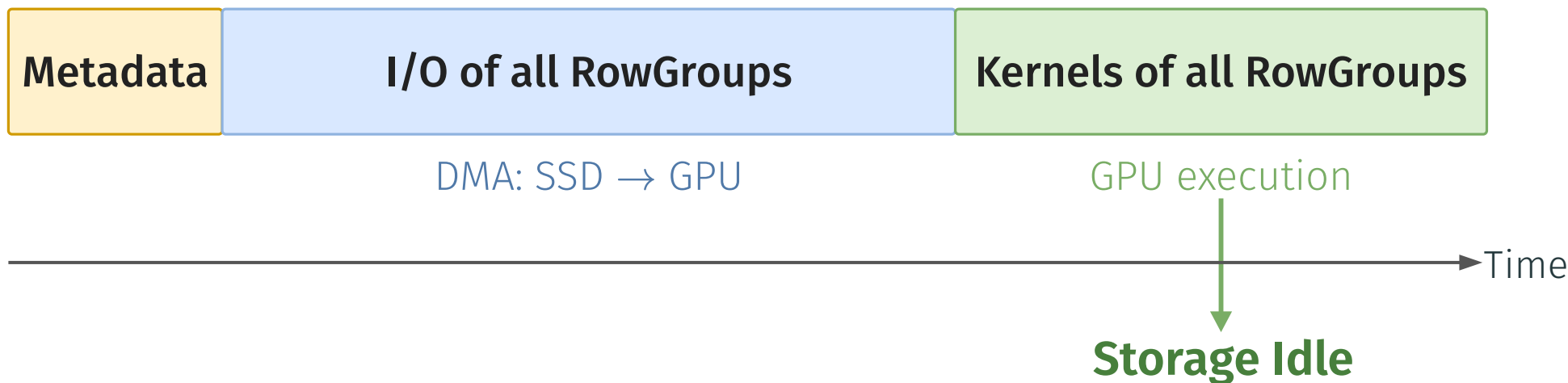
Inefficient Storage Access: Blocking Parquet Reader

`read_parquet()` in one CPU thread:



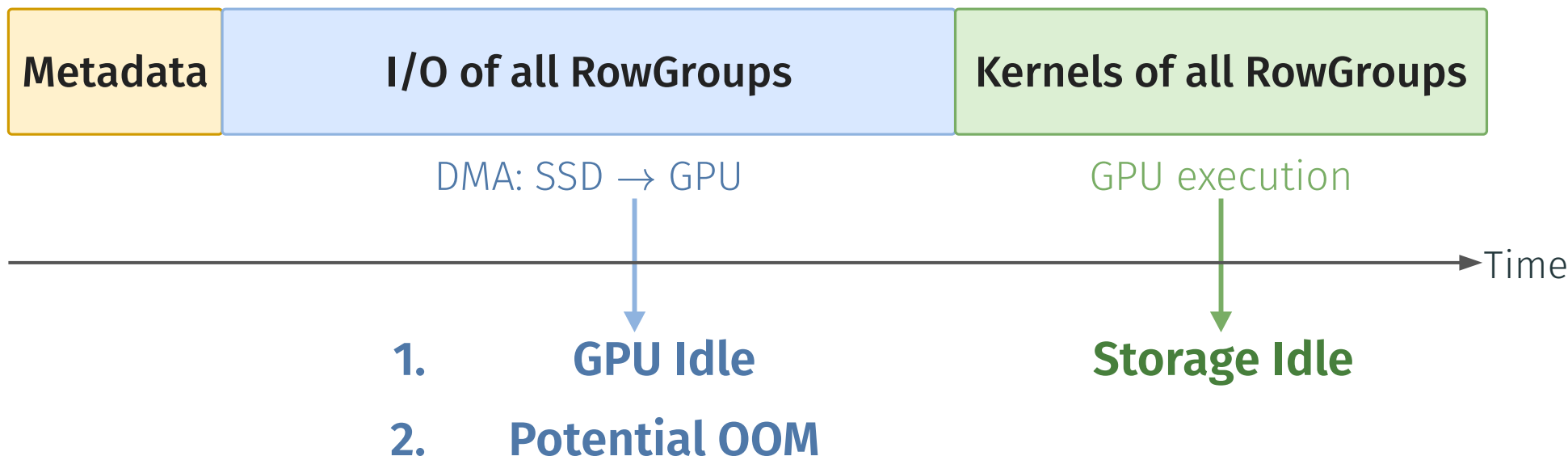
Inefficient Storage Access: Blocking Parquet Reader

`read_parquet()` in one CPU thread:

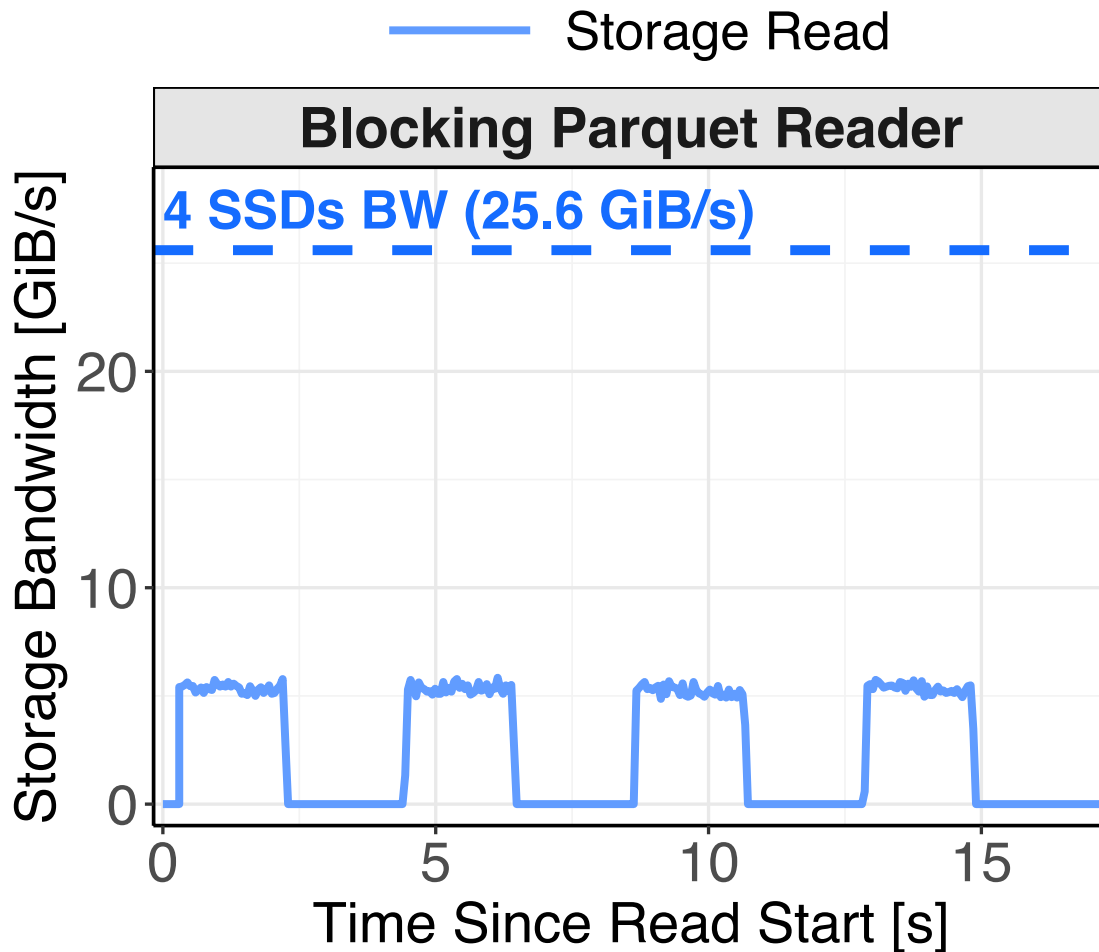


Inefficient Storage Access: Blocking Parquet Reader

`read_parquet()` in one CPU thread:

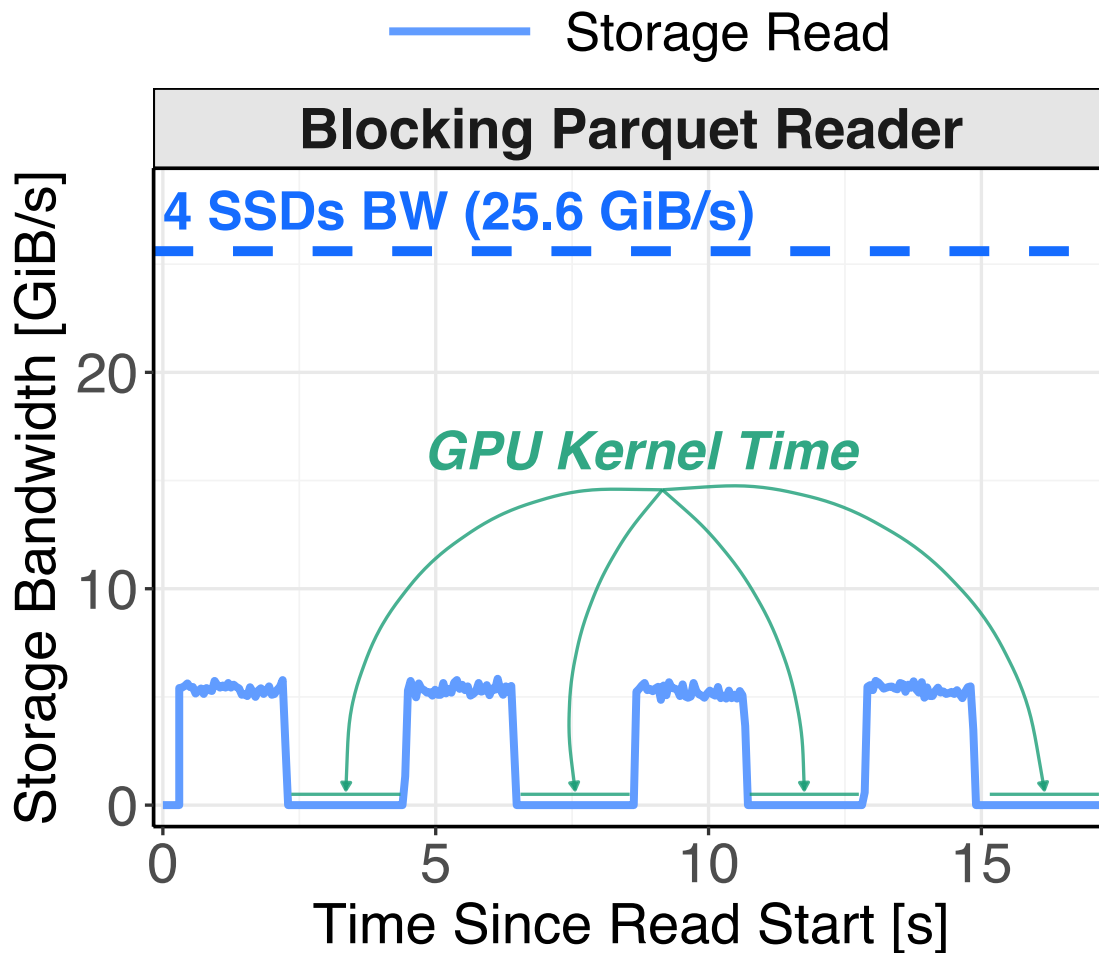


Inefficient Storage Access: Blocking Parquet Reader



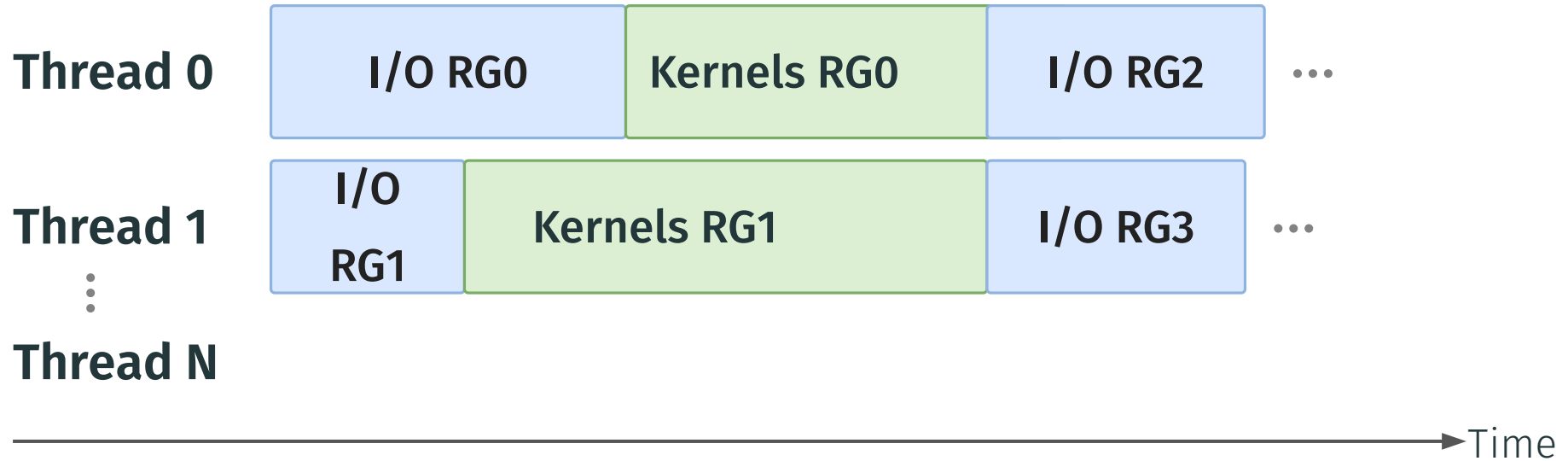
Storage microbenchmark with one GPU (A100, 40GB) and 4 NVMe SSD in PCIe 4.0.

Inefficient Storage Access: Blocking Parquet Reader

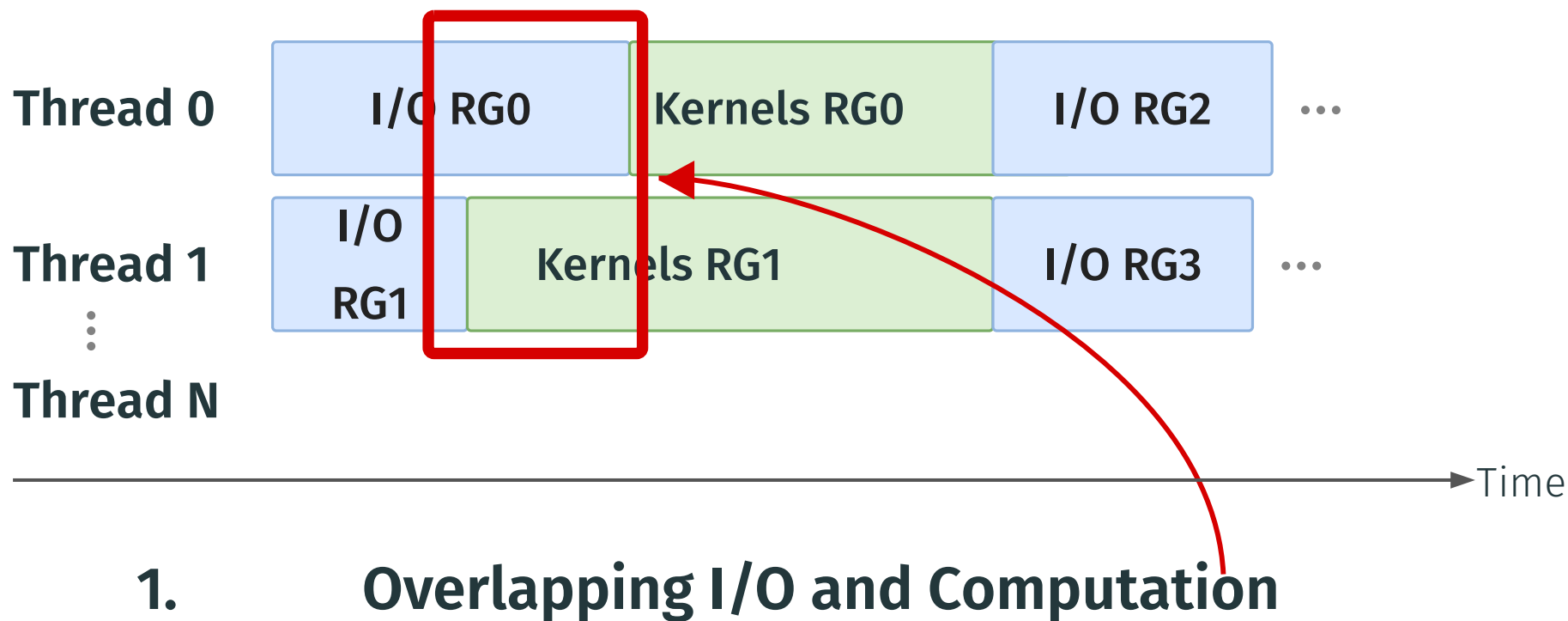


Storage microbenchmark with one GPU (A100, 40GB) and 4 NVMe SSD in PCIe 4.0.

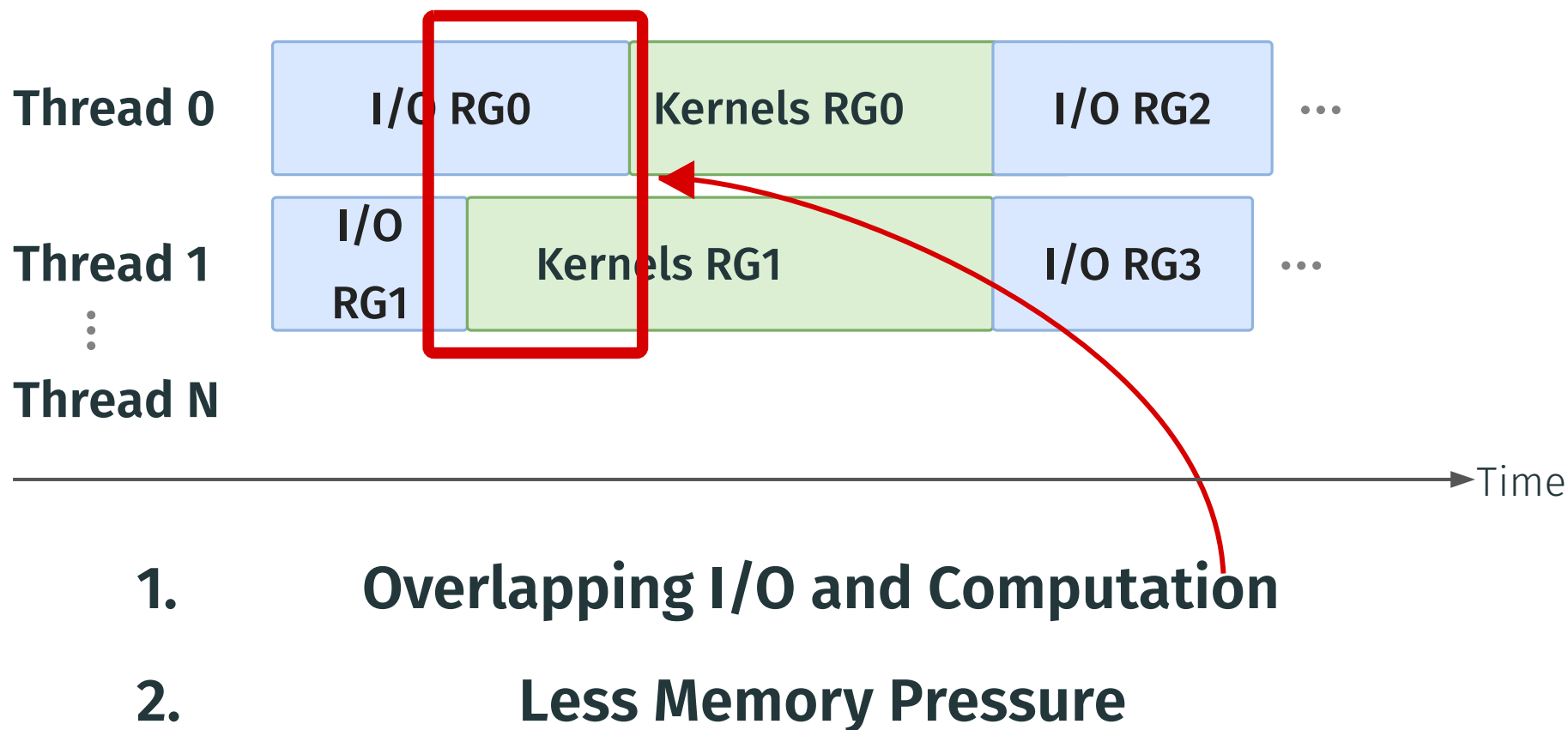
Efficient Storage Access: Chunking



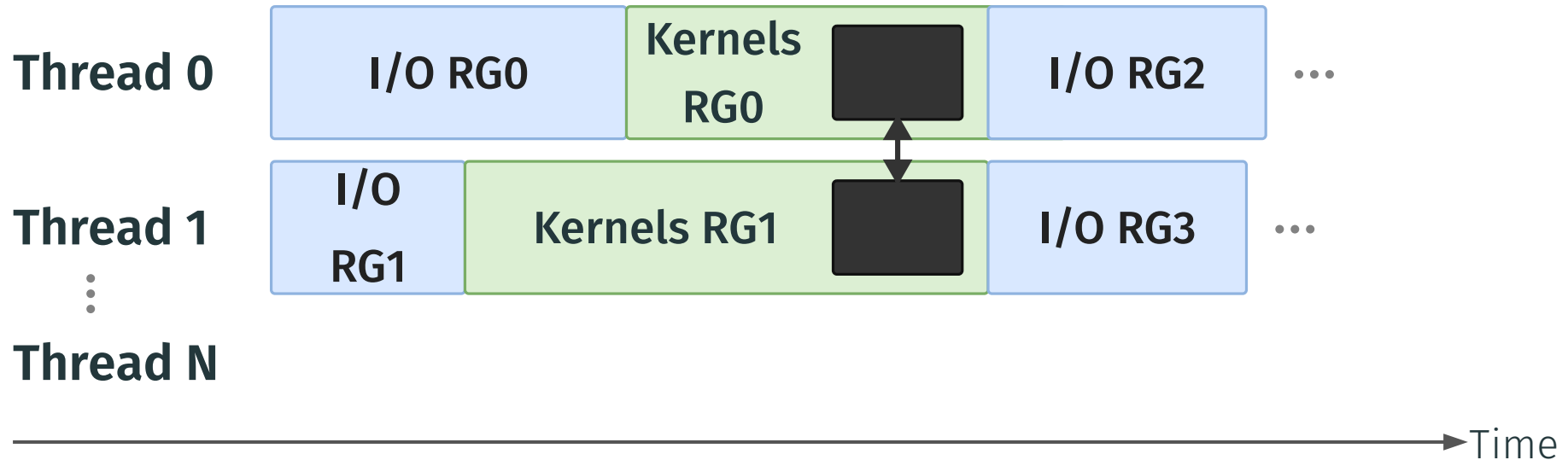
Efficient Storage Access: Chunking



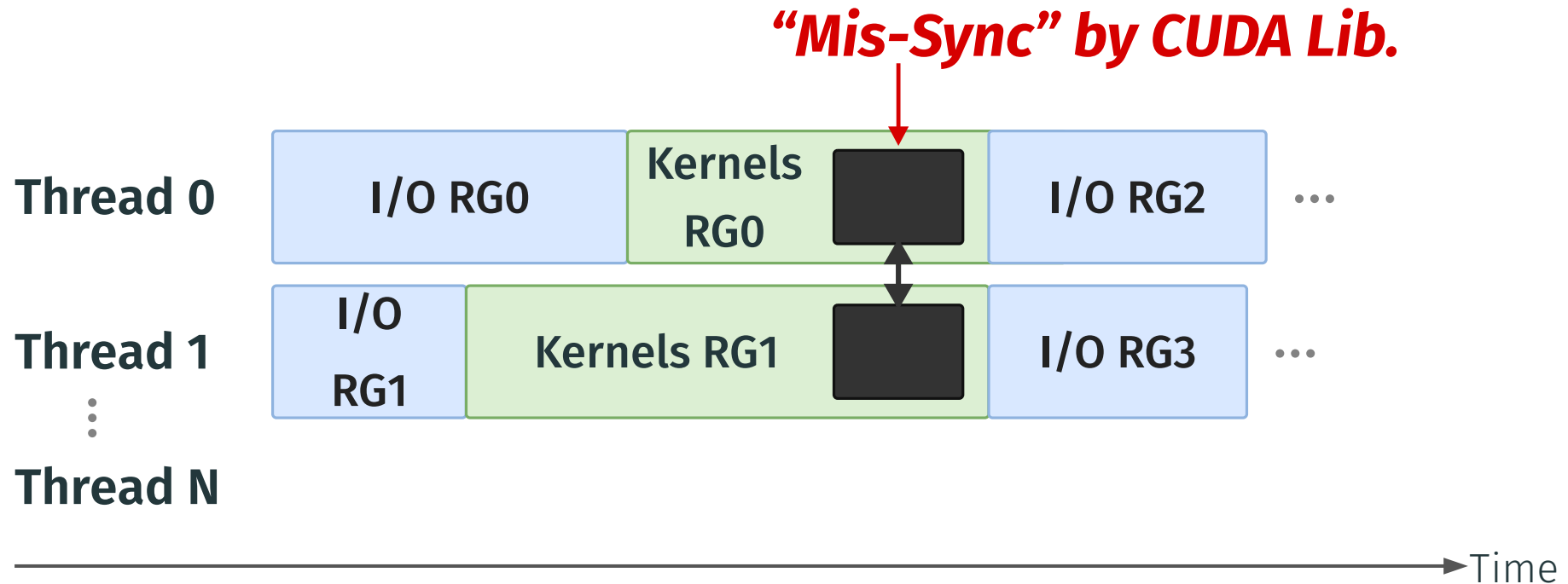
Efficient Storage Access: Chunking



Efficient Storage Access: Mis-Sync Elimination

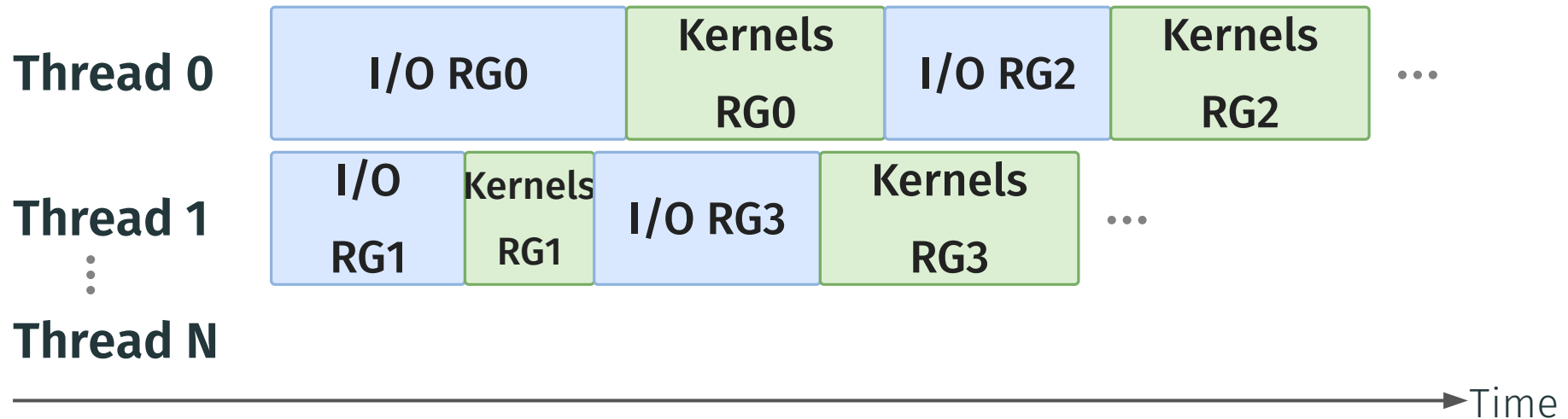


Efficient Storage Access: Mis-Sync Elimination

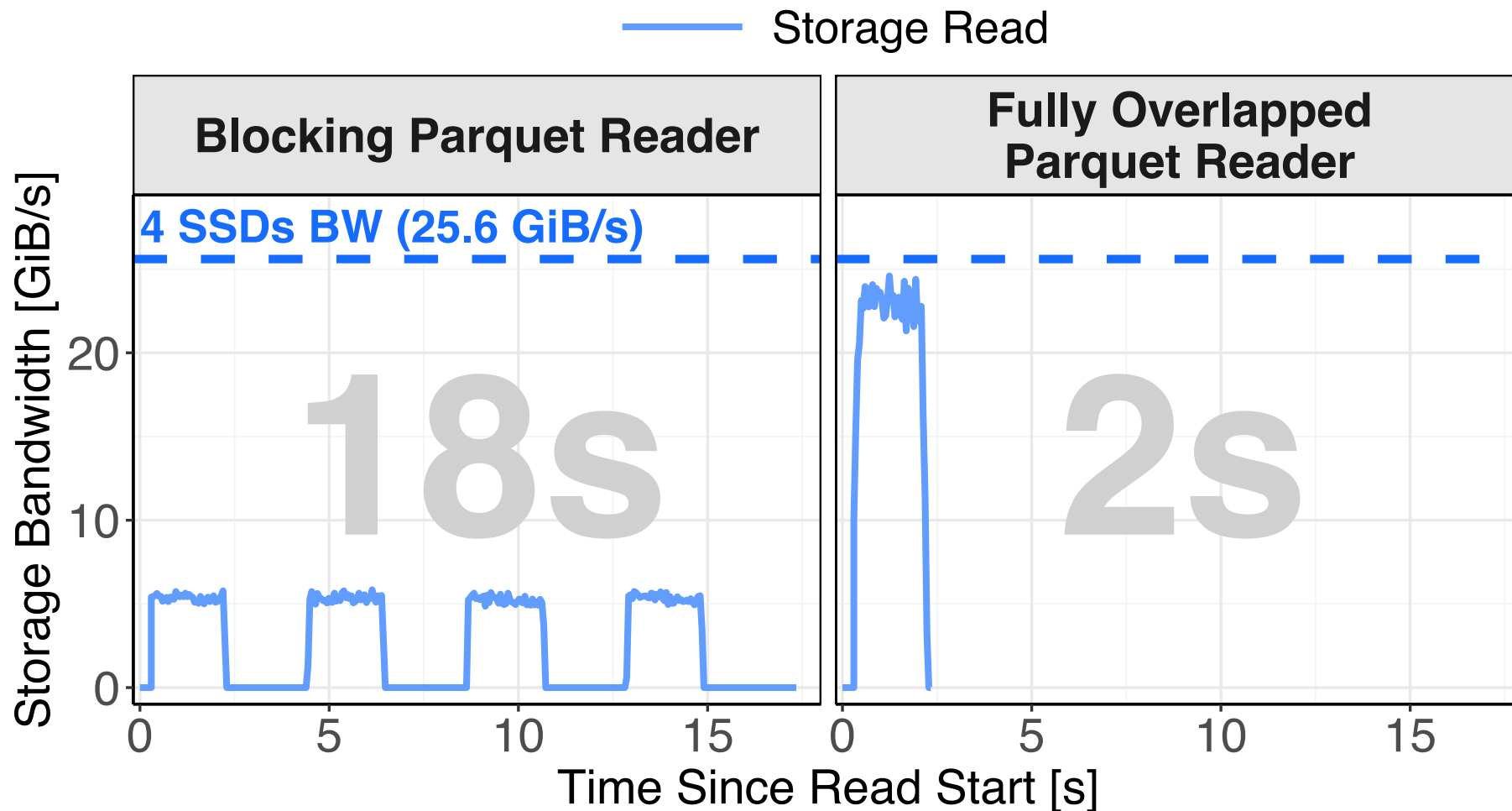


Efficient Storage Access: Mis-Sync Elimination

Fully Overlapped Parquet Reader

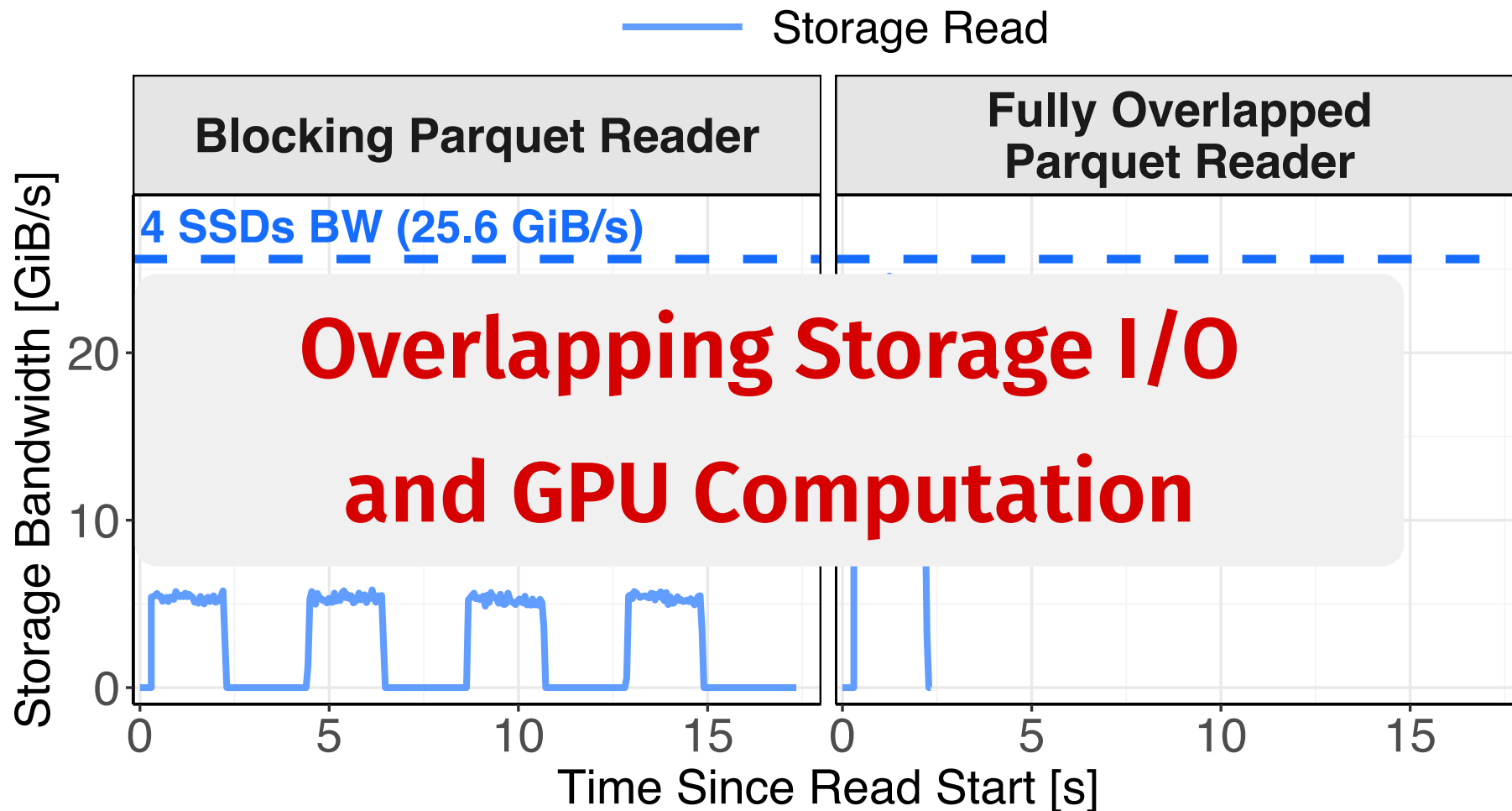


Efficient Storage Access: Mis-Sync Elimination



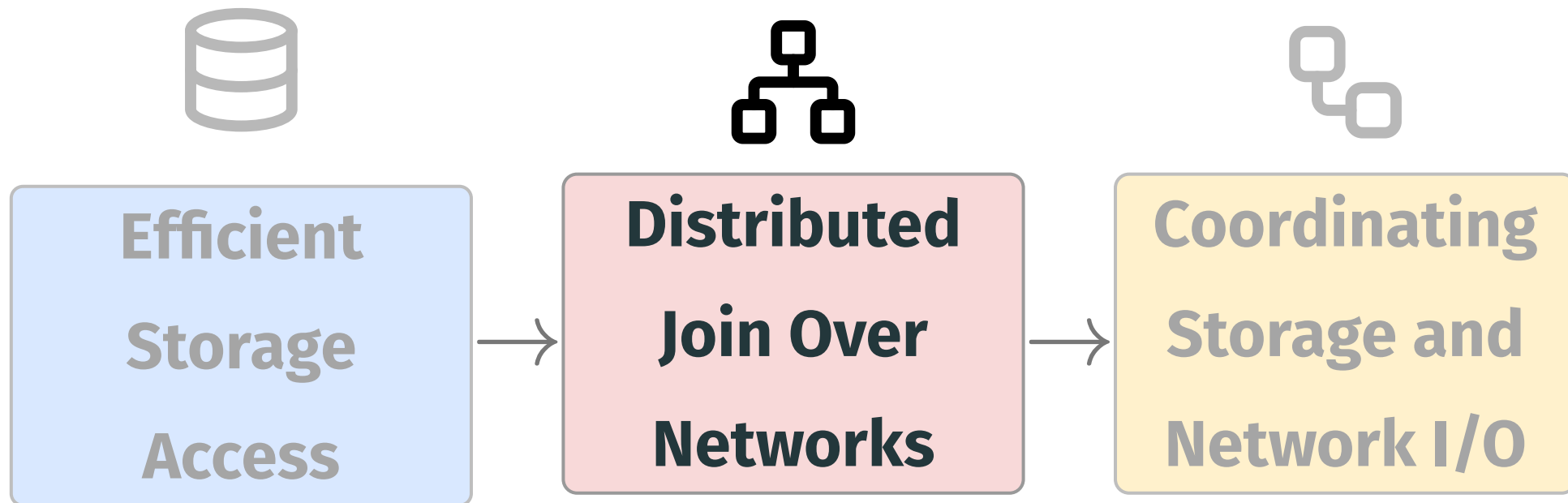
Storage microbenchmark with one GPU (A100, 40GB) and 4 NVMe SSD in PCIe 4.0.

Efficient Storage Access: Mis-Sync Elimination

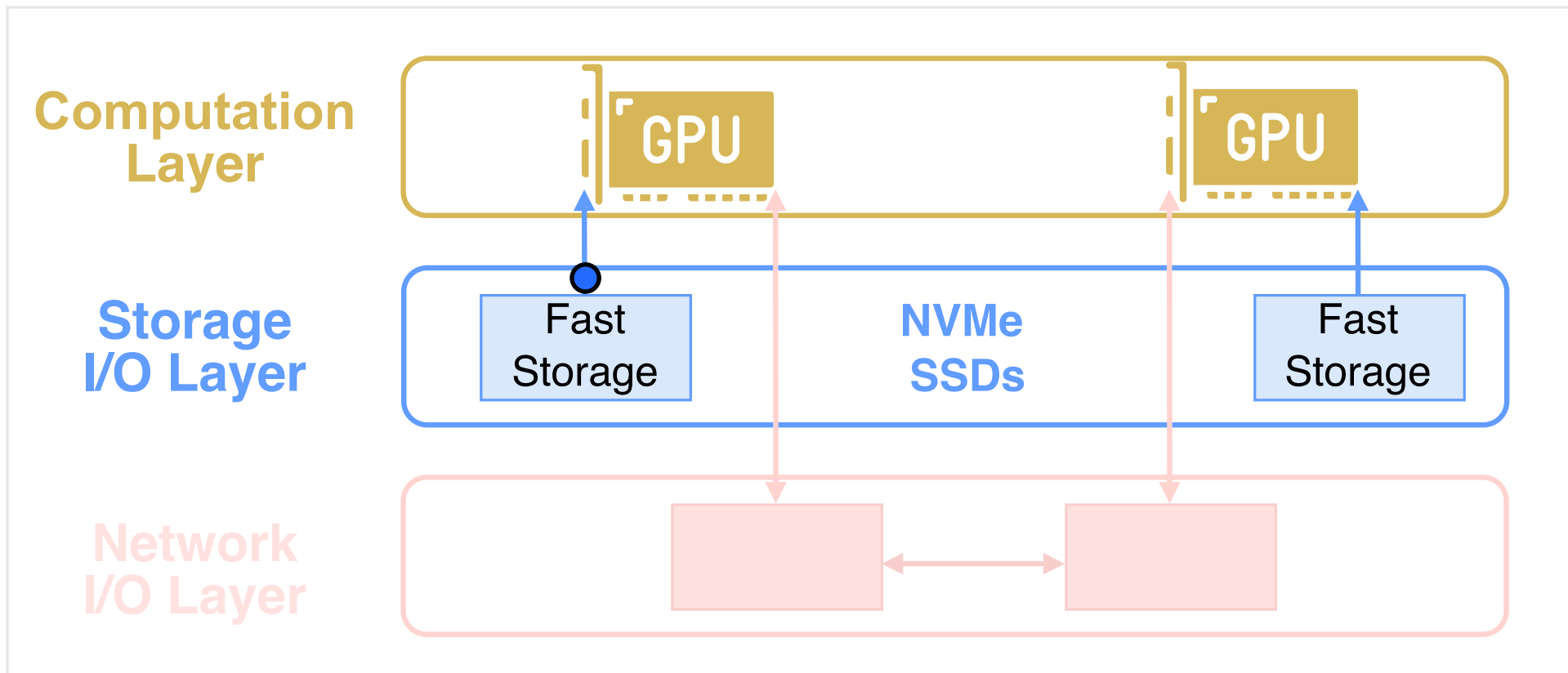


Storage microbenchmark with one GPU (A100, 40GB) and 4 NVMe SSD in PCIe 4.0.

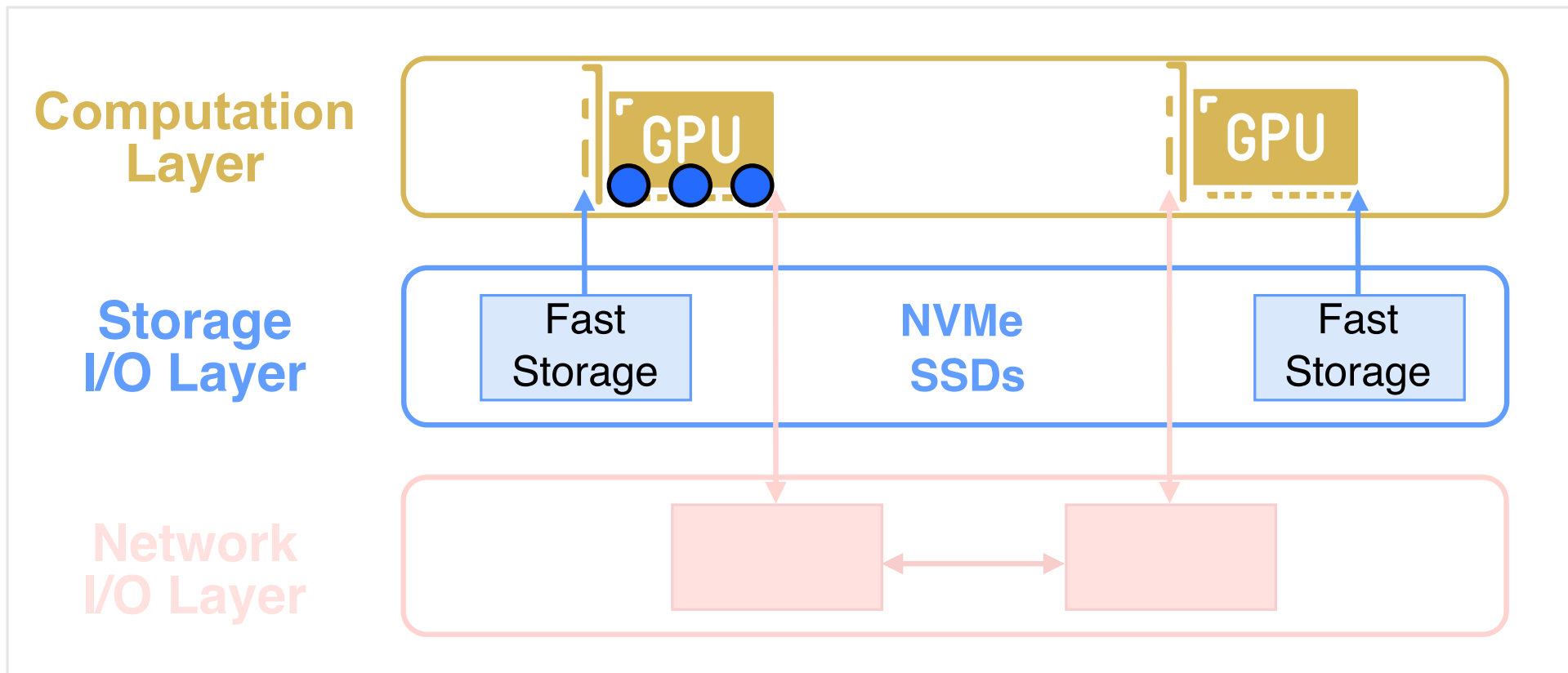
PystachIO: Network



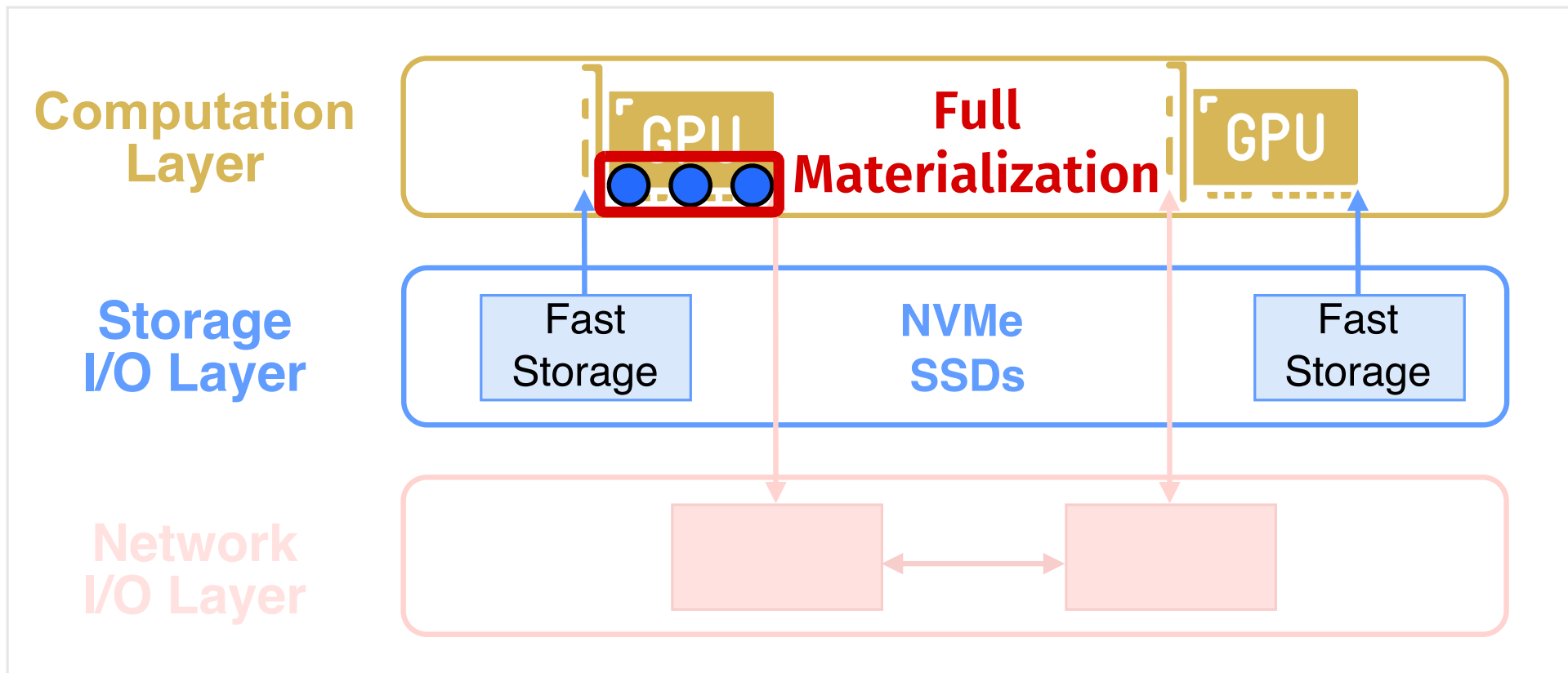
Run a Distributed Query



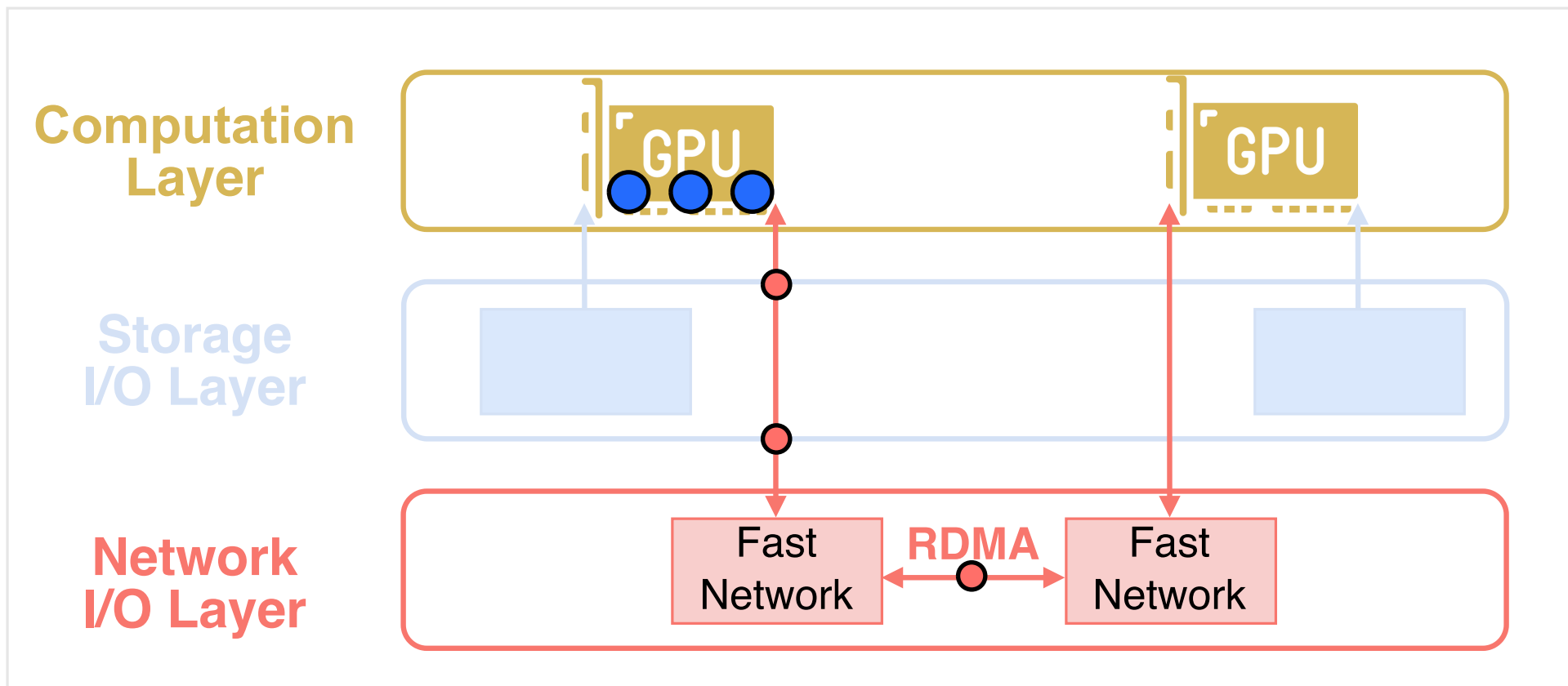
Run a Distributed Query



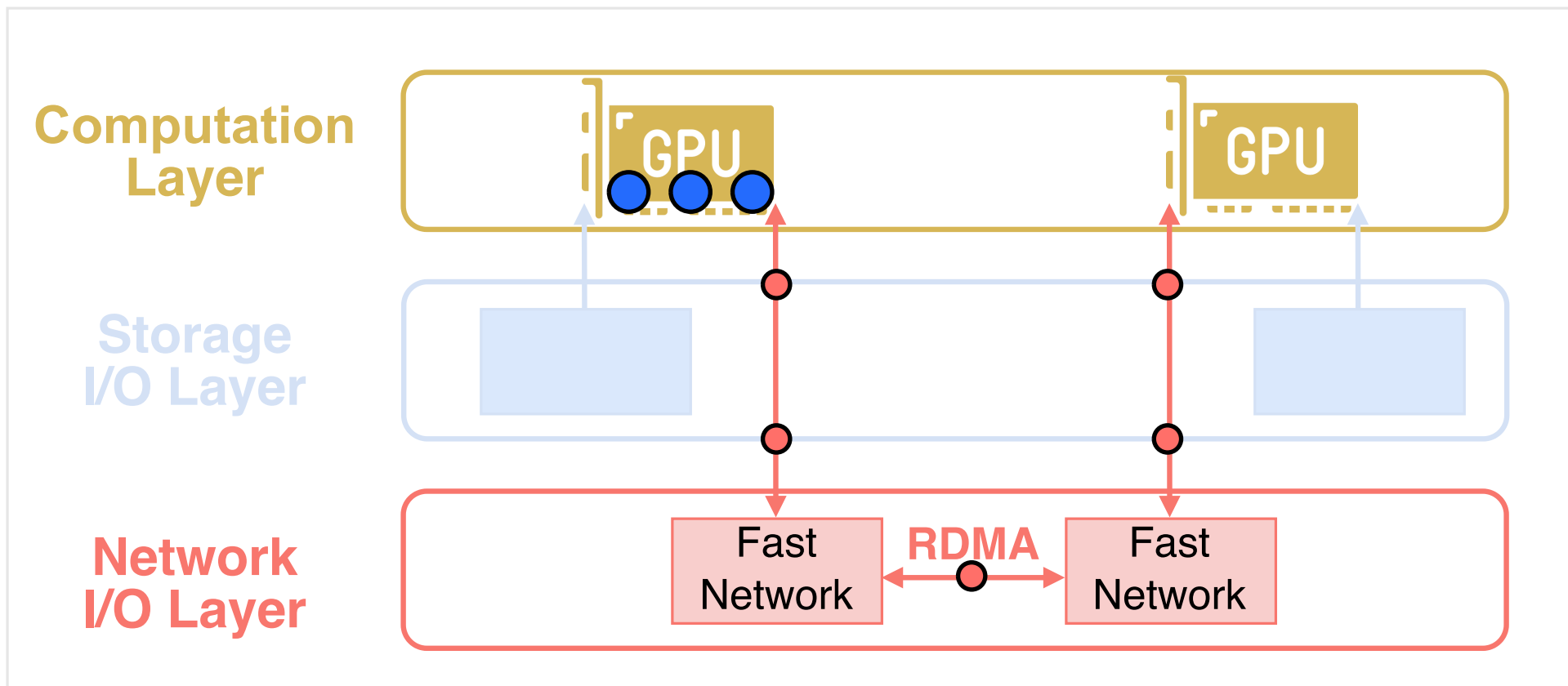
Run a Distributed Query



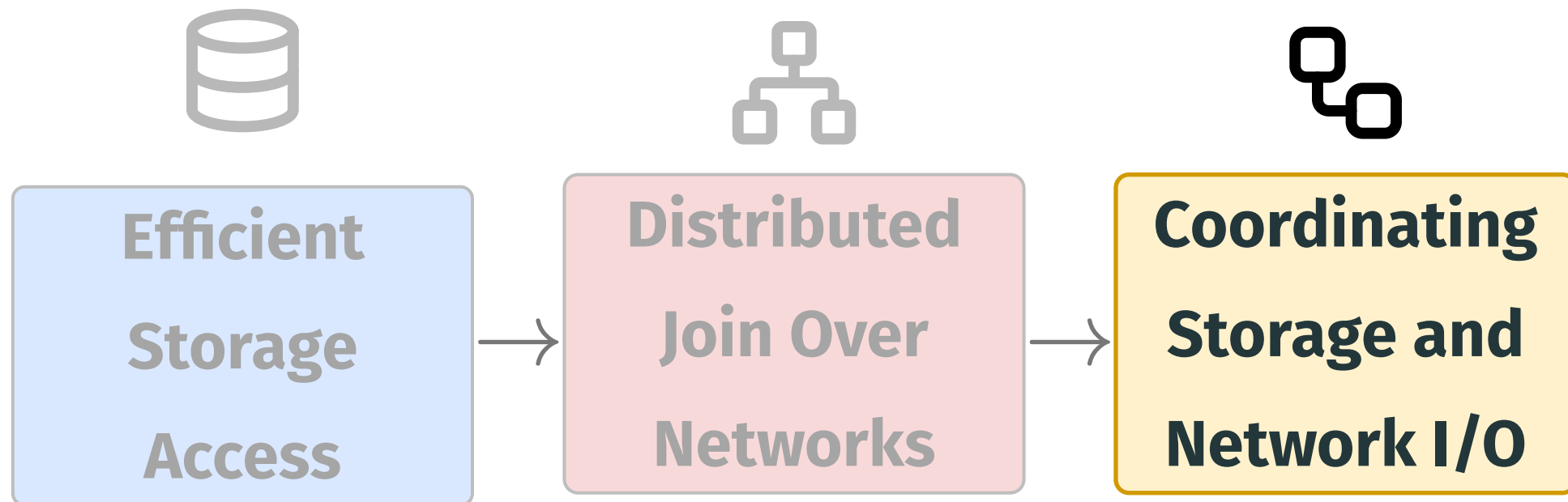
Run a Distributed Query



Run a Distributed Query



PystachIO: Coordinating I/O



Fast Storage, then Fast Network

- Storage optimized
- Network optimized

Storage

Networks & Join

Time

Fast Storage, then Fast Network

- Storage optimized
- Network optimized

Storage

Networks & Join

→ Time

- Slow Runtime
- High Risk of OOM
- Only One I/O Device Active

Fast Storage, then Fast Network

- Storage
- Network

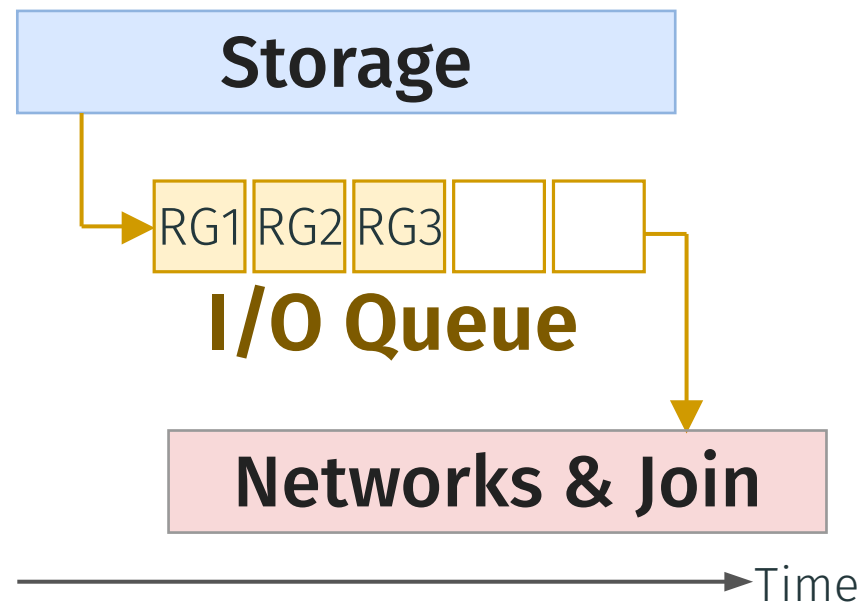
**Is It Possible to Overlap
Storage and Network?**

Networks & Join

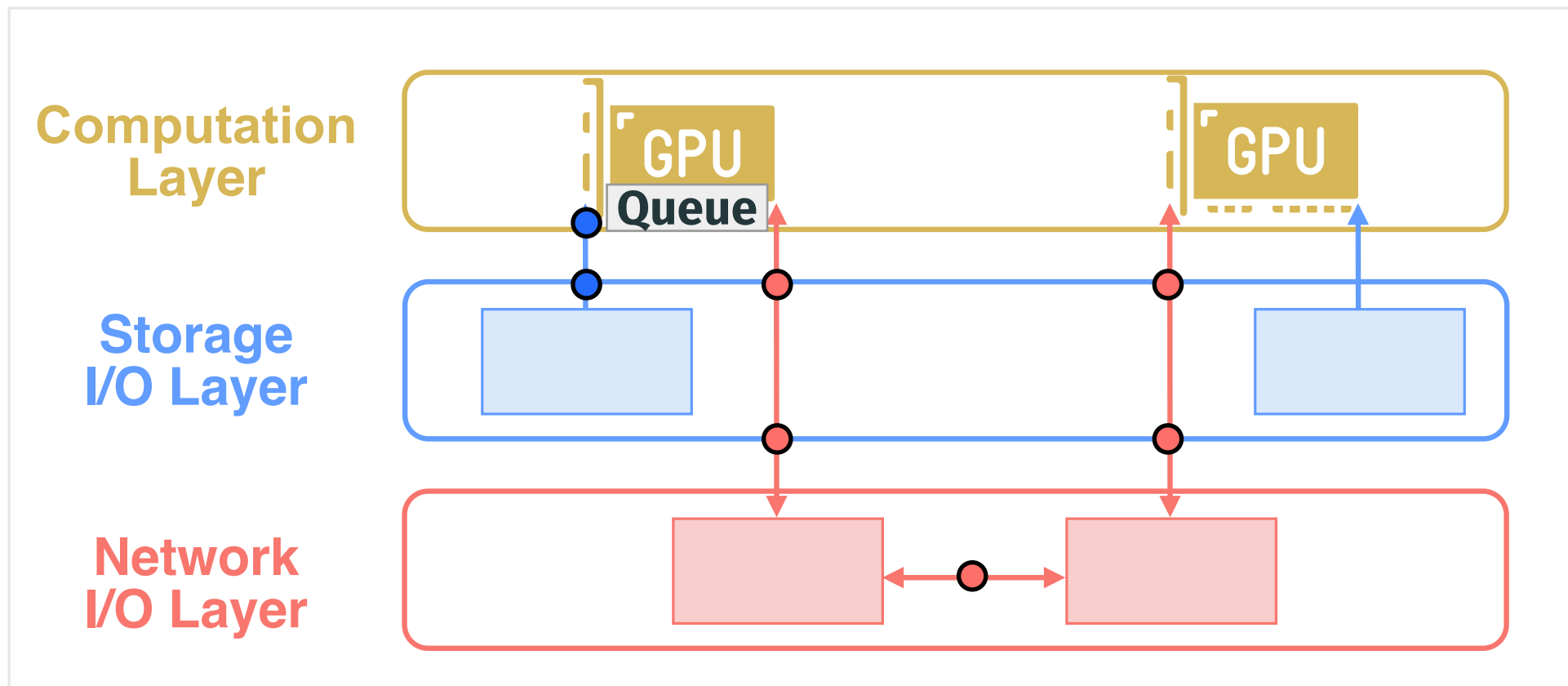
Time

I/O Queue For Coordination

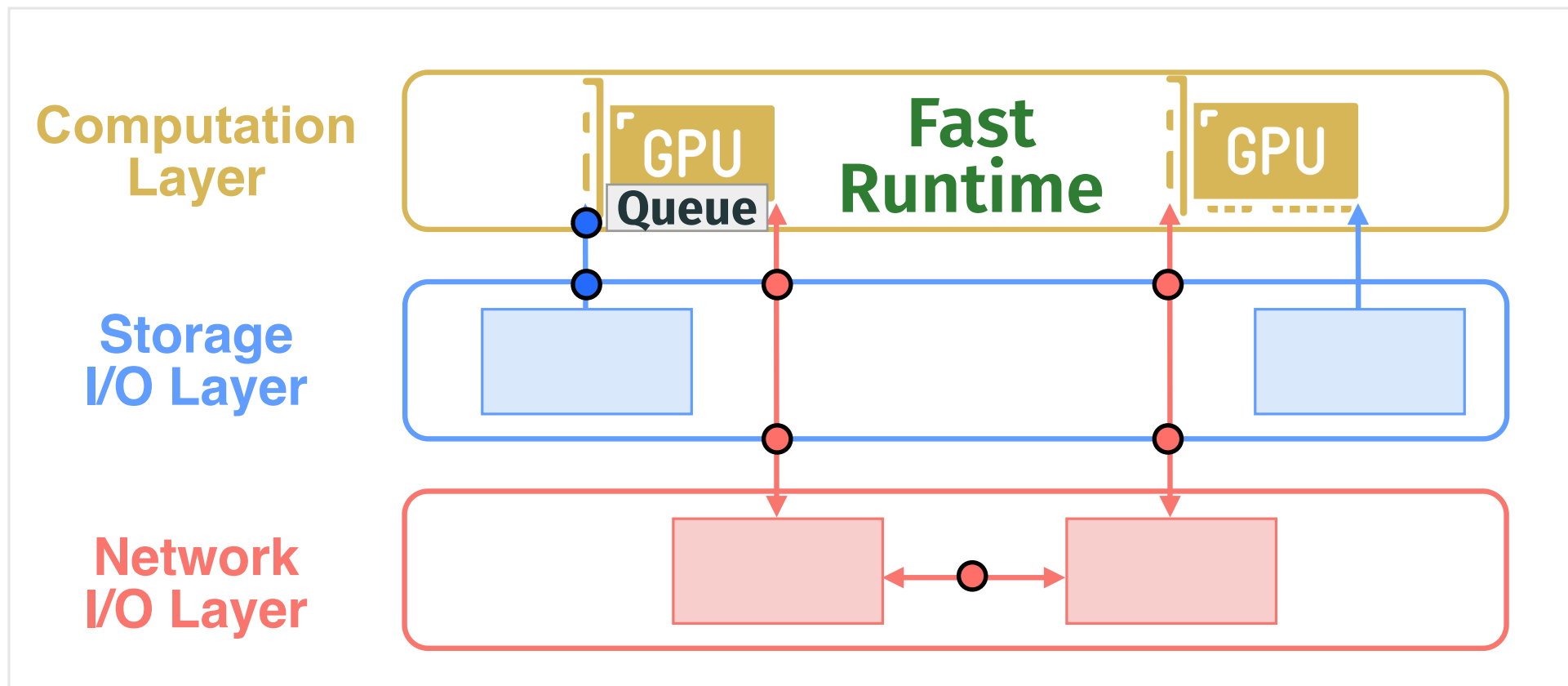
- Storage: enqueueing
- Network & Join: dequeuing
- Then network shuffling



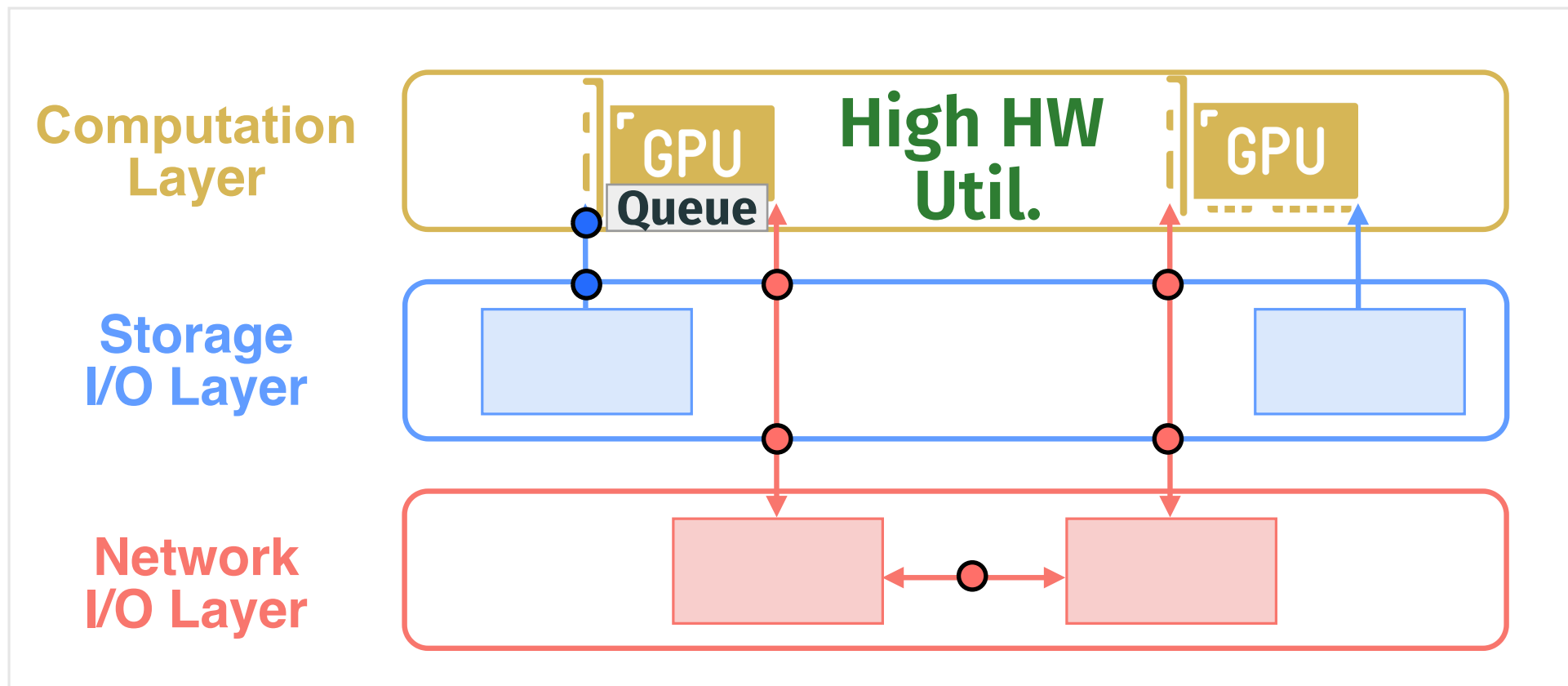
Overlapping Network, Storage, and Computation



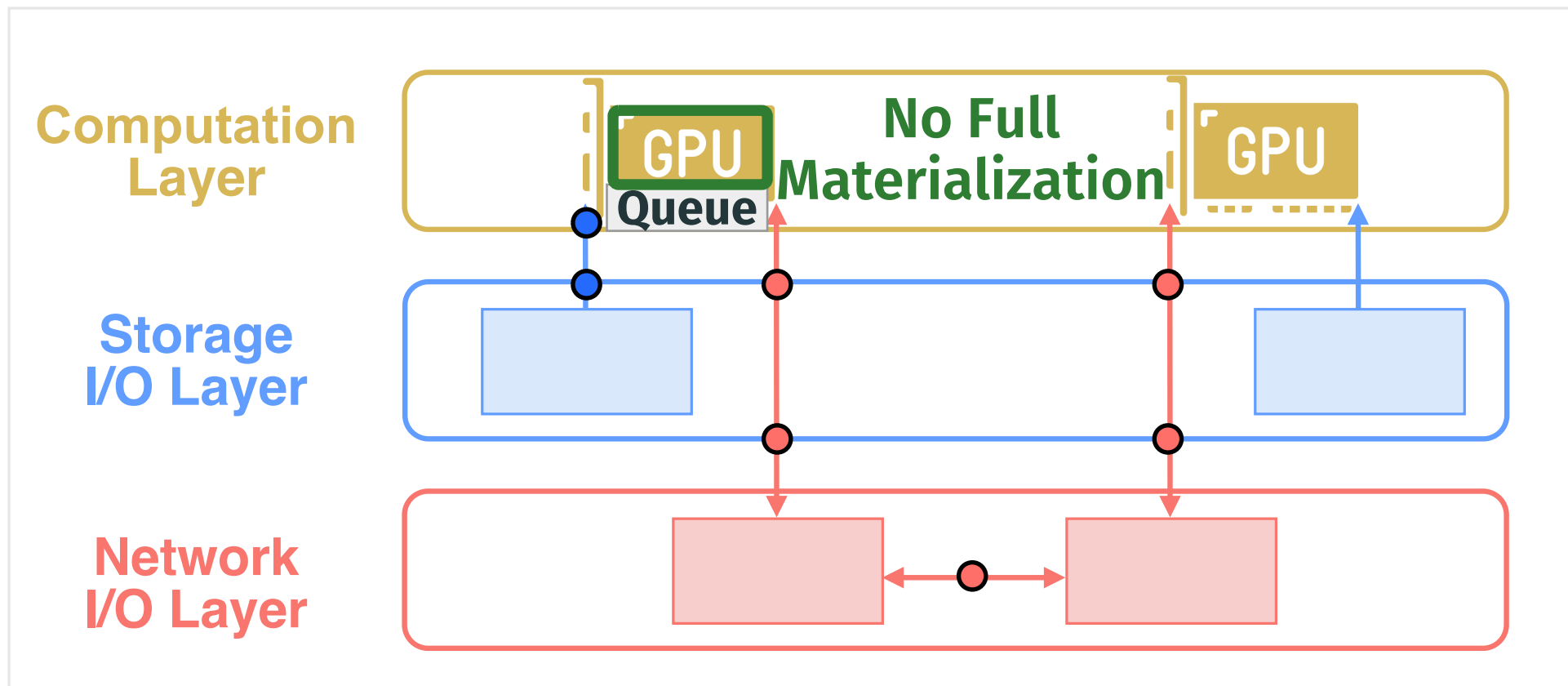
Overlapping Network, Storage, and Computation



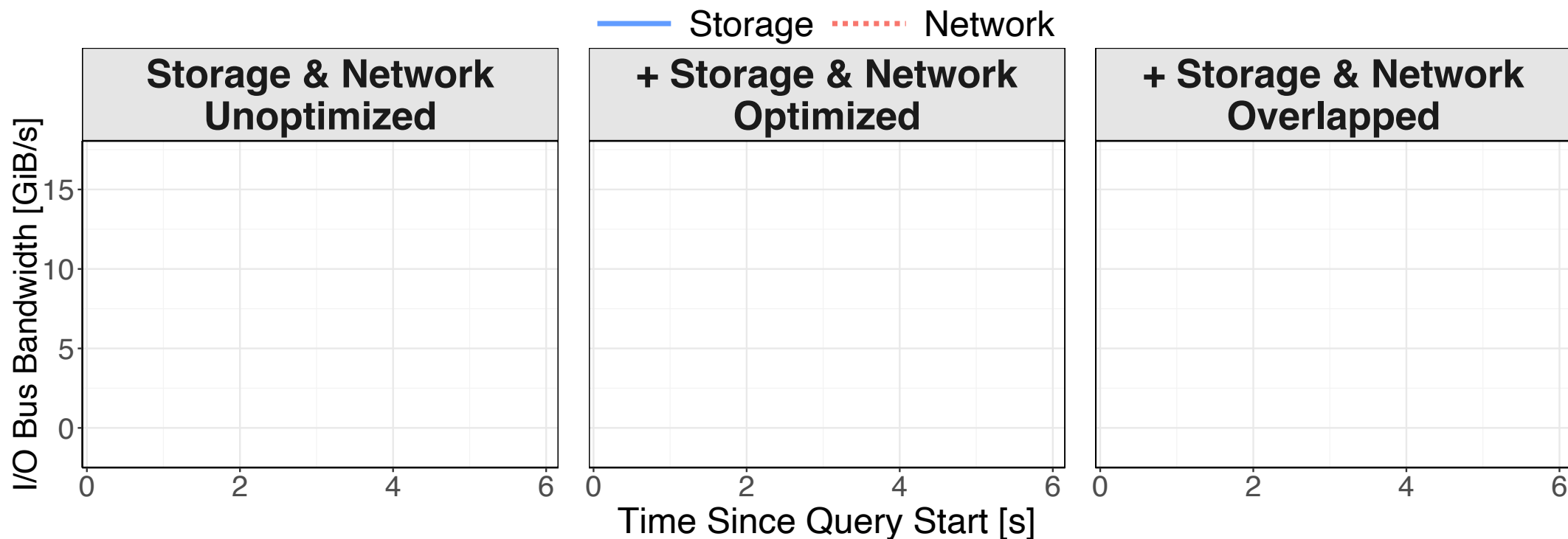
Overlapping Network, Storage, and Computation



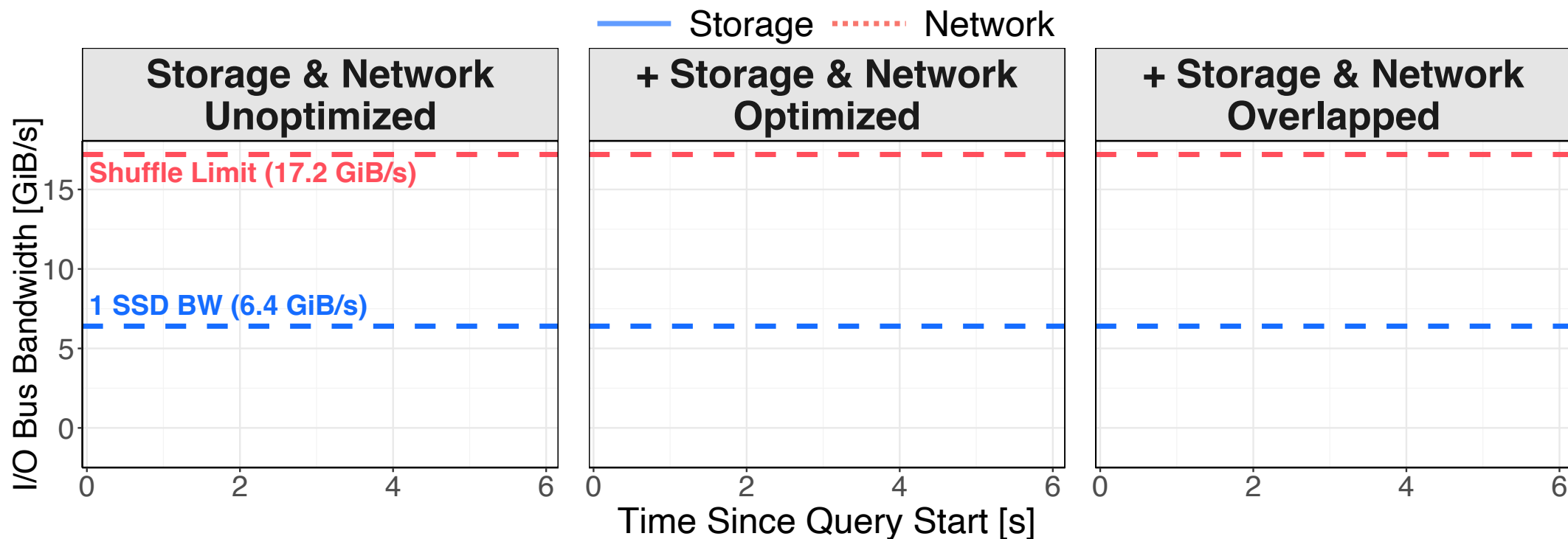
Overlapping Network, Storage, and Computation



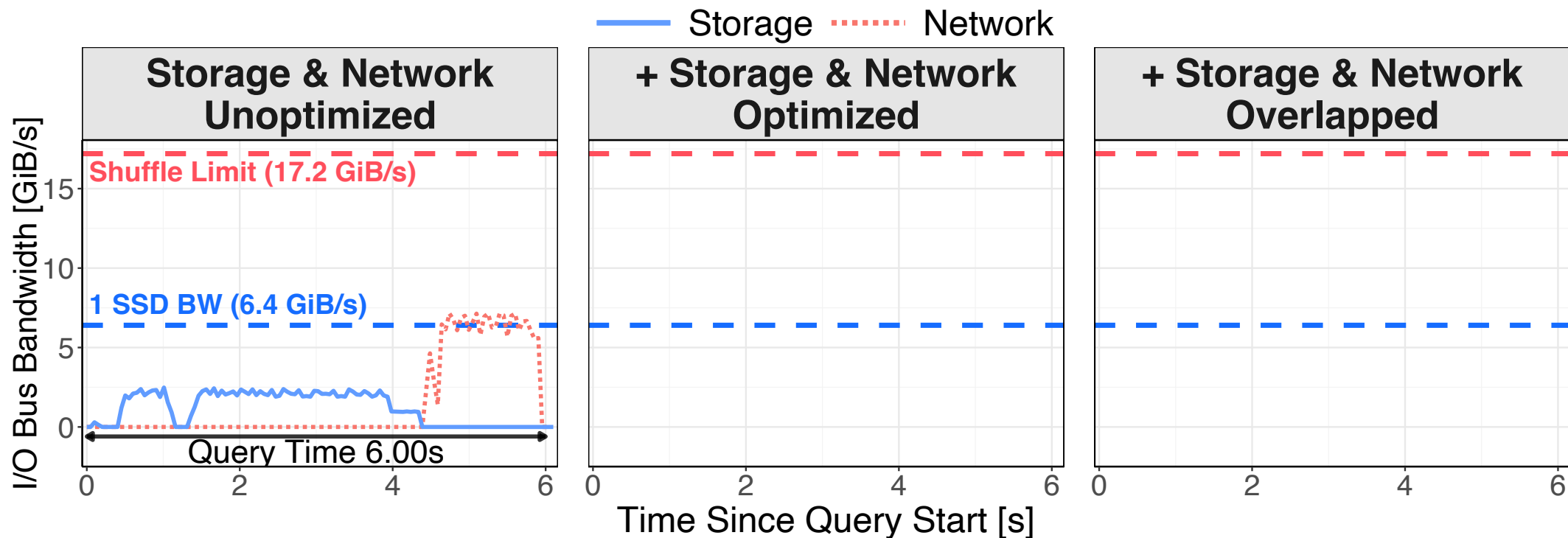
Overlapping Network, Storage, and Computation



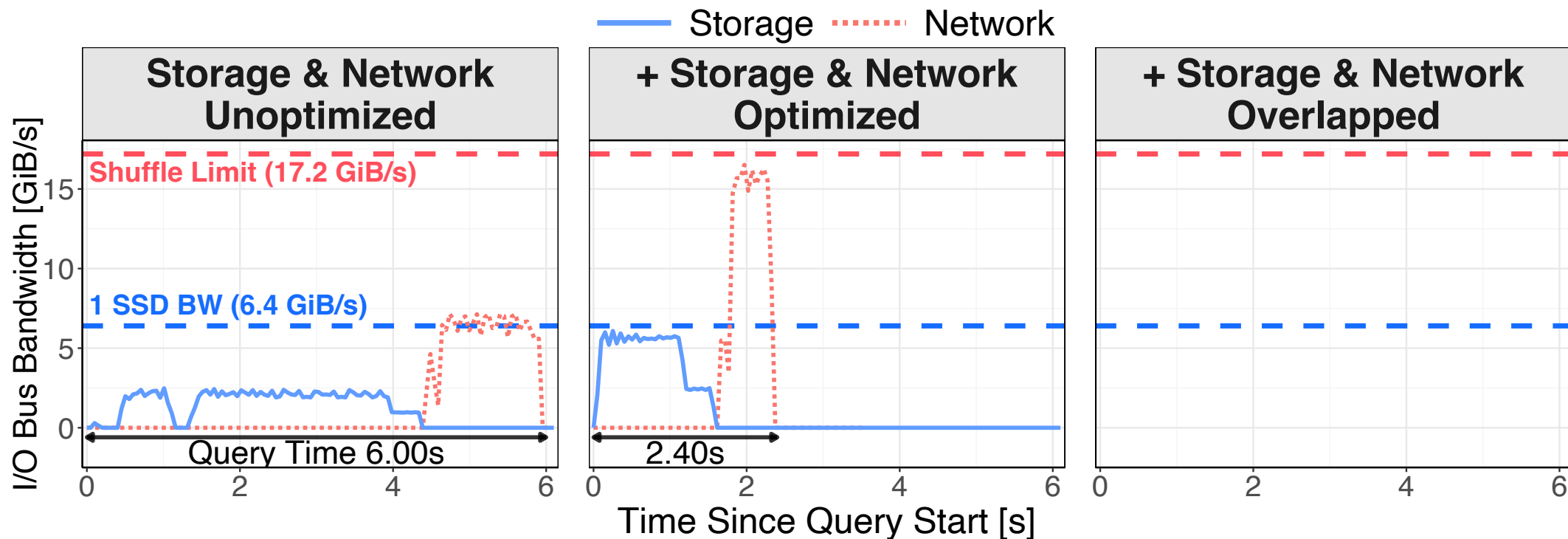
Overlapping Network, Storage, and Computation



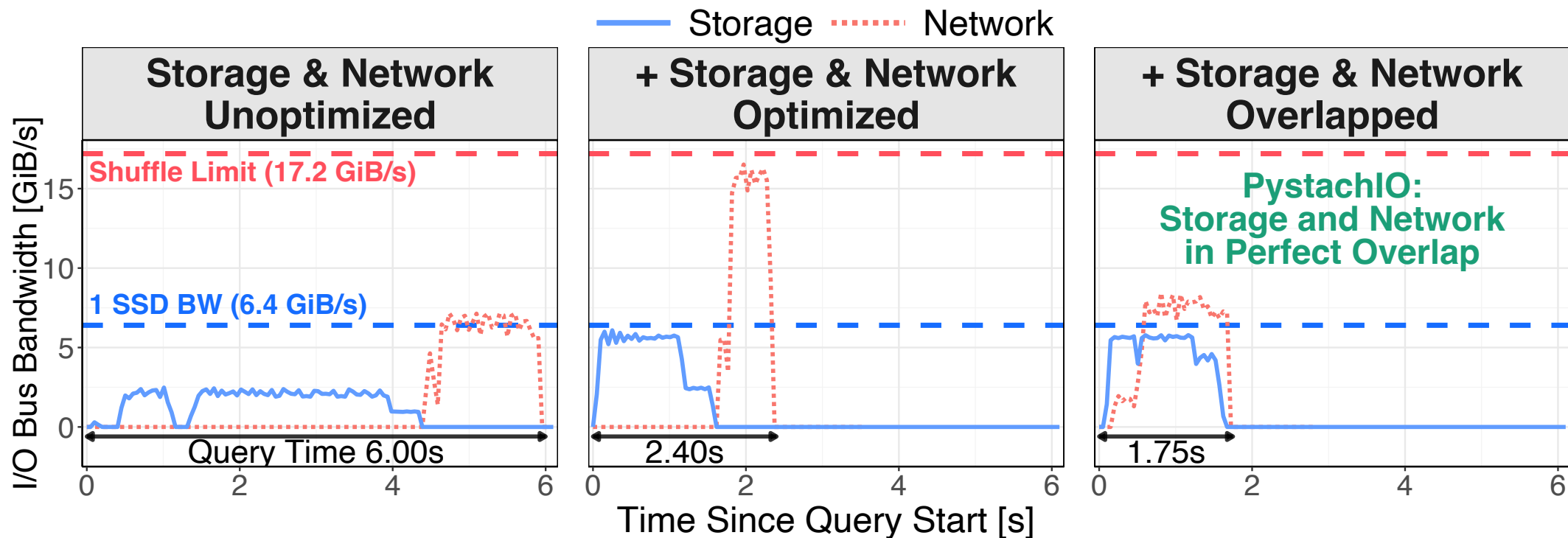
Overlapping Network, Storage, and Computation



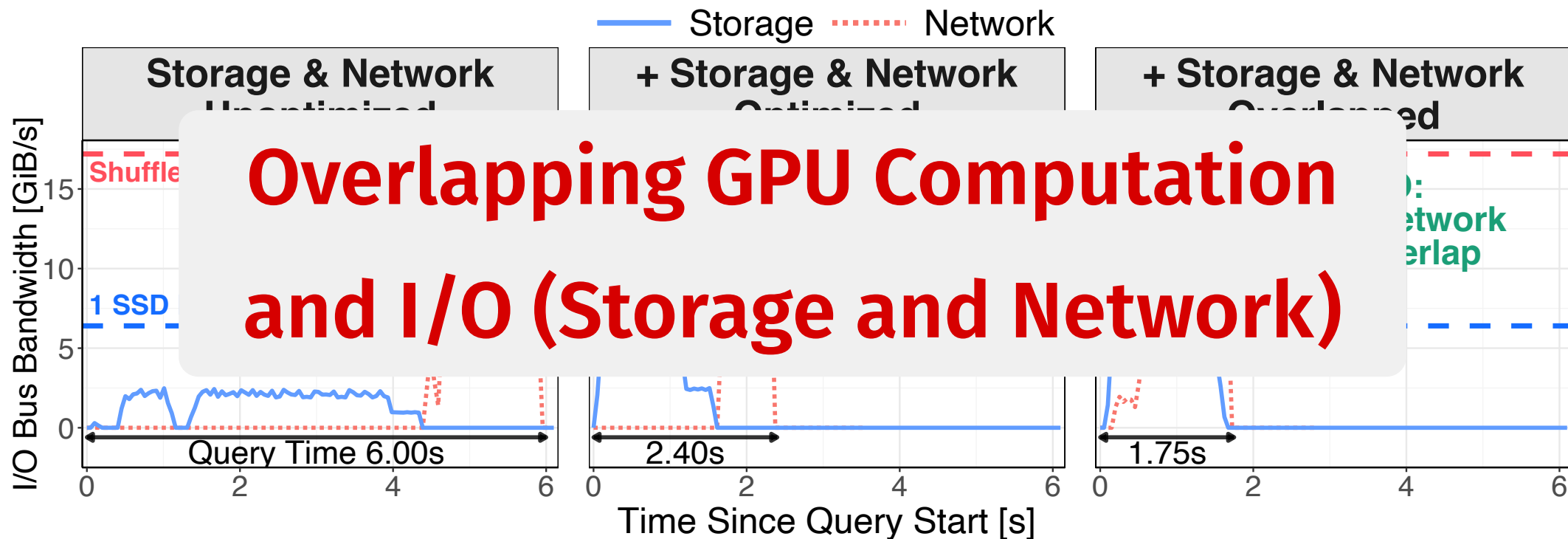
Overlapping Network, Storage, and Computation

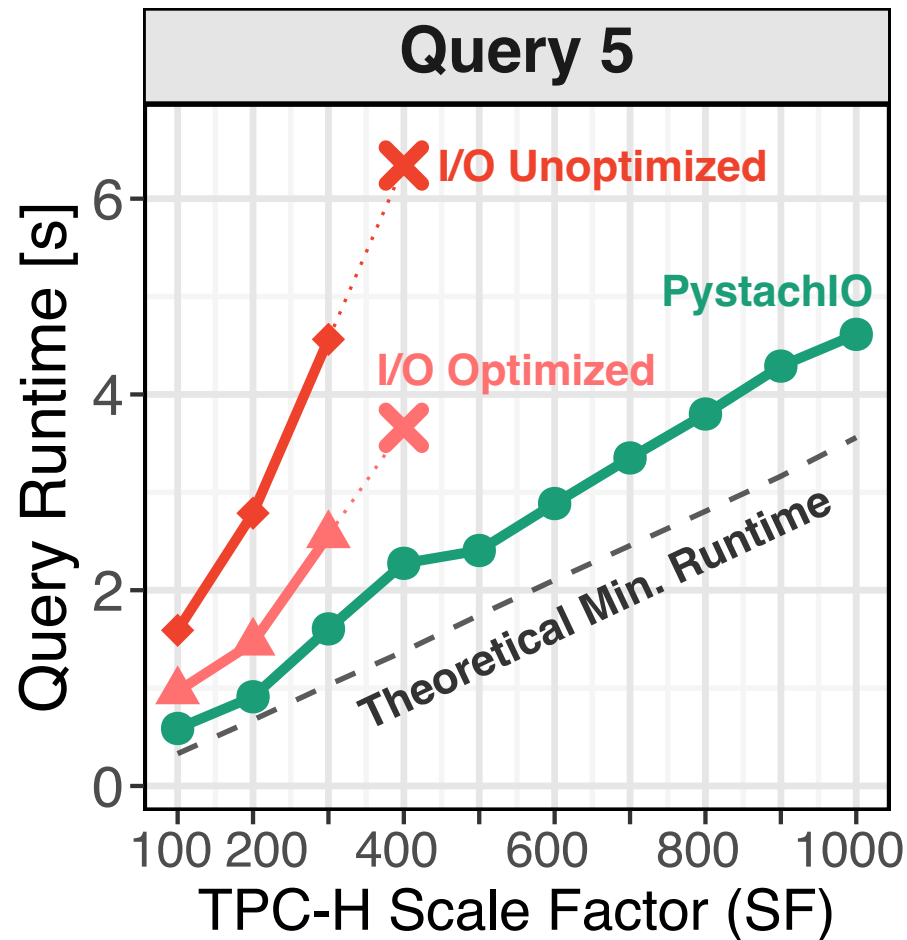


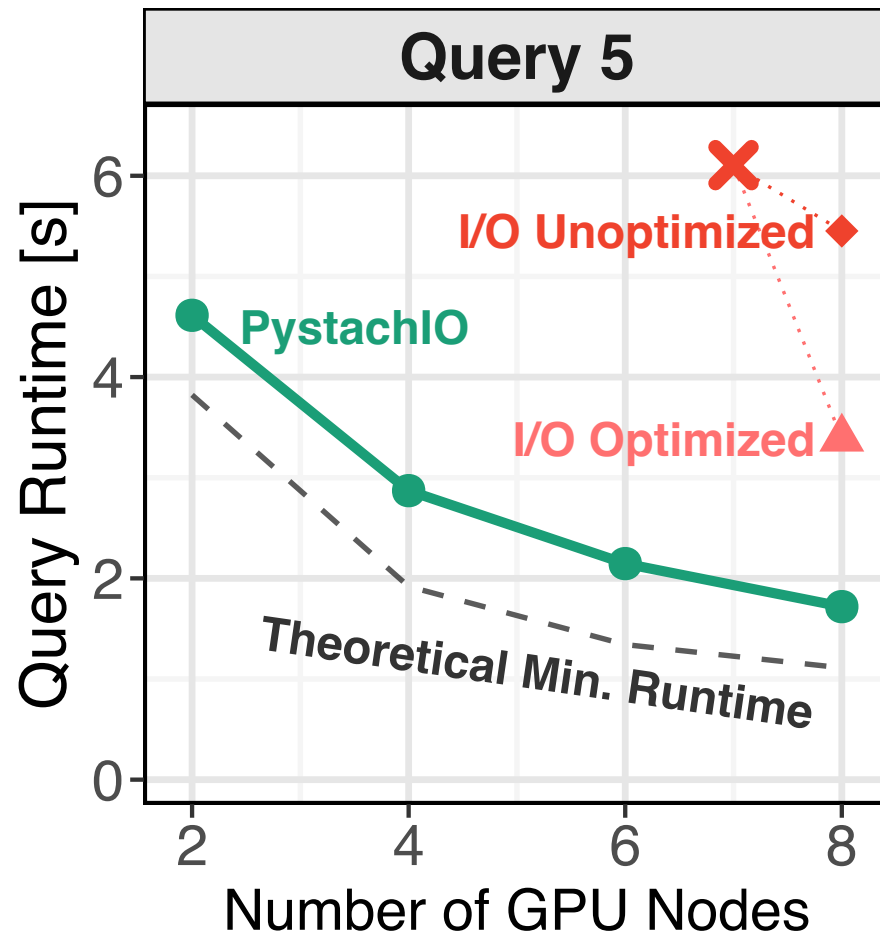
Overlapping Network, Storage, and Computation



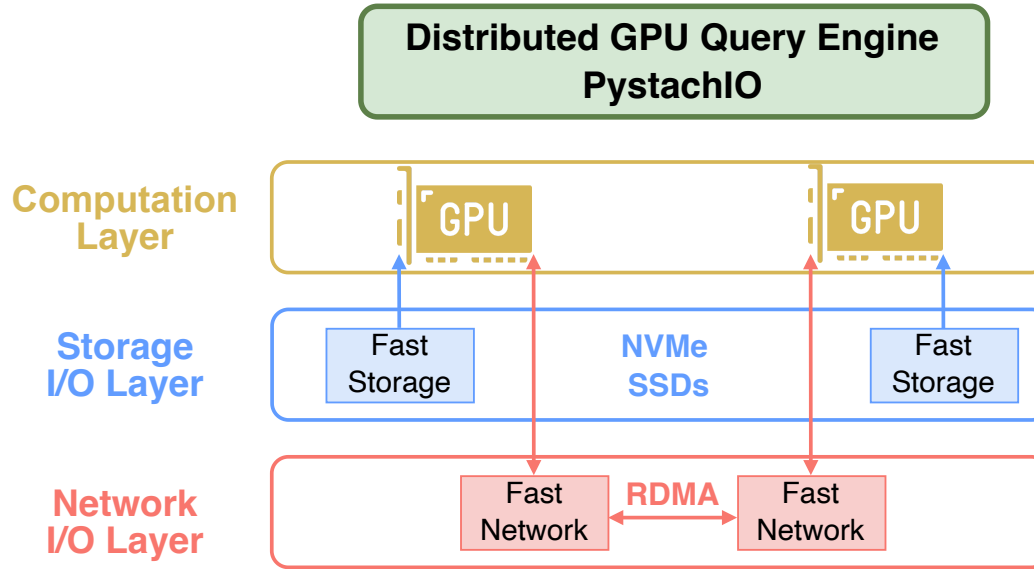
Overlapping Network, Storage, and Computation







Conclusion



- Efficient I/O
- Overlapping GPU & I/O
- I/O Coordination: Storage & Network



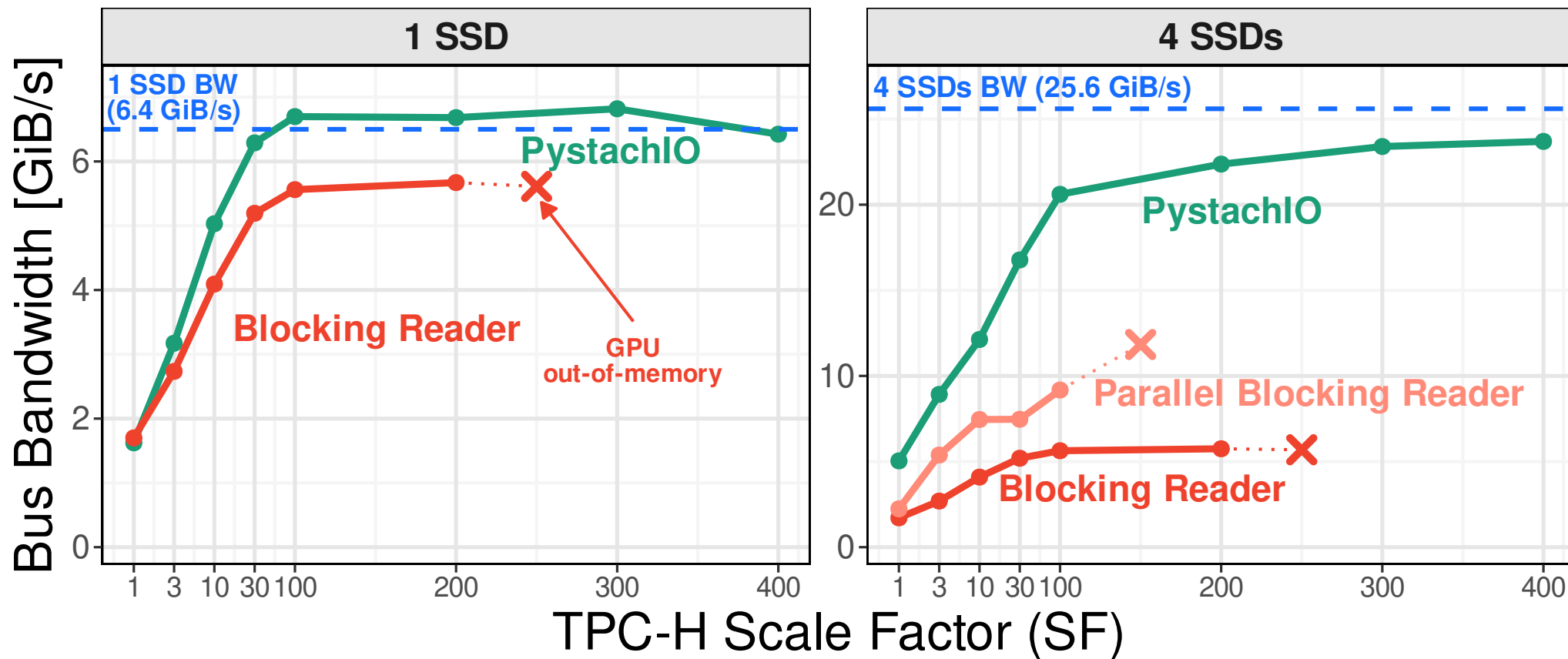
Paper



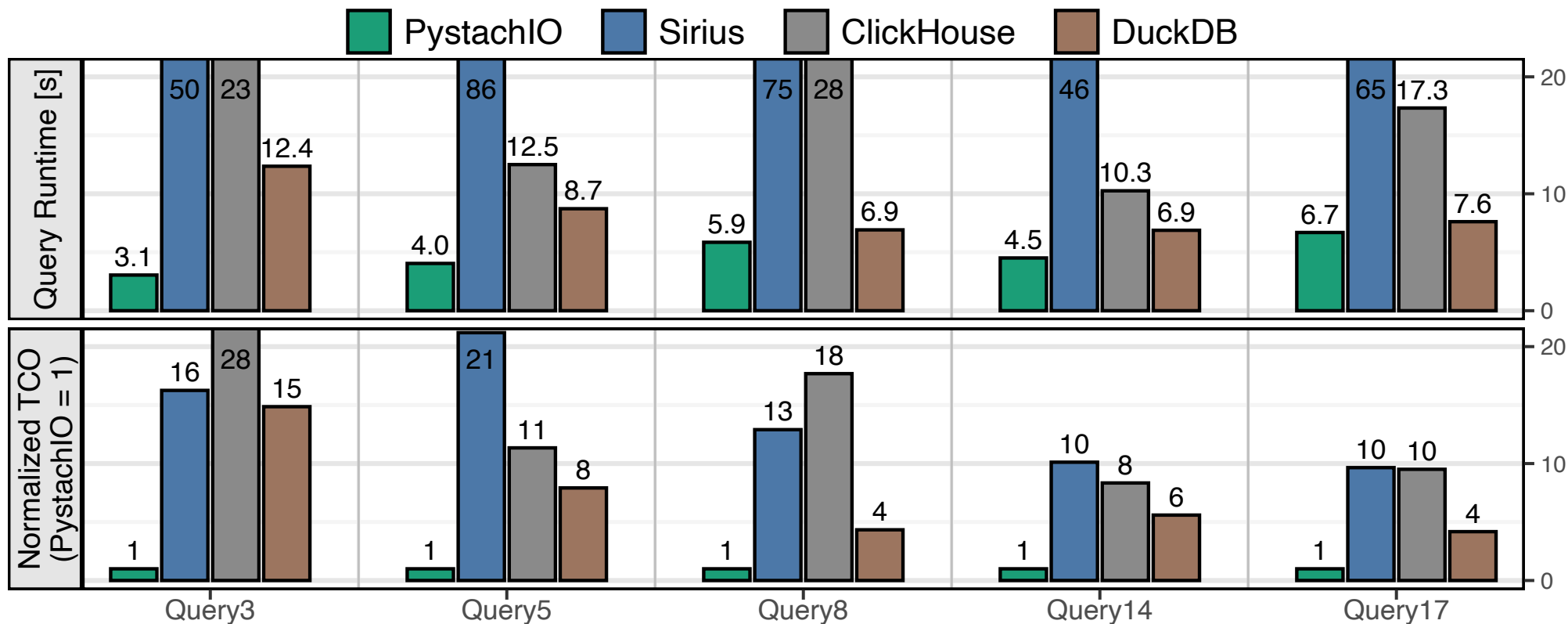
Code

Thanks For Your Attention!

[BACKUP] Storage Scalability: Dataset Size



[BACKUP] System Comparison

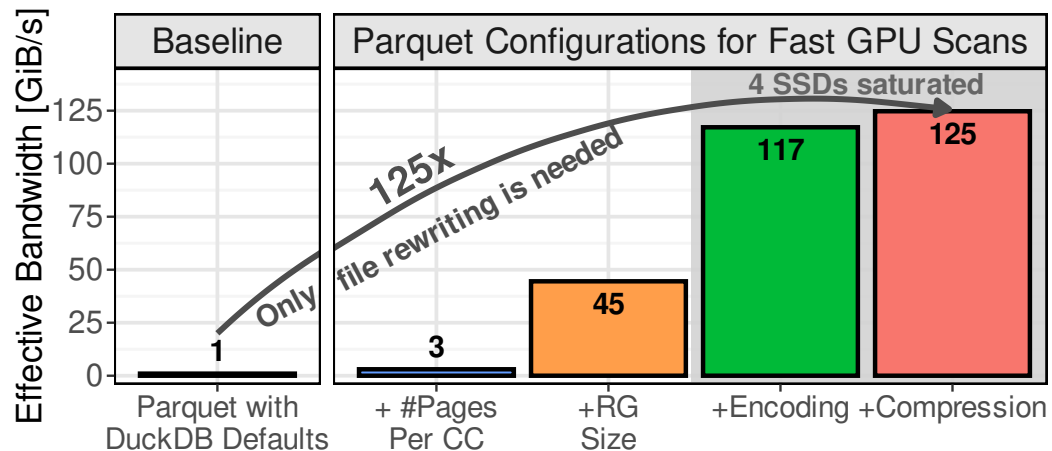


Instance Rental Cost:

NVIDIA A100 40GB GPU: \$1.99/h (Lambda Labs)

AMD EPYC CPU Node: \$7.30/h (AWS r7a.24xlarge, On-Demand)

[BACKUP] DaMoN'26 Conclusion



- GPU-aware Parquet configuration matters: 4 Insights
- Parquet performs well once optimized for GPUs
- GPU-optimized config. $\neq$ CPU-optimized config.
- Understand & optimize Parquet **before** replacing it!

