



Exceptional service in the national interest

TEMPI: An Interposed MPI Library with a Canonical Representation of MPI Datatypes

Carl Pearson¹, Kun Wu², I-Hsin Chung³, Jinjun Xiong³, Wen-Mei Hwu⁴

¹Sandia National Labs / ²University of Illinois Electrical and Computer Engineering / ³IBM T. J. Watson Research / ⁴Nvidia Research

HPDC - Jun 24, 2021

Work completed at University of Illinois prior to joining Sandia National Labs

SAND2021-6846 C

Sandia National Laboratories is a multimission laboratory managed and operated by National Technology and Engineering Solutions of Sandia LLC, a wholly owned subsidiary of Honeywell International Inc. for the U.S. Department of Energy's National Nuclear Security Administration under contract DE-NA0003525.



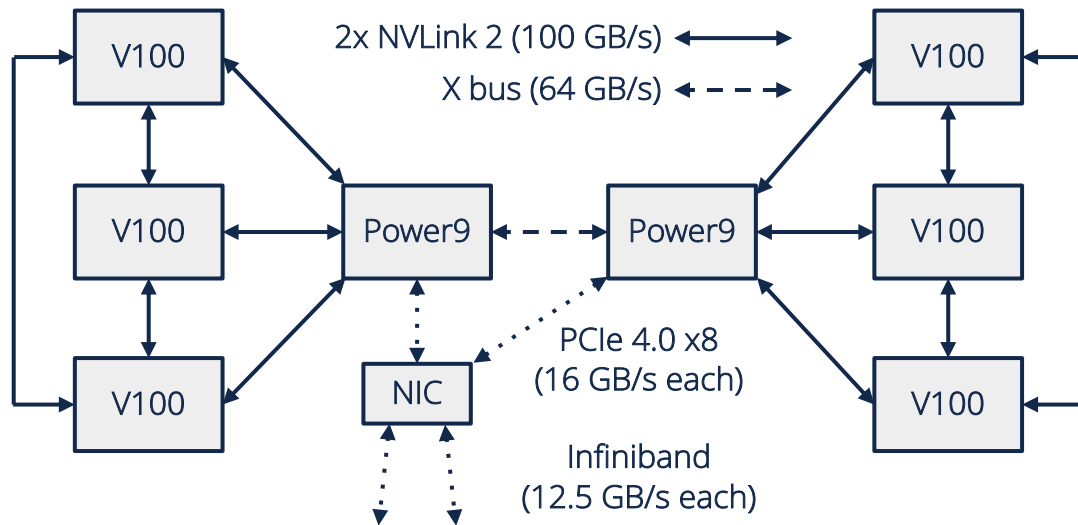


Outline

- Distributed GPU stencil, non-contiguous data
- Equivalence of strided datatypes and minimal representation
- GPU communication methods
- Deploying on managed systems
- Large messages and MPI datatypes
- Translation and canonicalization
- Automatic model-driven transfer method selection
- Interposed library implementation



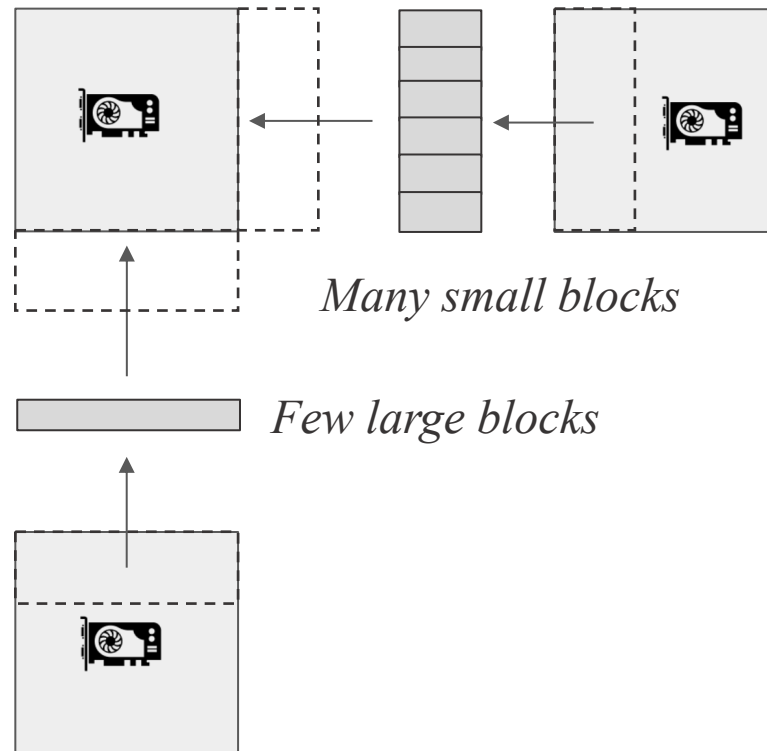
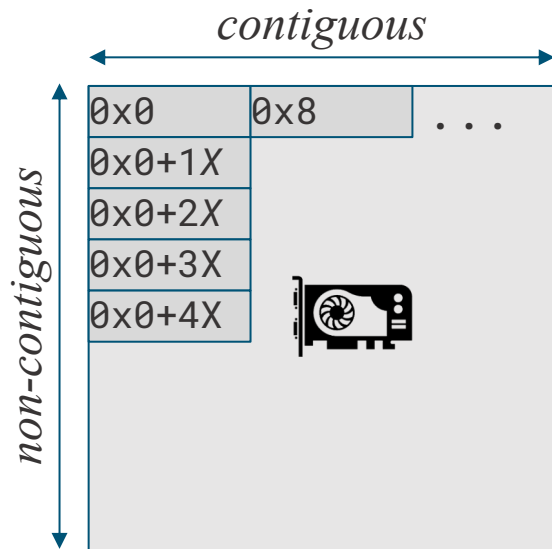
OLCF Summit Node



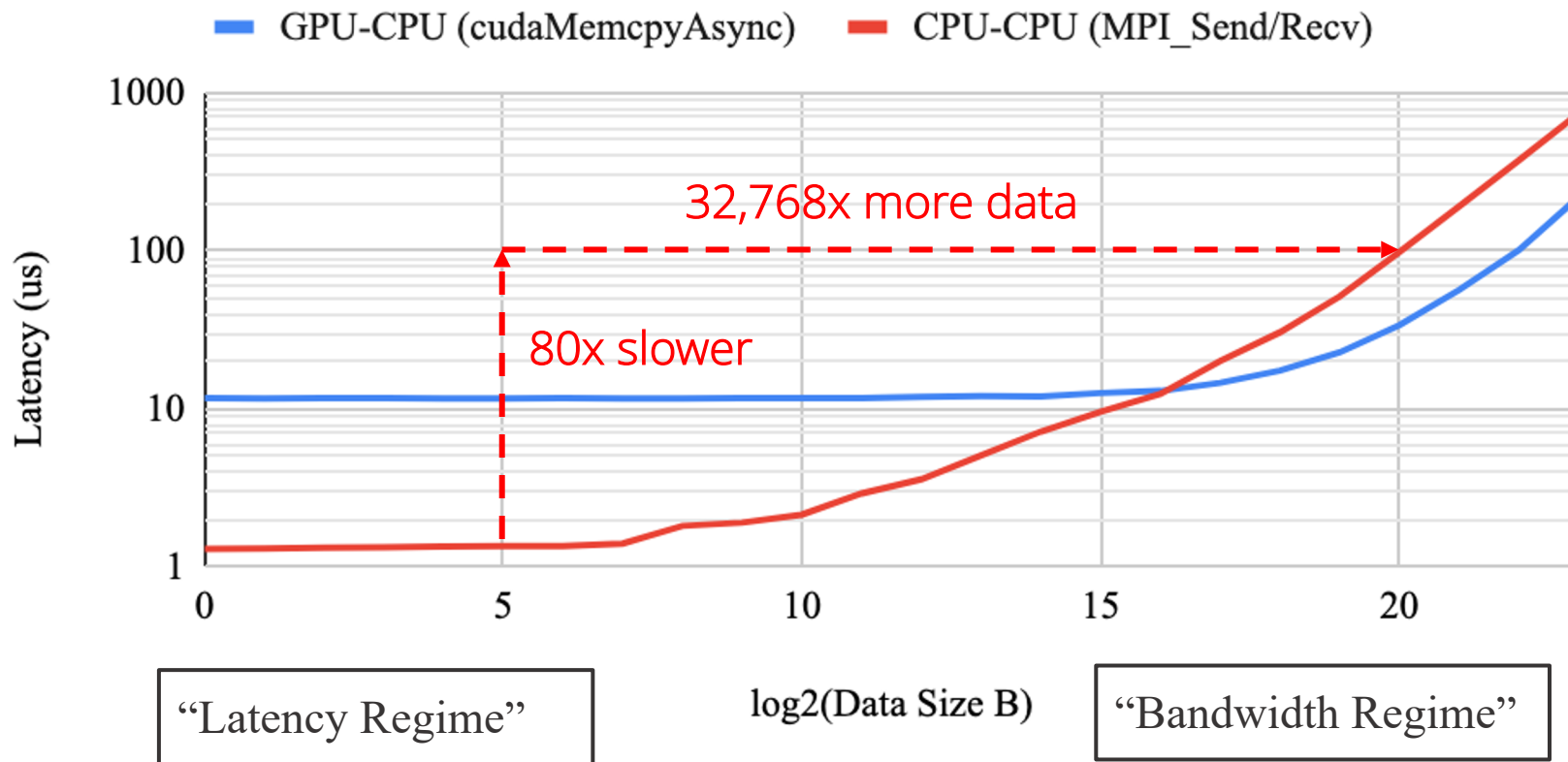
Summit Node
(bidirectional bandwidth)



Stencil Communication and non-contiguous Data

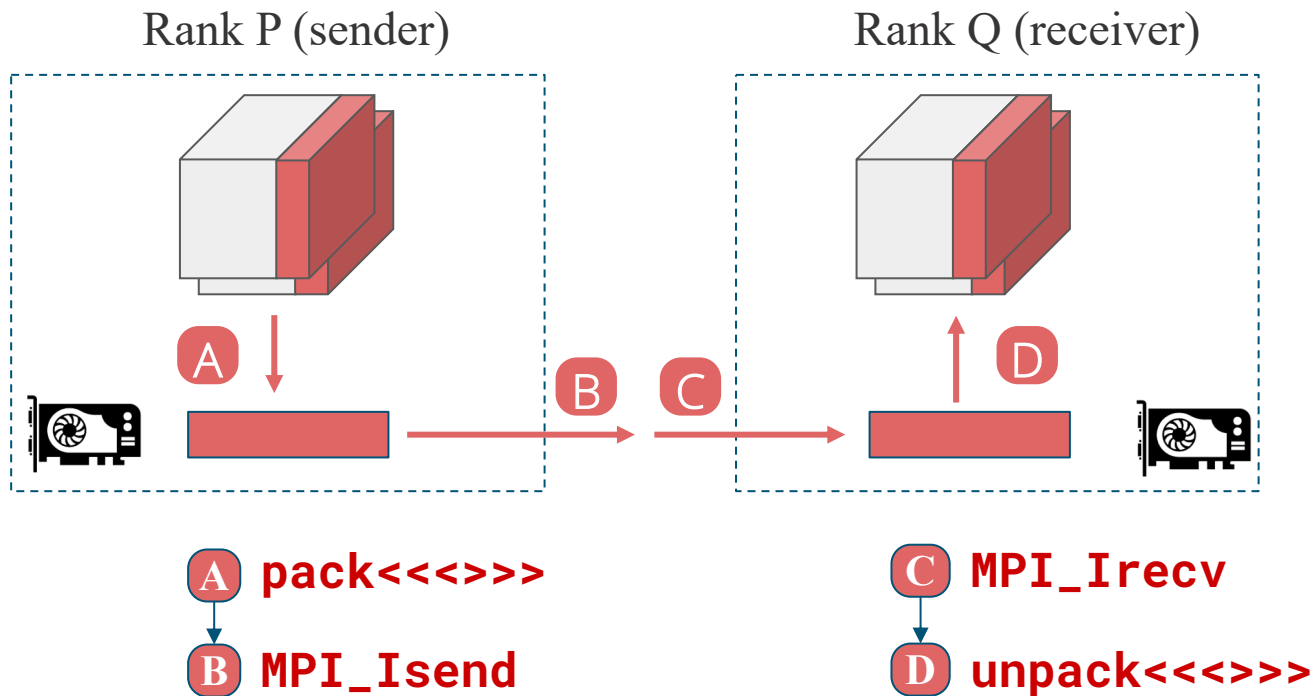


OLCF Summit Latency vs Transfer Size





CUDA-Aware MPI + Packing

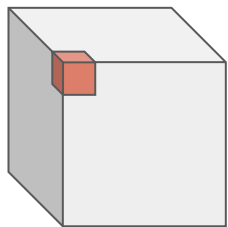


MPI Derived Datatypes

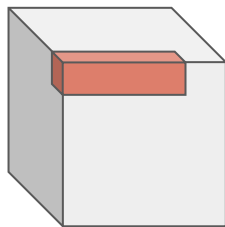
"MPI_Type_contiguous(...)"

"MPI_Type_vector(...)"

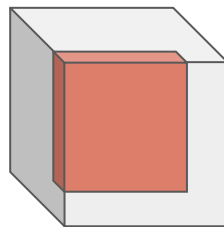
"MPI_Type_vector(...)"



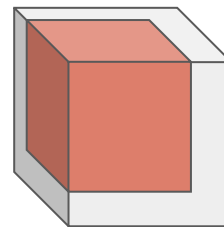
1 MPI_BYTE



row of
contiguous
bytes

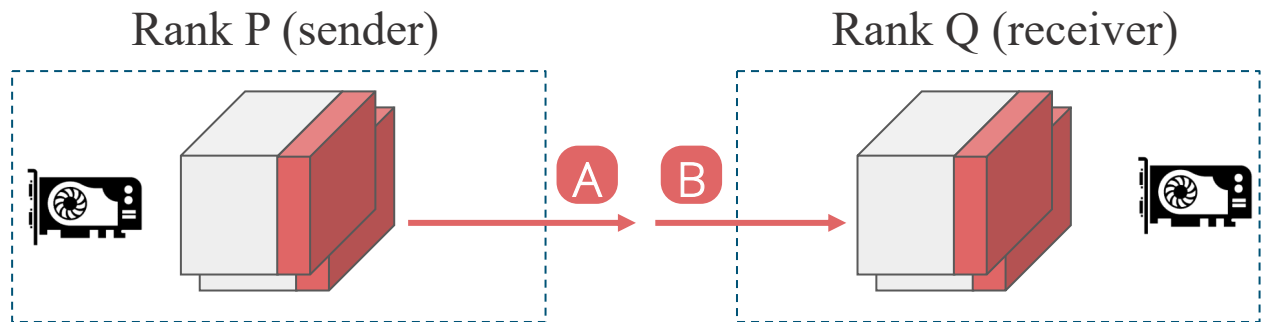


plane of non-
contiguous
rows



cuboid of
non-contiguous
planes

CUDA-Aware MPI + MPI Derived Datatypes



Setup (once):

MPI_Type...

MPI_Type...

...

MPI_Type_commit

MPI_Type...

MPI_Type...

...

MPI_Type_commit

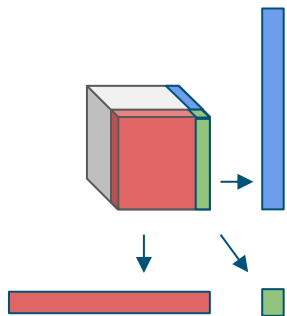
Each halo exchange:

A MPI_Isend

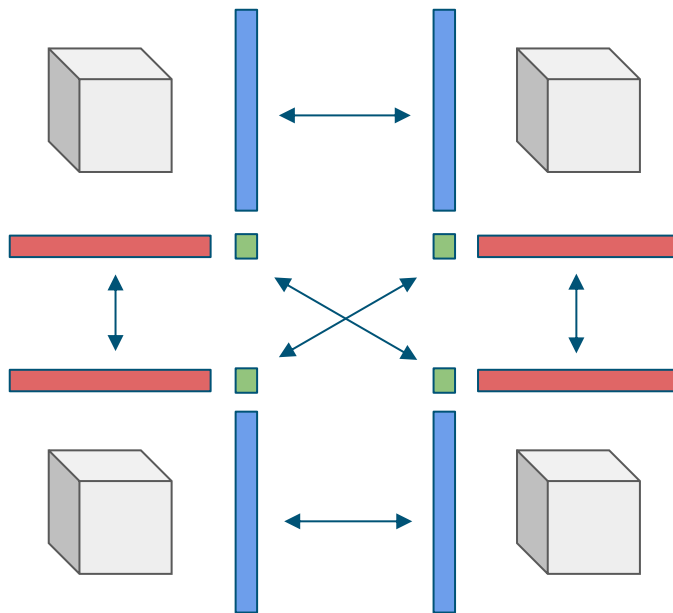
B MPI_Irecv



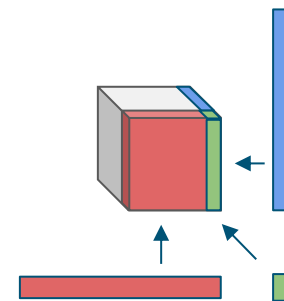
“alltoallv” Halo Exchange



MPI_Packs

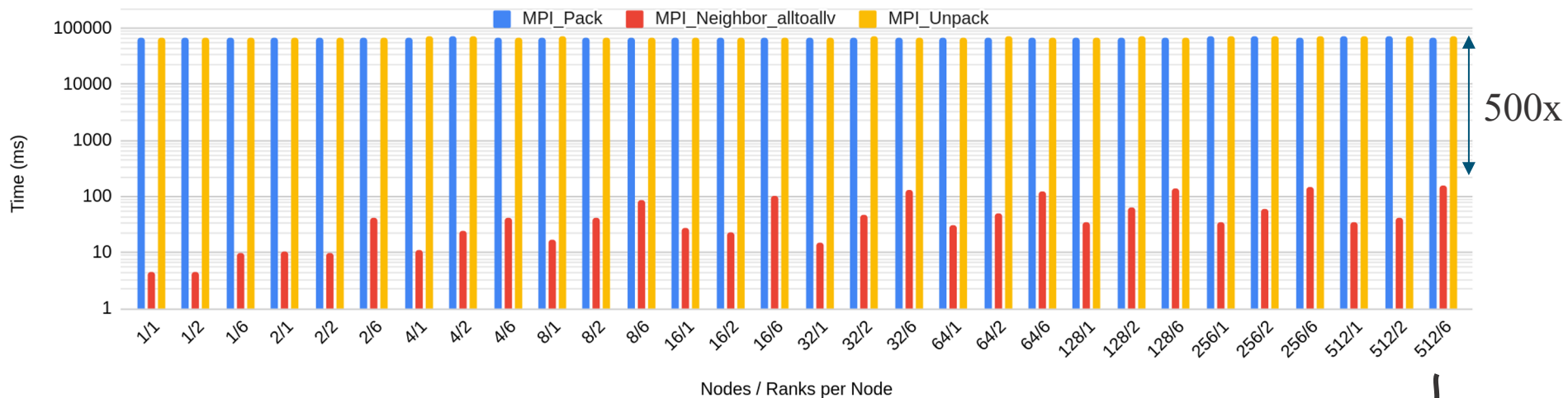


MPI_Neighbor_alltoallv



MPI_Unpacks

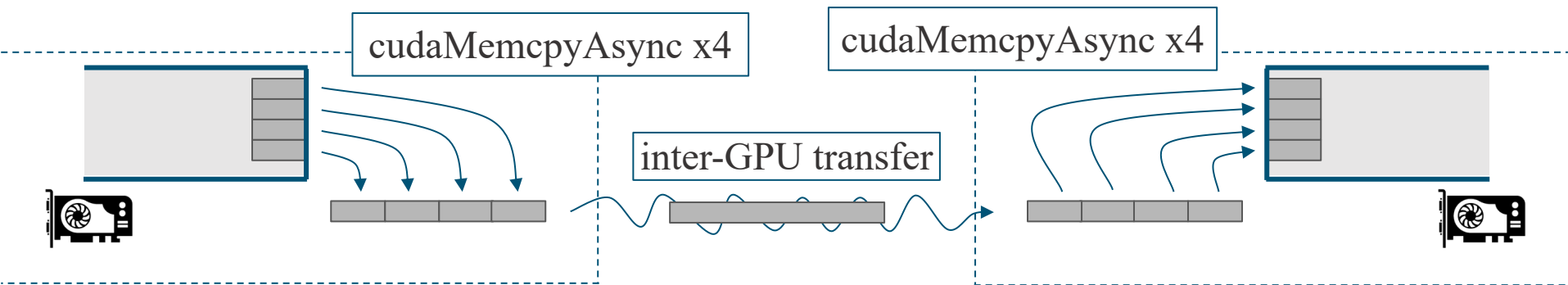
“alltoallv” with MPI derived types



- $\text{MPI_Neighbor_alltoallv} = \sim 500 \text{ MB/s/rank}$
- $\text{MPI_Pack} / \text{MPI_Unpack} = \sim 1 \text{ MB/s}$



MPI_Send (Spectrum MPI)



“Latency Regime”



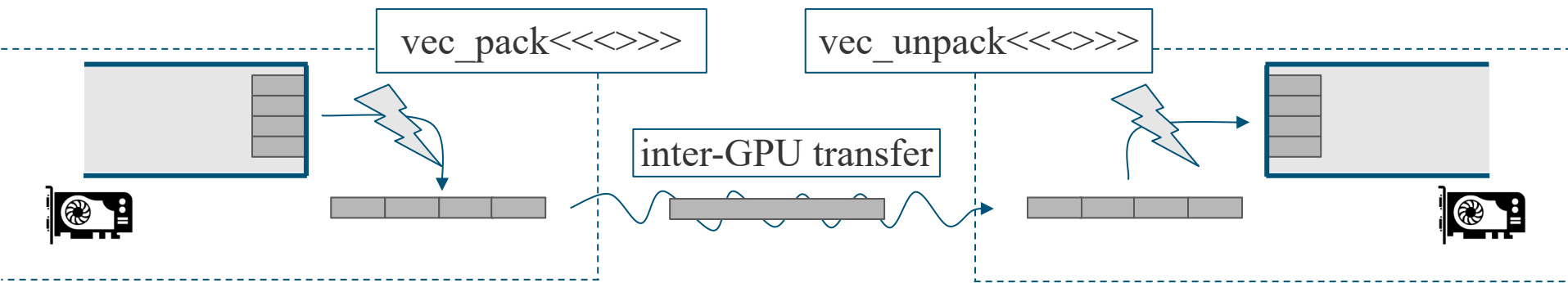
“Bandwidth Regime”



“Latency Regime”

Better MPI_Send (common foundation)

- GPU kernels to pack non-contiguous data



“Bandwidth Regime”



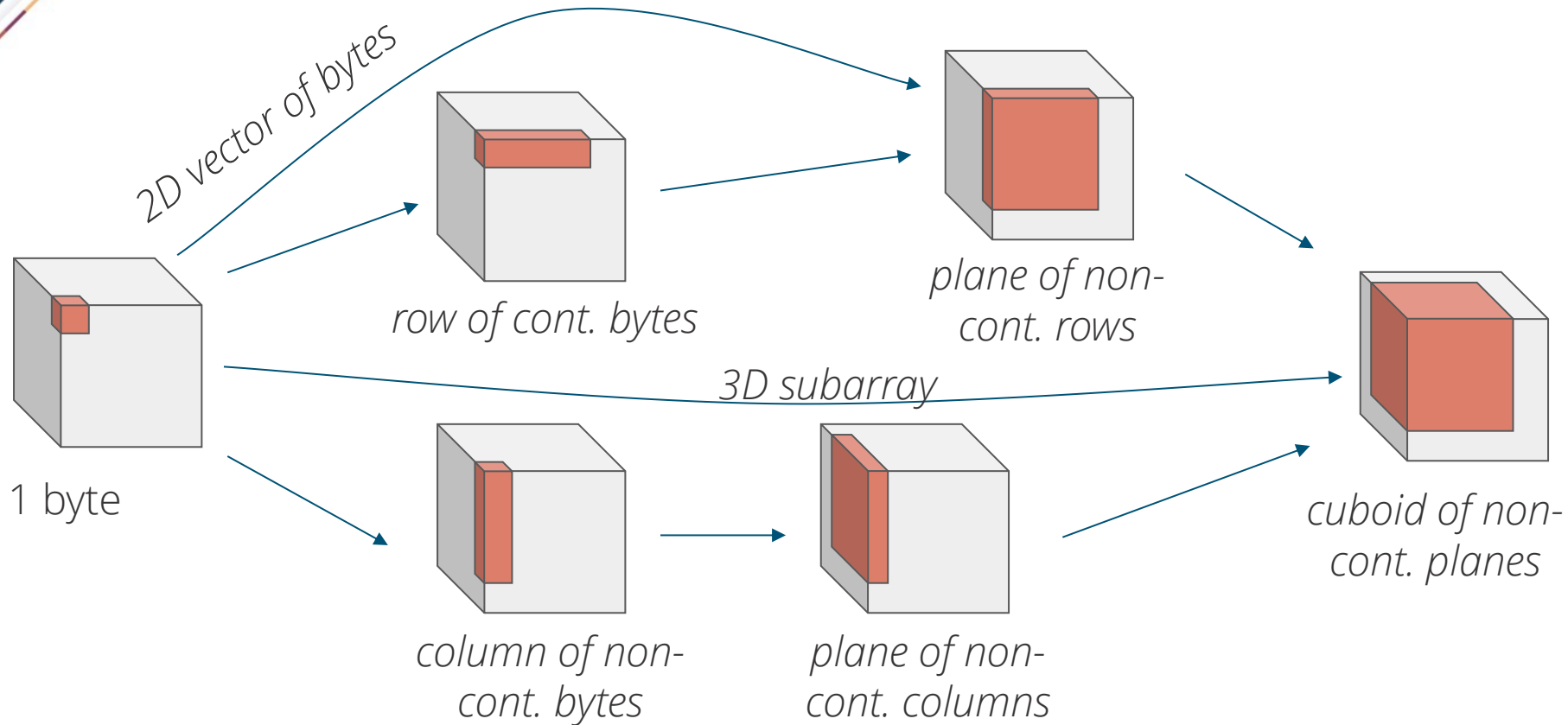
“Bandwidth Regime”



“Bandwidth Regime”



MPI Derived Datatype Equivalence





MPI_Type_commit()



Translation

Convert MPI Derived Datatype into internal representation (IR)



Canonicalization

Convert semantically-equivalent IR to simplified form



Kernel Selection

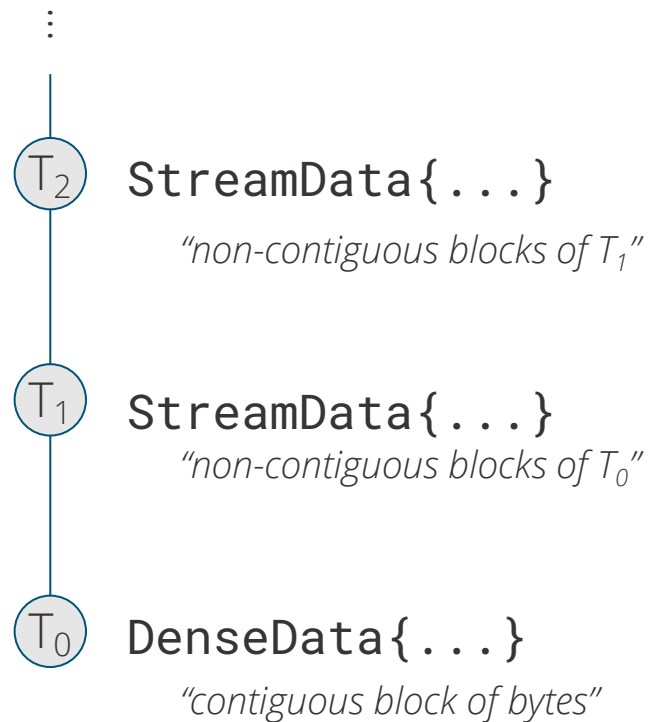
Choose packing/unpacking kernel for future operations

IR - “Internal Representation”

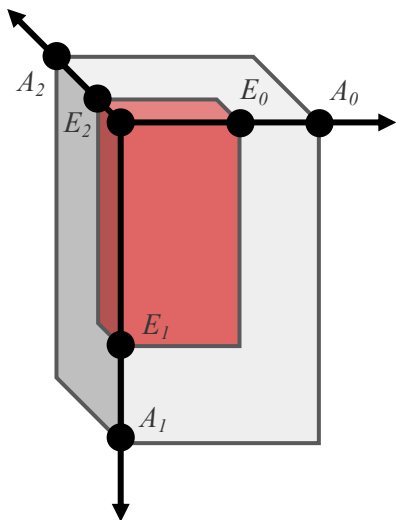
```
StreamData {  
    integer offset; // offset (B) of the first element  
    integer stride; // pitch (B) between element  
    integer count;  // number of elements  
}
```

```
DenseData {  
    integer offset; // offset (B) of the first byte  
    integer extent; // number of bytes  
}
```

Hierarchy of StreamData, rooted at DenseData



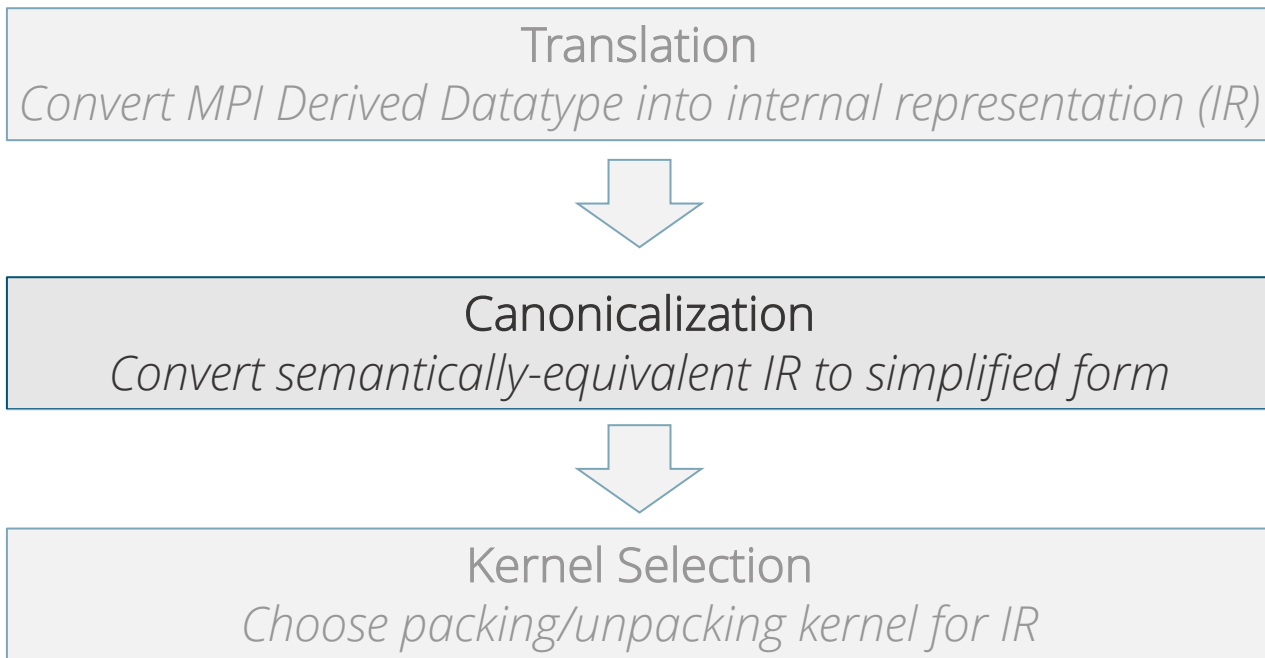
Example



Type	cuboid	StreamData{offset:0, count: E_2 , stride: $A_1 * A_0$ }
		StreamData{offset:0, count:1, stride: $E_1 * A_0$ }
	plane	StreamData{offset:0, count: E_1 , stride: A_0 }
		StreamData{offset:0, count:1, stride: E_0 }
	row	StreamData{offset:0, count: E_0 , stride:1}
		StreamData{offset:0, count:1, stride:1}
	MPI_BYTE	DenseData{offset:0, extent: 1}

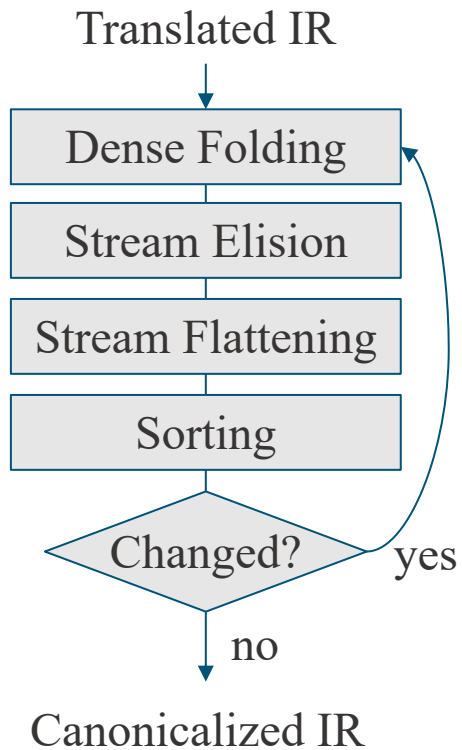


TEMPI Datatype Handling



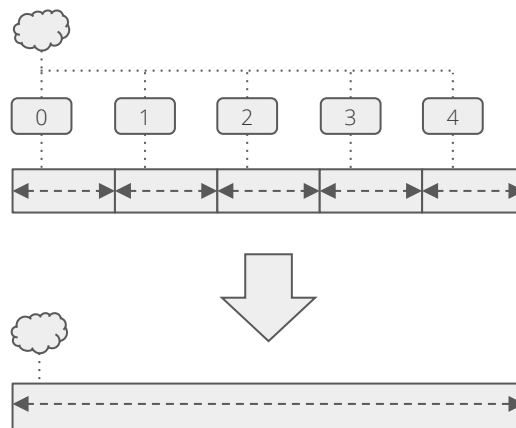
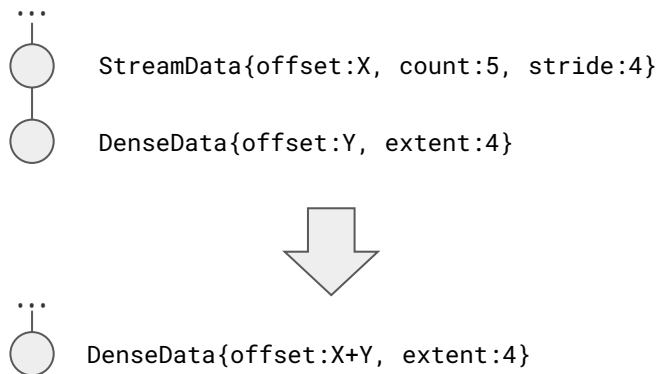


Canonicalization



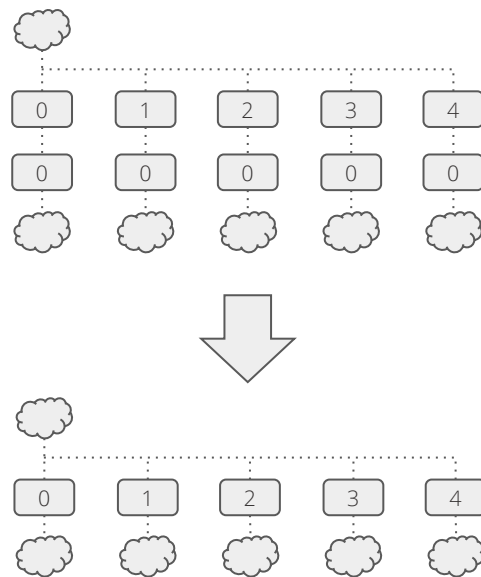
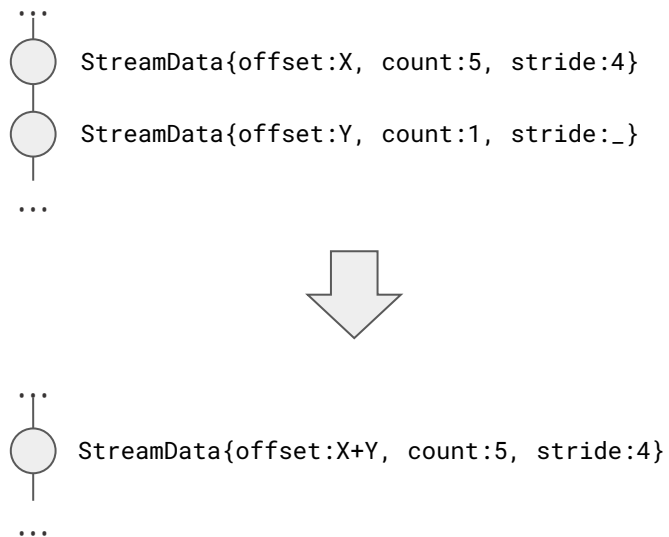
Canonicalization: Dense Folding

- `StreamData` is contiguous `DenseData` (parent of MPI named type)



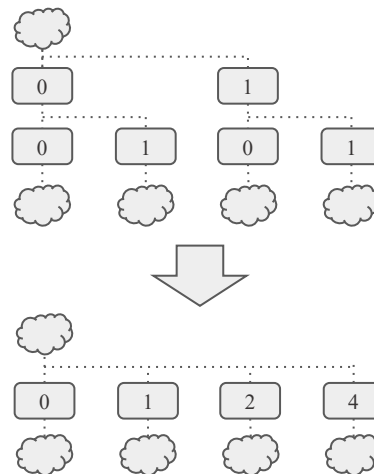
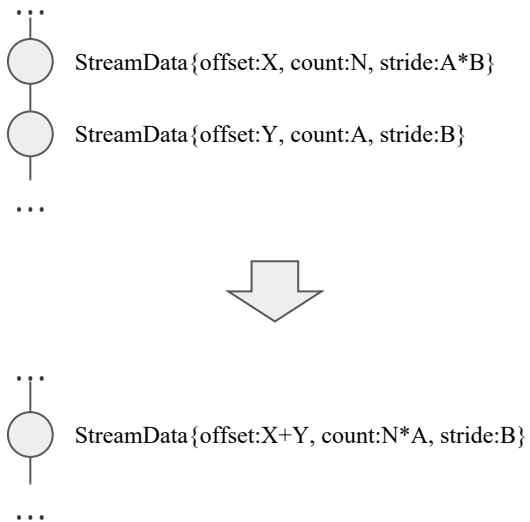
Canonicalization: Stream Elision

- Remove `StreamData` with `count = 1` (MPI Vector blocks commonly have one element)



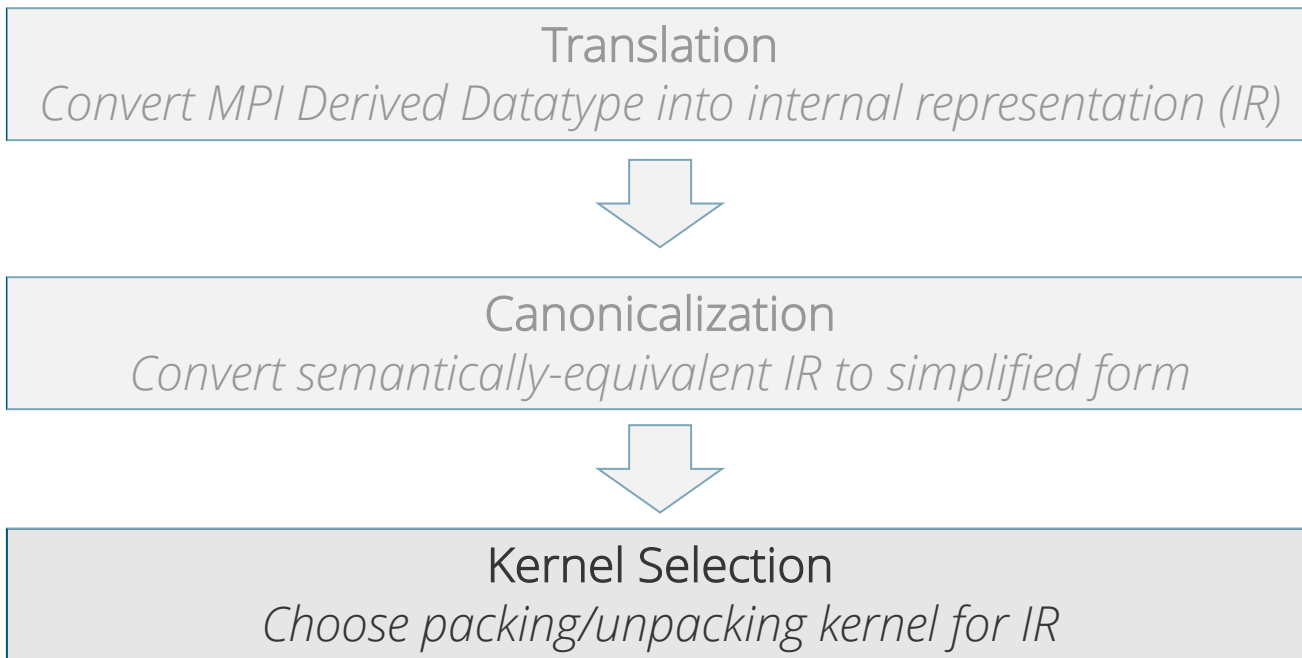
Canonicalization: Stream Flattening

- e.g. two vectors of three vs. one vector of six

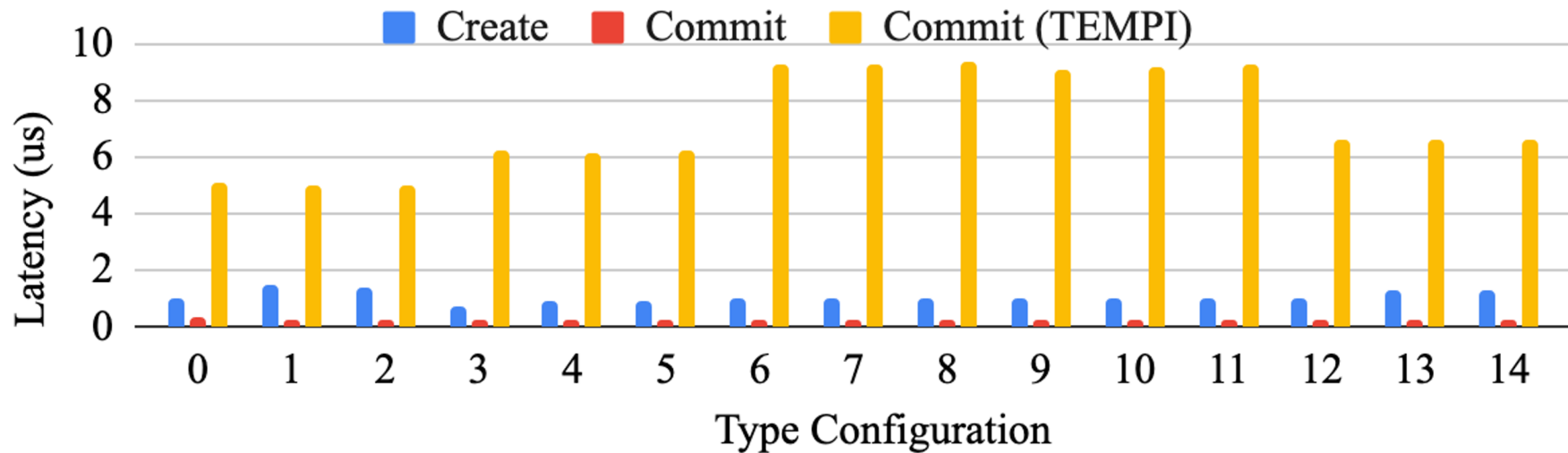




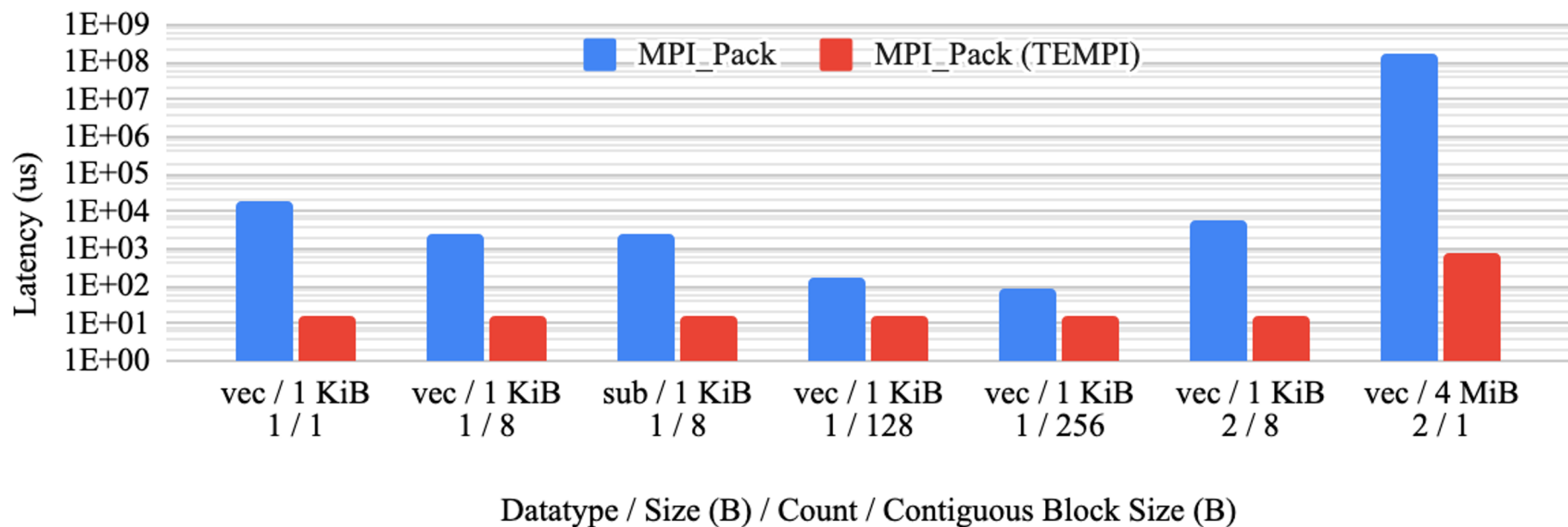
TEMPI Datatype Handling



MPI_Type_commit Time

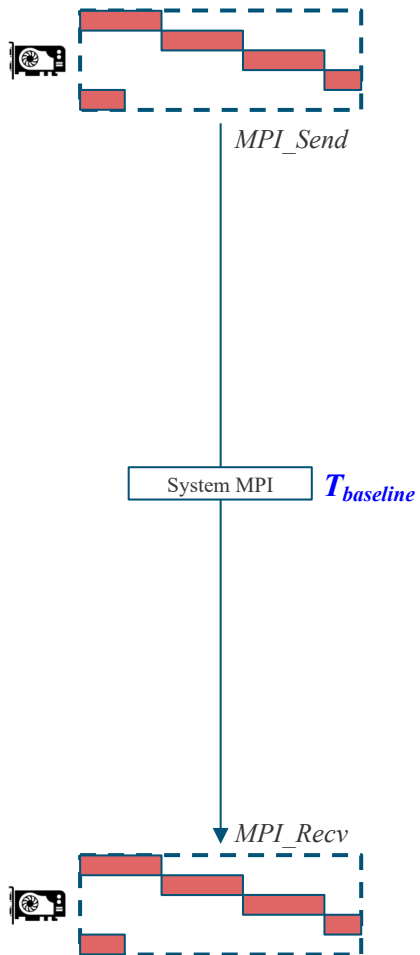


MPI_Pack Results



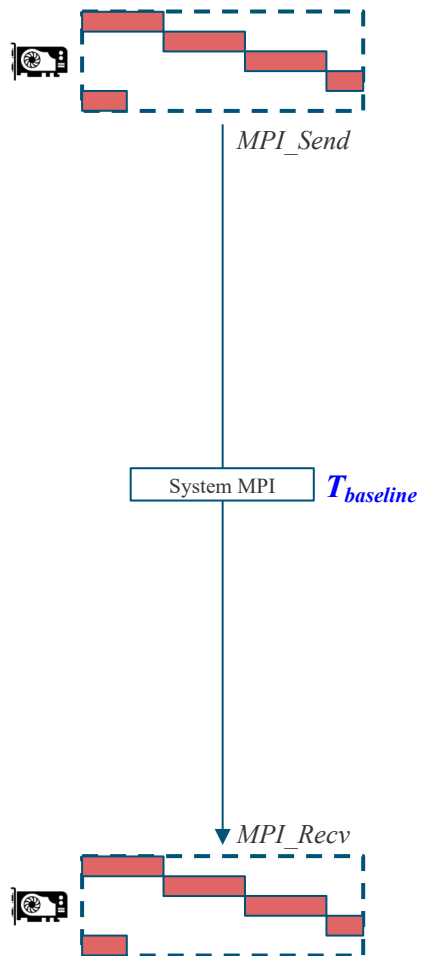


CUDA-Aware System MPI

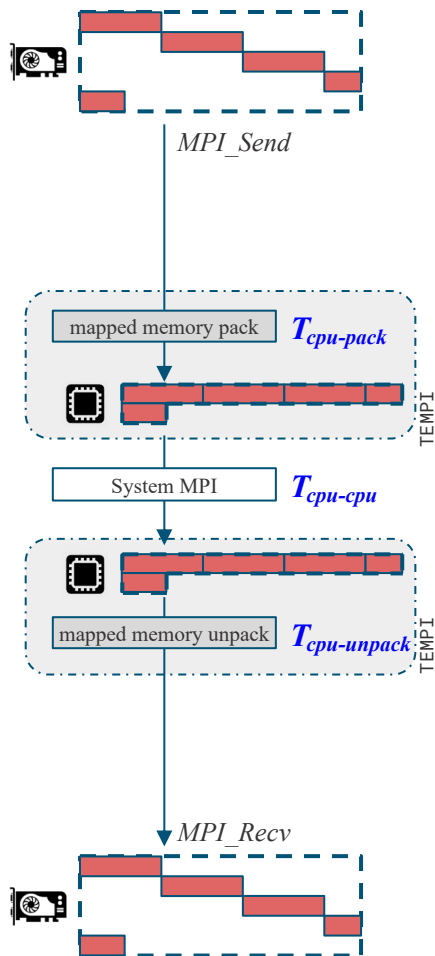




CUDA-Aware System MPI

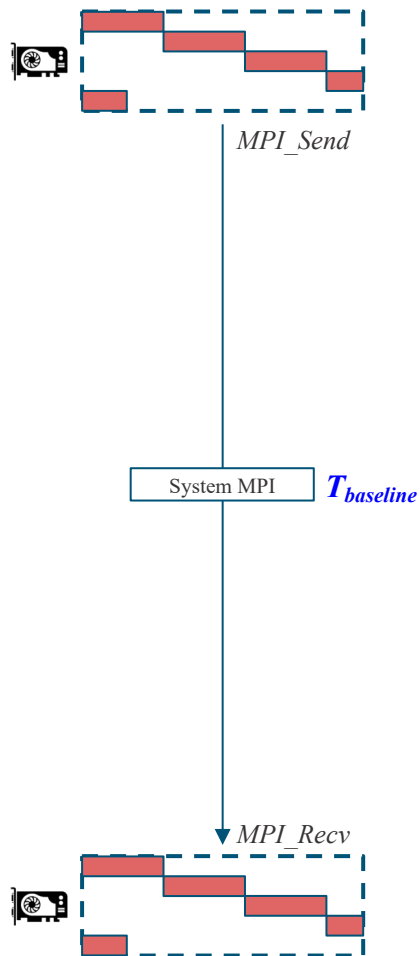


TEMPI "One-shot"

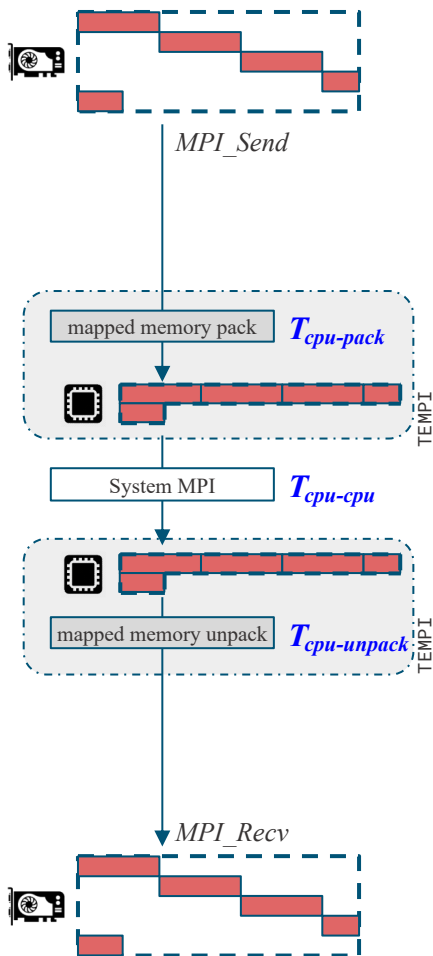




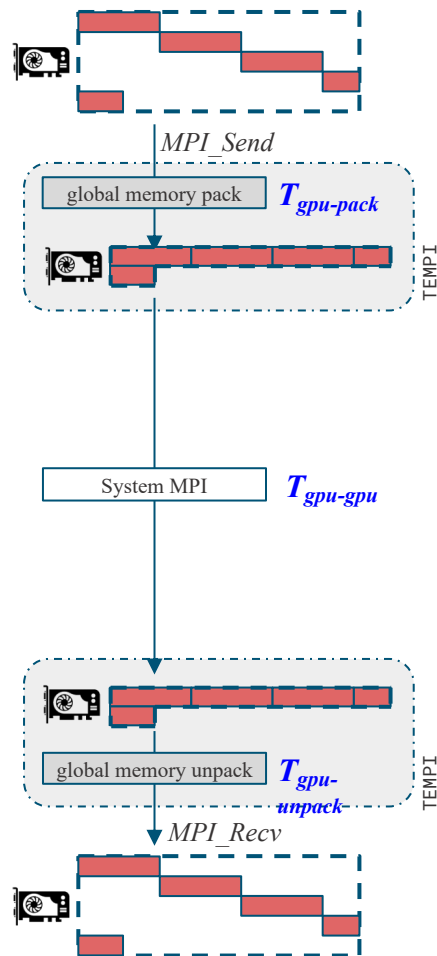
CUDA-Aware System MPI



TEMPI "One-shot"



TEMPI "Device"

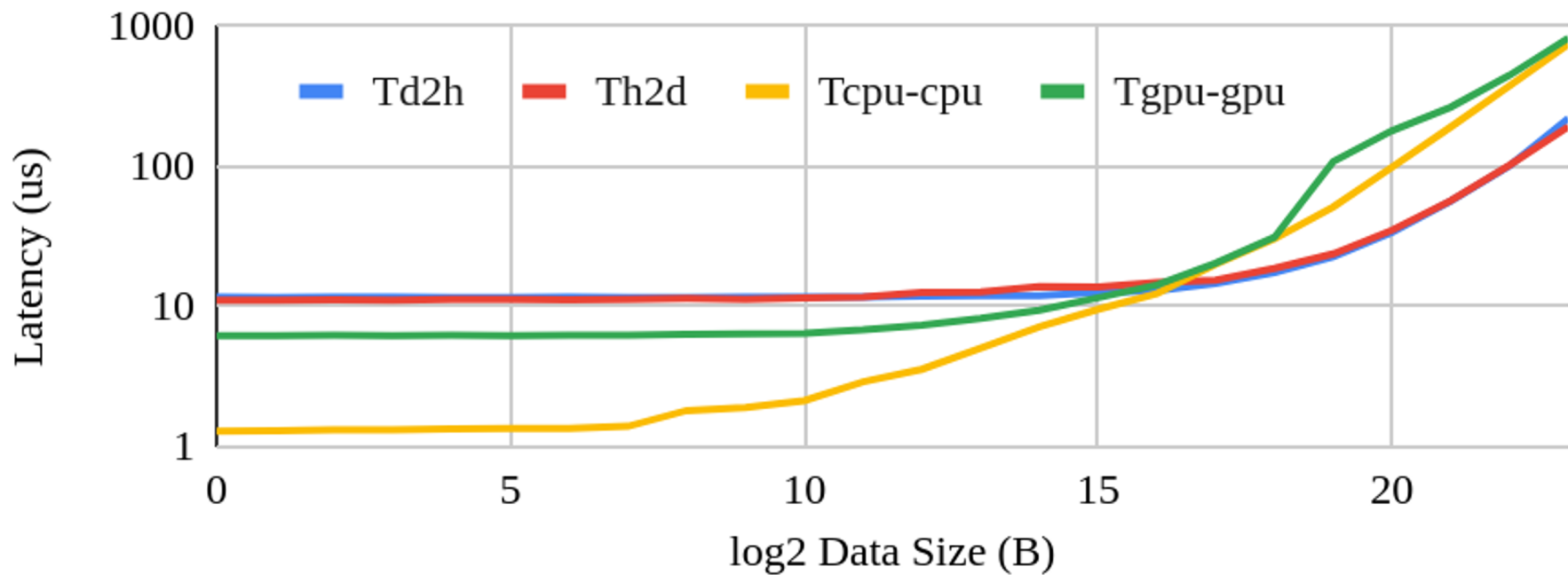


$$T_{device} = T_{gpu-pack} + T_{gpu-gpu} + T_{gpu-unpack}$$

$$T_{oneshot} = T_{host-pack} + T_{cpu-cpu} + T_{host-unpack}$$

$$T_{staged} = T_{gpu-pack} + T_{d2h} + T_{cpu-cpu} + T_{h2d} + T_{gpu-unpack}$$

Contiguous Transfer Measurements



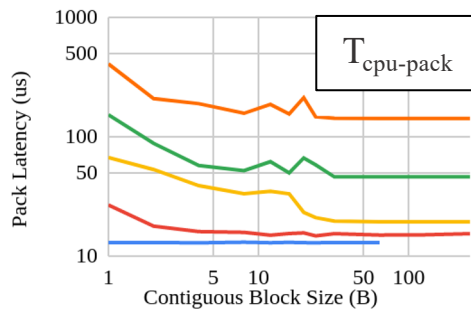
T depends on # bytes transferred



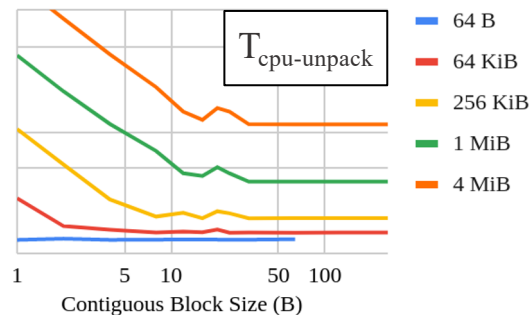
One-shot and Device Pack and Unpack

ONE-SHOT

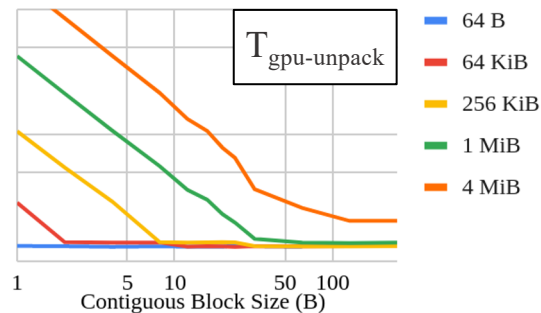
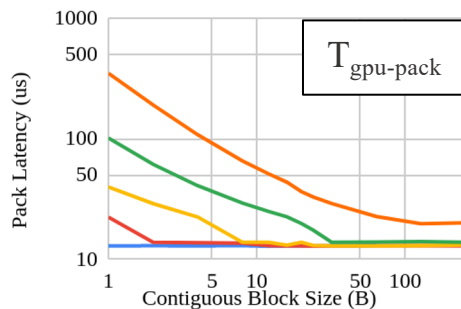
PACK



UNPACK

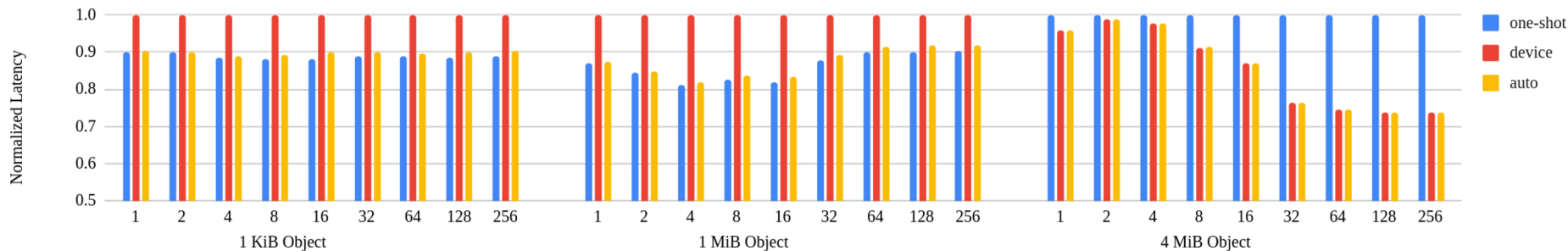


DEVICE



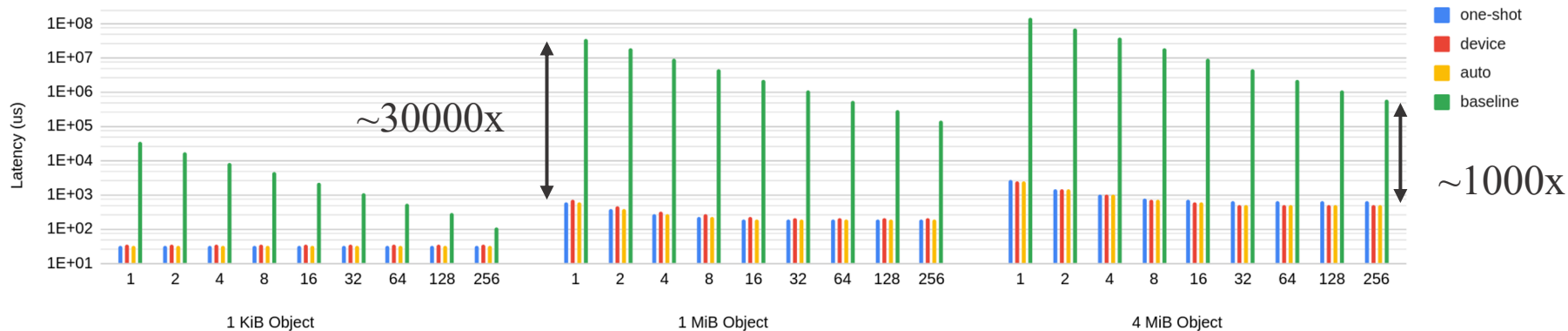
T depends on object size and block size

MPI_Send Microbenchmark



Minimal runtime performance-modeling overhead
Performance model reliably chooses faster method

MPI_Send/Recv Latency for 2D objects with different block sizes



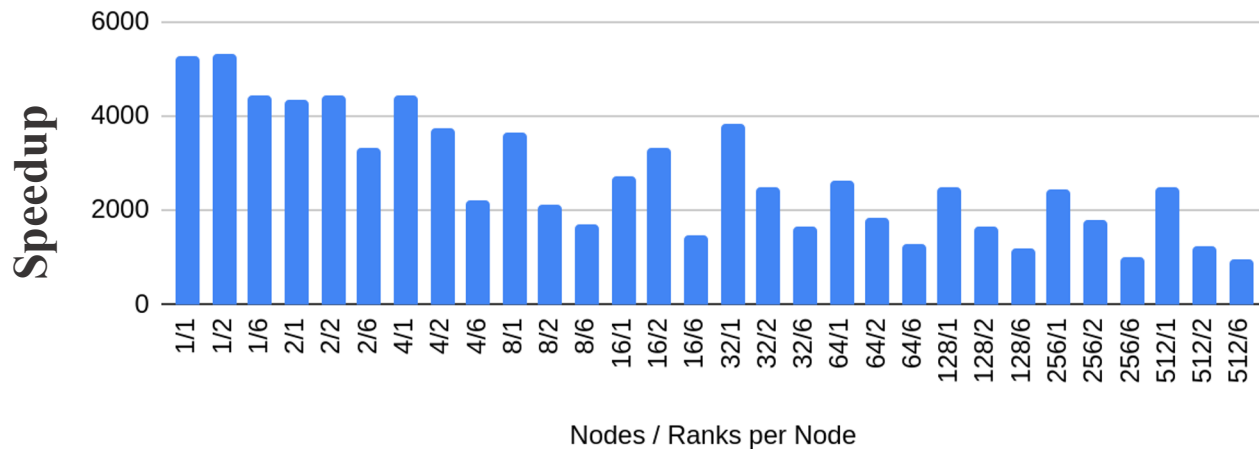
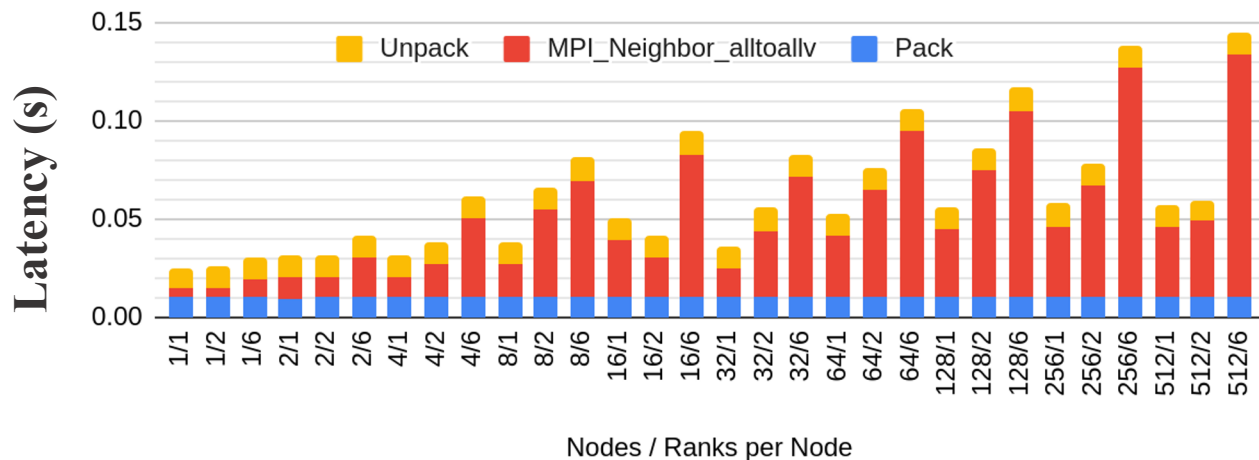
Large latency improvement from datatype handling
Small additional improvement from automatic method selection



Halo Exchange

With TEMPI, non-contiguous packing no longer dominates

Large speedup in 3D halo exchange



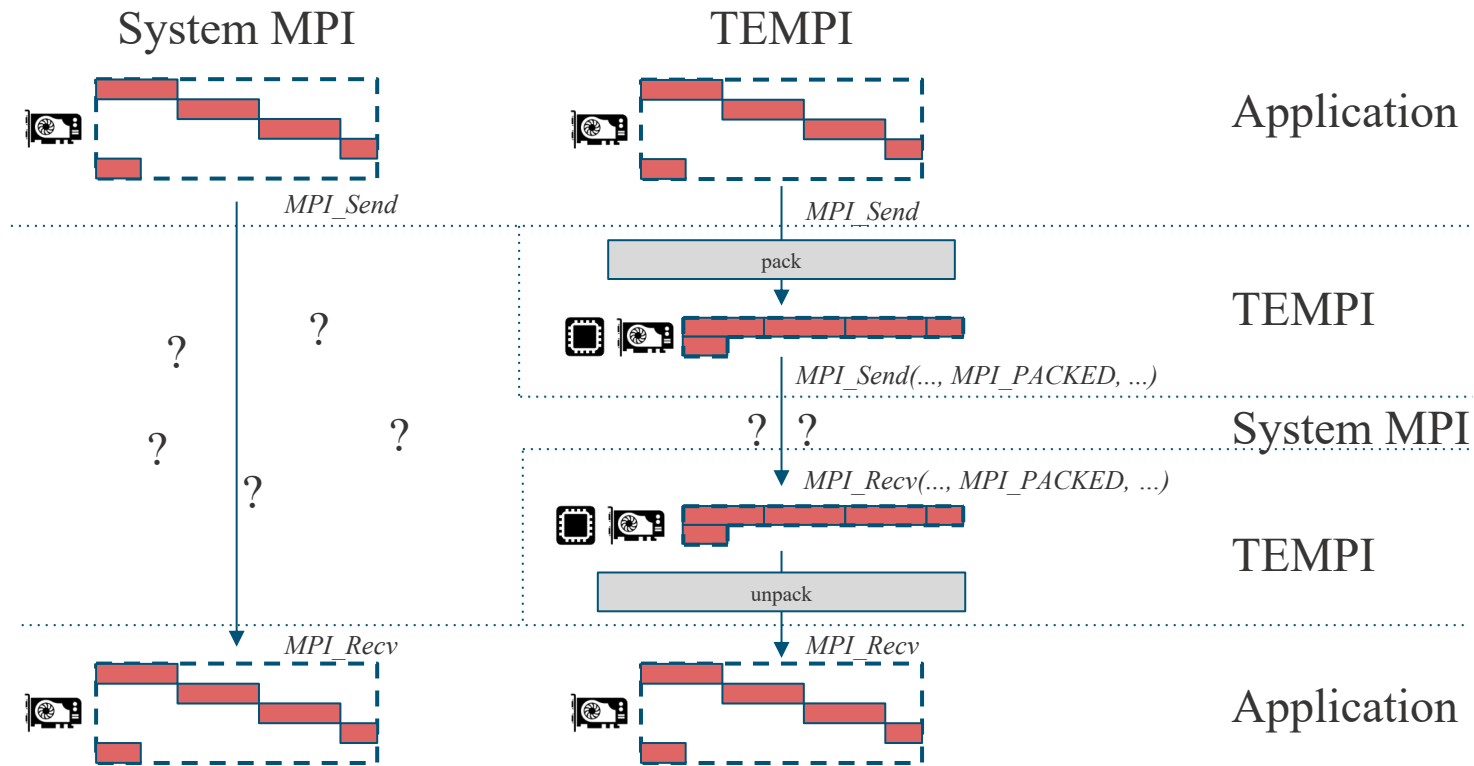


A Practical Challenge: Deploying MPI experiments on a Supercomputer

- Large-scale systems are tightly controlled
 - Can't just make whatever changes you want
 - Usually one MPI (or maybe two) are deployed on the system
 - Rarely bugfixed (if ever)
 - Even more rarely are new features added
 - Difficult or impossible to make experimental modifications
-
- MPI has a well-defined standard
 - Take advantage of this + how OS loads libraries to inject modifications



TEMPI's Architecture: Interposer Library



app.c

```
#include <mpi.h>

int main(int argc, char **argv) {
    MPI_Init(&argc, &argv);
}
```

1



```
gcc app.c -o app \
-I /mpi/include \
-L /mpi/lib \
-l mpi
```

2



app

```
...
callq 7260
<MPI_Init@plt>
...
```

4

8



system MPI

mpi.h

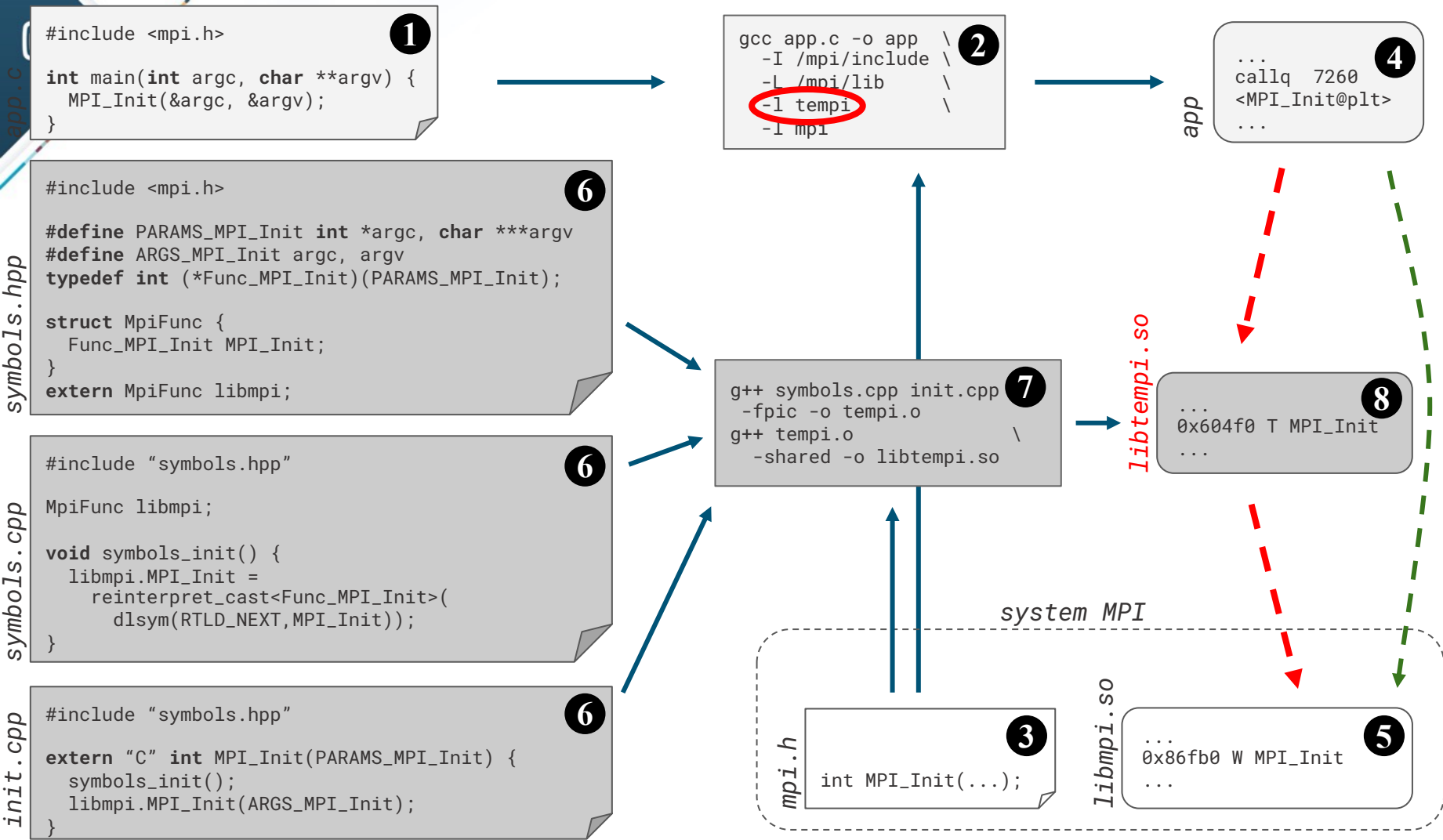
```
int MPI_Init(...);
```

3

libmpi.so

```
...
0x86fb0 W MPI_Init
...
```

5





Conclusions & Future Work

- Canonicalization approach works
 - 👍 Speedup for unmodified applications on OLCF summit
 - 👍 Any strided datatype
- Simple performance model to select GPU data transfer method
 - 👍 speedup
 - 👎 < 2x speedup
- Interposed library approach nice for experiments & prototype deployment
 - 👍 Easy to use without system privileges
 - 👎 Limited integration with existing MPI
- OpenMPI, MPICH, MVAPICH have other strategies
 - Use if available - covers some or all of the same problems
 - Comparison of MPI datatype performance would be useful to community
- github.com/cwpearson/tempi



Acknowledgements

This work is supported by IBM-ILLINOIS Center for Cognitive Computing Systems Research (C3SR) - a research collaboration as part of the IBM AI Horizon Network.

This research used resources of the Oak Ridge Leadership Computing Facility at the Oak Ridge National Laboratory, which is supported by the Office of Science of the U.S. Department of Energy Contract No. DE-AC05-00OR22725.

This work utilizes resources supported by the National Science Foundation's Major Research Instrumentation program, grant #1725729, as well as the University of Illinois at Urbana-Champaign



Thank you

cwpears@sandia.gov

github.com/cwpearson/tempi

This work was completed prior to joining Sandia National Labs.