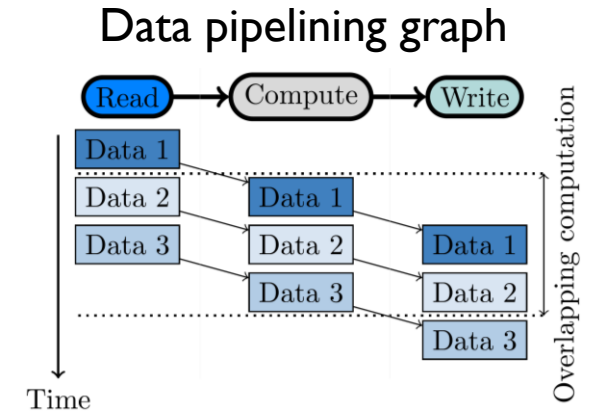
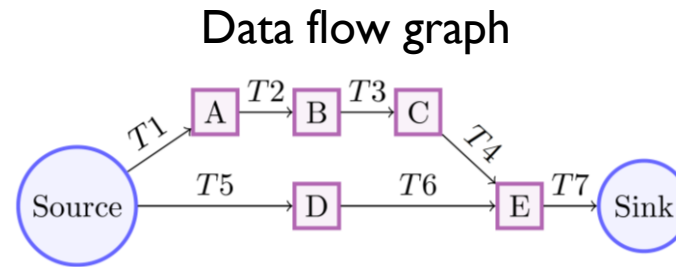


HiHAT - Q&A

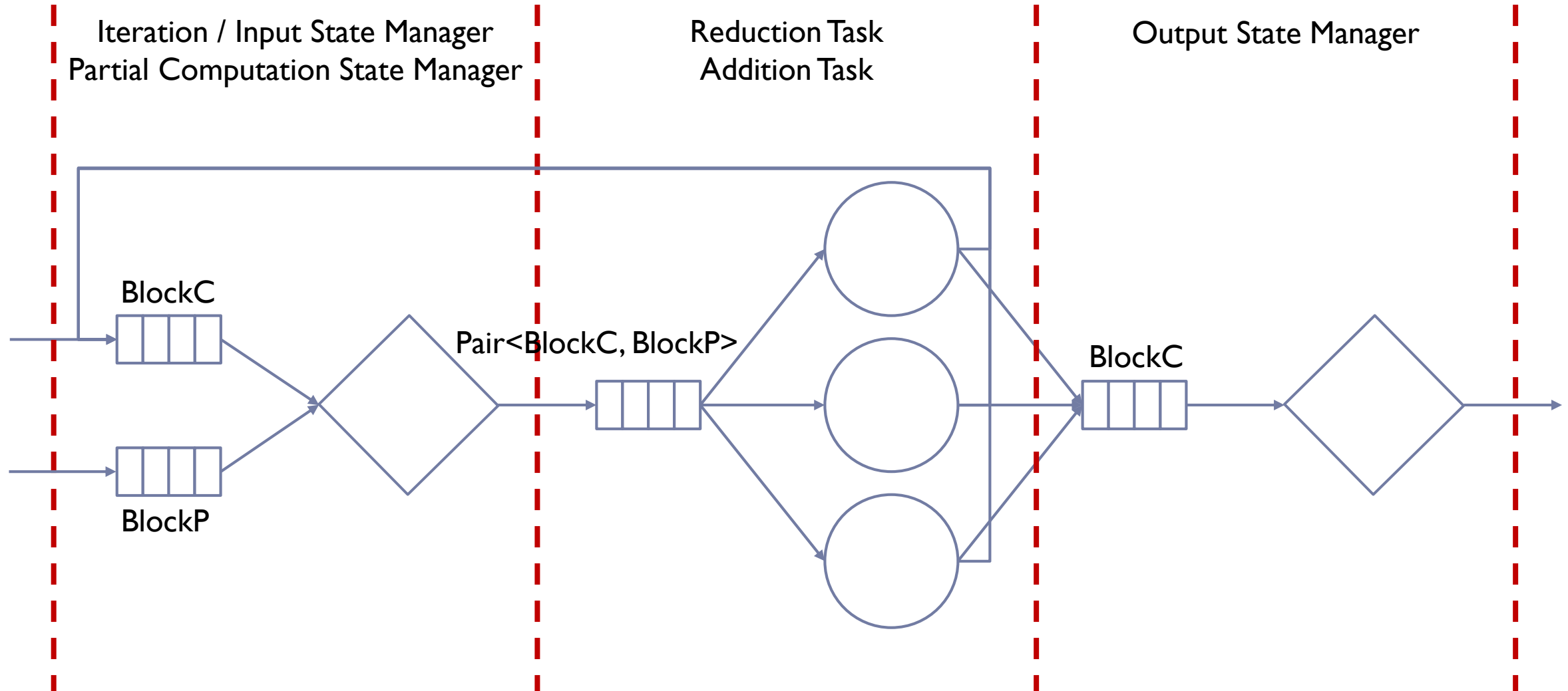
Alexandre Bardakoff, Timothy Blattner, Walid Keyrouz
NIST | ITL | SSD | ISG

Recaps

- ▶ Data Flow
- ▶ Data pipelining
- ▶ Persistent kernel: tasks lives forever on a runtime loop until the graph is done
- ▶ One input is enough to fire a task
- ▶ Coarse grain parallelism



Example Parallel Reduction w/ matrix multiplication



Example Parallel Reduction w/ matrix multiplication

Partial Computation State Manager

```
template<class Type, Order Ord = Order::Row>
class PartialComputationStateManager
    : public hh::StateManager<
        std::pair<std::shared_ptr<MatrixBlockData<Type, 'c', Ord>>, std::shared_ptr<MatrixBlockData<Type, 'p',
Ord>>>>,
        MatrixBlockData<Type, 'c', Ord>, MatrixBlockData<Type, 'p', Ord>
    > {
public:
    explicit PartialComputationStateManager(std::shared_ptr<PartialComputationState<Type, Ord>> const &state) :
        hh::StateManager<std::pair<std::shared_ptr<MatrixBlockData<Type, 'c', Ord>>,
            std::shared_ptr<MatrixBlockData<Type, 'p', Ord>>>>,
            MatrixBlockData<Type, 'c', Ord>,
            MatrixBlockData<Type, 'p', Ord>>>("Partial Computation State Manager", state, false) {}

    bool canTerminate() override {
        this->state()->lock();
        auto ret = std::dynamic_pointer_cast<PartialComputationState<Type, Ord>>(this->state())->isDone();
        this->state()->unlock();
        return ret;
    }
};
```

Example Parallel Reduction w/ matrix multiplication

Partial Computation State Manager

```
template<class Type, Order Ord = Order::Row>
class PartialComputationState
: public hh::AbstractState<std::pair<std::shared_ptr<MatrixBlockData<Type, 'c', Ord>>, std::shared_ptr<MatrixBlockData<Type, 'p', Ord>>>,
  MatrixBlockData<Type, 'c', Ord>, MatrixBlockData<Type, 'p', Ord>> {

void execute(std::shared_ptr<MatrixBlockData<Type, 'c', Ord>> ptr) override {
    auto i = ptr->rowIdx(), j = ptr->colIdx();
    if (isPAvailable(i, j)) {
        auto res = std::make_shared<
            std::pair<std::shared_ptr<MatrixBlockData<Type, 'c', Ord>>, std::shared_ptr<MatrixBlockData<Type, 'p', Ord>>>
        >();
        res->first = ptr;
        res->second = partialProduct(i, j);
        this->push(res);
        --ttl_;
    } else {
        blockMatrixC(ptr);
    }
}

void execute(std::shared_ptr<MatrixBlockData<Type, 'p', Ord>> ptr) override { //Same Algorithm for P with C & P switch}

bool isDone() { return ttl_ == 0; };
};
```

Example Parallel Reduction w/ matrix multiplication

Addition Task

```
template<class Type, Order Ord = Order::Row>
class AdditionTask : public hh::AbstractTask<
    MatrixBlockData<Type, 'c', Ord>,
    std::pair<std::shared_ptr<MatrixBlockData<Type, 'c', Ord>>, std::shared_ptr<MatrixBlockData<Type, 'p', Ord>>>> {

public:
    void execute(std::shared_ptr<std::pair<std::shared_ptr<MatrixBlockData<Type, 'c', Ord>>,
        std::shared_ptr<MatrixBlockData<Type, 'p', Ord>>>> ptr) override {

        auto c = ptr->first;
        auto p = ptr->second;
        assert(c->blockSizeWidth() == p->blockSizeWidth());
        assert(c->blockSizeHeight() == p->blockSizeHeight());

        if constexpr (Ord == Order::Row) {
            for (size_t i = 0; i < c->blockSizeHeight(); ++i) {
                for (size_t j = 0; j < c->blockSizeWidth(); ++j) {
                    c->blockData()[i * c->leadingDimension() + j] += p->blockData()[i * p->leadingDimension() + j];
                }
            }
        } else {
            for (size_t j = 0; j < c->blockSizeWidth(); ++j) {
                for (size_t i = 0; i < c->blockSizeHeight(); ++i) {
                    c->blockData()[i * c->leadingDimension() + i] += p->blockData()[j * p->leadingDimension() + i];
                }
            }
        }

        delete[] p->blockData();
        this->addResult(c);
    }
};
```

Example Parallel Reduction w/ matrix multiplication

Output State Manager

```
template<class Type, Order Ord = Order::Row>
class OutputState : public hh::AbstractState<MatrixBlockData<Type, 'c', Ord>, MatrixBlockData<Type, 'c', Ord>> {
private:
    std::vector<size_t>
        ttl_ = {};

    size_t
        gridHeightTTL_ = 0,
        gridWidthTTL_ = 0;

public:
    OutputState(size_t gridHeightTtl, size_t gridWidthTtl, size_t const &ttl)
        : ttl_(std::vector<size_t>(gridHeightTtl * gridWidthTtl, ttl)),
          gridHeightTTL_(gridHeightTtl), gridWidthTTL_(gridWidthTtl) {}

    void execute(std::shared_ptr<MatrixBlockData<Type, 'c', Ord>> ptr) override {
        auto i = ptr->rowIdx(), j = ptr->colIdx();
        --ttl_[i * gridWidthTTL_ + j];
        if (ttl_[i * gridWidthTTL_ + j] == 0) {
            this->push(ptr);
        }
    }
};
```

User Specification of Types for Nodes

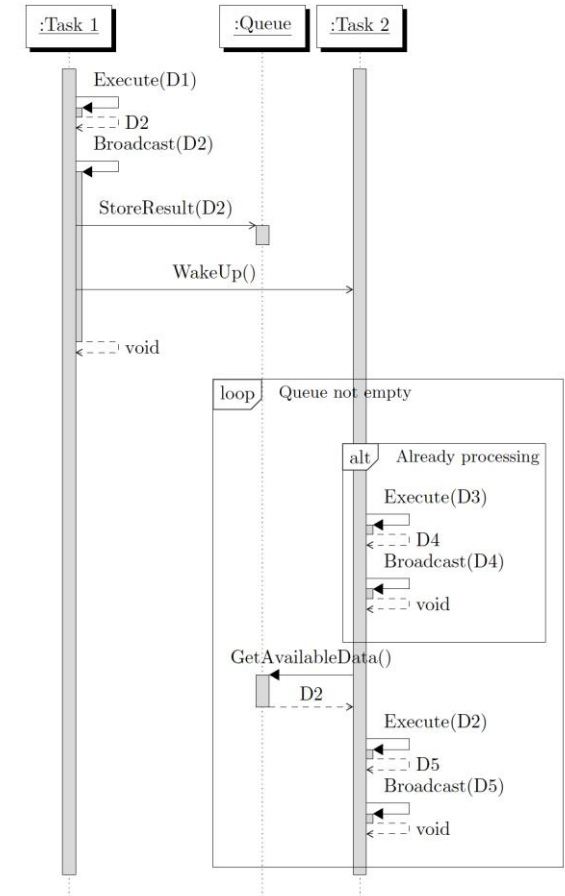
Q: Why is identification of input and output nodes by user necessary? Because compile time only?

- ▶ Every node is specialized by its input & output template types
- ▶ To instantiate a node, these templates must be defined
- ▶ This allow us to check with traits if the nodes are compatible with their use

Push vs. Pull Model & Work Creation

Q: Push doesn't preclude a pull model. How do tasks create data and generate work?

- ▶ Tasks inherit from a pure abstract class, Execute
- ▶ Specialized tasks overload from Execute class:
`virtual void execute(std::shared_ptr<Input>) = 0;`
- ▶ This execute method is called by the library when the corresponding input (of type Input) is available.
- ▶ Tasks have a method:
`void addResult(std::shared_ptr<TaskOutput> output)`
 - ▶ Explicitly called by developer in execute's implementation
 - ▶ Adds the “output” data to the queue between the task's node and a successor node



UML sequence diagram task I/O

Multiple output data

Q: Can a task have many out edges that get triggered prior to the completion of the whole task? For example, could split vector task have let a batch increment task (that no longer subdivides) start working as soon as $1/10^{\text{th}}$ of its work, i.e. the first

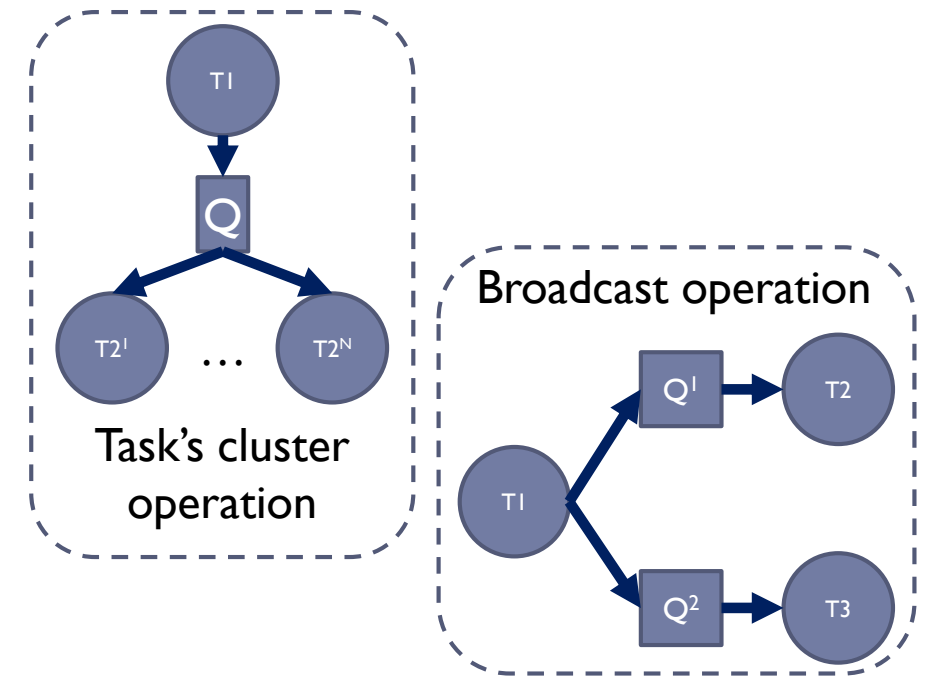
- ▶ If an “out edge” is activated, all are
- ▶ A task in “execute” can call “addResult” to add multiple data items to out queues/edges
- ▶ A task starts when an input is available. On waking up, it examines all its input queues for available data and invokes the corresponding “execute” method

Data movement

Q: Could pushData have been simply a data movement node?

- ▶ Not sure about what “data movement node” is
- ▶ The data are always wrapped in a shared_ptr
- ▶ So what is copied from a node to another is the shared_ptr

- ▶ When a task is duplicated, i.e., associated to multiple threads, all duplicated tasks are linked to the same queue
- ▶ If two tasks are linked to the same task, there are 2 queues, and the shared_ptr is put in each queue (broadcast)



First block time / Average block time

Q: I missed the talk, but I'd be interested to know where the gap comes from between “first block time” and “average block time”. Is the first exceptional, or are all blocks slightly different times (so last block is faster than average)? Is this an implementation detail or an effect of the runtime?

- ▶ The first is exceptional because it relates to how fast the next computation can begin.
- ▶ Our approach relates to the underlying data pipelining approach. As soon as data is available, the next computation can begin.
- ▶ Average block time in our experiments relates to the production rate of blocks after the first one.
- ▶ This is an effect of the runtime.
- ▶ No guarantees on the rate.

Termination Action/Condition

Q: finishPushingData tells it to not expect any more data beyond what's already been received, so that worker threads can quit. Relevant for persistent kernels.

- ▶ Exactly
- ▶ finishPushingData sends *terminate* to the graphs and will terminate all nodes in a breadth-first way
- ▶ By default, when *terminate* is sent, the graph waits for all inputs to be consumed before terminating

Dataflow representation

Q: What are the multiple conditions that could trigger dynamic execution, e.g. address (e.g., specific dependence) or generic type? How does a postdominator act on a stimulus from a then or else clause of a preceding (runtime) conditional?

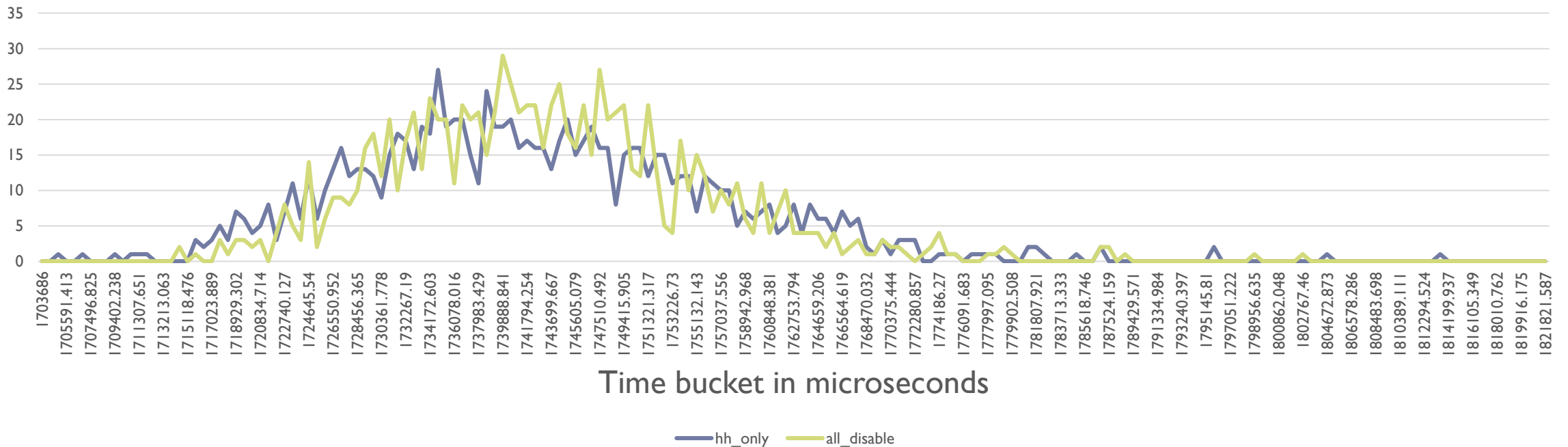
- ▶ Hedgehog is based on data flow
- ▶ Hedgehog is relying on data pipelining

Feedback cost

Q: How is feedback “costless?”

- ▶ We gather information only at node level.
 - ▶ Distributions with & w/o statistically indistinguishable.

Timer experiment for matrix multiplication of size (16k x 16k) with blocks of size (2k x 2k)



Group of nodes

Q: What are group nodes? Do they relate to hierarchy?

- ▶ I don't know what you relate to
- ▶ A cluster of task, or a group of task are the set of task plus its clone
- ▶ Construct of the graph as a node
- ▶ Execution pipeline

Graph and GPU binding

Q: “Bind a graph to a GPU”. Does this imply graphs do not span GPUs? Do they have external incoming/outgoing edges in order to synchronize with graphs on other GPUs?

- ▶ Our idea | Graph = | GPU
- ▶ If same algorithm on n GPUs → Graph duplication → Execution pipeline
- ▶ Execution pipeline will:
 - ▶ Duplicate a graph and associate the graph to a GPU
 - ▶ Have a rule to redirect an execution pipeline input to the right graph so the right GPU
- ▶ If data need to be shared amongst different GPU during the same algorithm, a state & state manager can be used amongst the graph's copies
 - ▶ Device Id can be stored in data, depending on boundaries for the algorithm, can be used to initiate copies between GPU memory addresses
 - ▶ CUDA tasks can also automatically enable peer access

Shutdown and GPU

Q: “Shutdown virtual method to break cycles”. This has implications on the underlying hardware which may introduce inefficiencies (for a GPU, this would be the case as upcoming work would not be able to be pre-fetched). A shutdown conditional node would solve this (i.e. prefetch disabled only for the successor to the shutdown node) but limits the locations where shutdown can occur. There are many other approaches: what does Hedgehog have in mind here?

▶ Tim any thoughts?

Asynchronous pre-fetch

Q: My own experiments on transparent multi-GPU execution indicate that performance is very sensitive to data locality. Is there a pinning operation for the data? Do you assume the pre-fetch is persistent? Would data migration be beneficial at all? What is the granularity with which these pre-fetches can happen?

- ▶ Hedgehog by itself only provides a pool of available data. The developer decides how data are pre-fetched
- ▶ No pinning
- ▶ Peer access
- ▶ Structure the dataflow into the graph to maximize locality, no guarantees
- ▶ Depends on domain decomposition for granularity, done by user
 - ▶ Inner vs outer product of GPU

State management

- ▶ **What does the system for managing state look like?**
- ▶ **What is per-node and what is system wide?**
- ▶ **Per above, how are triggering conditions specified? Can that be integrated into inter-node comms like MPI via OMP tasking?**

- ▶ **A state Manager:**
 - ▶ Is a specialized *single-threaded* task
 - ▶ That can't be duplicated
 - ▶ Need a state to be constructed
 - ▶ When `execute` is invoked:
 1. Lock the state
 2. Send the input data to the state
 3. Gather the output data from the state
 4. Unlock the state
 5. Send the output data from the state to the rest of the graph
- ▶ We do not represent system state
 - ▶ State is local, only between a StateManager and its inputs and outputs
 - ▶ Task1 → StateManager → Task2
- ▶ We do not have a way to represent the global state of the computation, the state manager is only a graph's node, so will deal only with local information
- ▶ Tim ? Walid →
- ▶ OMP can be done within a task if needed
- ▶ MPI requires strict ordering for data transfers
 - ▶ Hedgehog is fully asynchronous, so very difficult
 - ▶ MPI communication can be done outside of the graph, but incurs significant overhead

Memory management

- ▶ **Can alloc/free be visible both as a node in a graph and as a code operation inside a graph?**
- ▶ **Are special interfaces needed to make alloc/free visible to this dep system, especially for alloc/free in code inside an execution action?**
- ▶ **Is the totally memory accessible under admin control?**
- ▶ **Is the same dependence system used to manage memory availability as for execution and data deps?**
- ▶ **Can in or out arcs regarding memory availability occur from/to the middle of an execution action?**
- ▶ Memory manager is a tool that the tasks use
- ▶ A node is connected to a memory manager
 - ▶ Can have multiple independent memory managers for multiple tasks
- ▶ Memory manager will deal with alloc/free, and Node will interact with the memory manager
- ▶ There are no special interfaces as the alloc/free are encapsulated into memory manager
- ▶ There is no admin control, user controls pool size
- ▶ A memory manager is attached to a state manager or a task, so the data is acquired in “execution” method
- ▶ If a memory pool is empty, the task that uses that memory pool will wait, other tasks can operate as usual as long as data is available

Data as a first-class citizen

- ▶ **In what ways is data treated as a first-class citizen?**
- ▶ Data is encapsulated into `shared_ptr`
 - ▶ Avoid not needed real data movement
- ▶ We tried to make data management easier for computation on really small computers, or limited-memory devices as GPUs
- ▶ Data in Hedgehog controls the scheduling and parallelism of our nodes and graphs
 - ▶ Rely on data pipelining

Headers only

- ▶ **HiHAT's dispatch is headers only for perf reasons**
- ▶ Hedgehog is header only because of:
 - ▶ The templates
 - ▶ Perf reasons

Compile time focus

- ▶ **Burdens the user with identifying input and output nodes?**
- ▶ Explicitness of nodes and graphs interfaces, what they accept and produce
 - ▶ Clear model
 - ▶ Helps collaboration
- ▶ Correctness check at compile time

Modern C++

- ▶ **SFINAE** – can be hard to interpret compile-time error messages
- ▶ **HiHAT** didn't take a templated approach since it used a **C ABI** and wasn't all header only.
- ▶ **Others** (Tal Ben-Nun did, and Mathias Noack talked about it) did a shim with **C++ templates**
- ▶ **Could have used traits**, since just used at compile time
- ▶ True about SFINAE
 - ▶ We just use it in the memory manager
 - ▶ Future usage of C++ Concepts with C++20 standards (little mistake on last time report)
- ▶ Traits can help having meaningful message with the usage of `std::static_assert`

Backup
