

CS 744: GANDIVA

Shivaram Venkataraman

Fall 2019

ADMINISTRIVIA

- Course project proposal
- Midterm

Machine Learning

Bismarck

Supervised learning, Unified Interface
Shared memory, Model fits in memory

Parameter Server

Large datasets, large models (PB scale)
Consistency model, Fault tolerance

Tensorflow

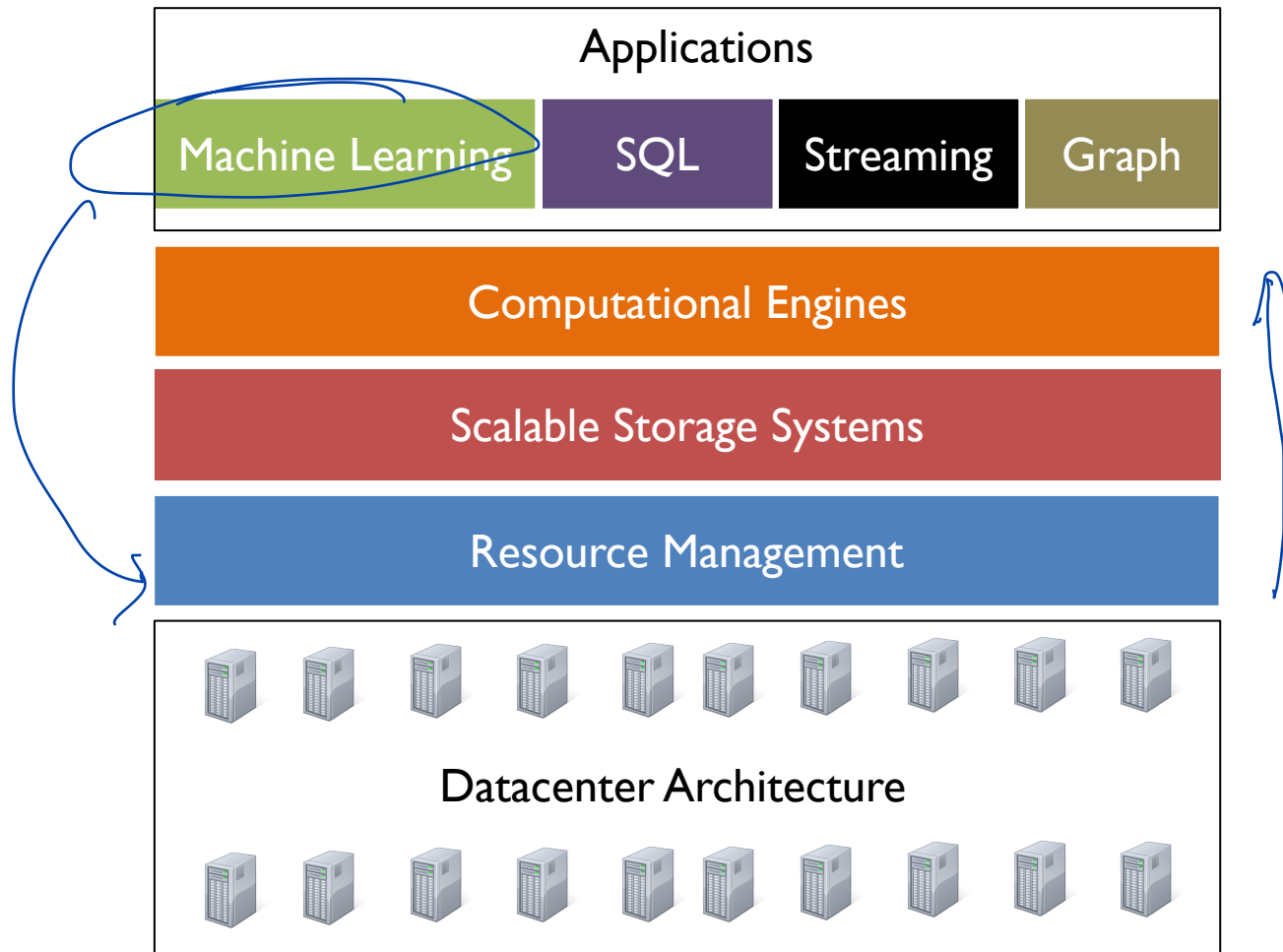
Need for flexible programming model
Dataflow graph, Heterogeneous accelerators

Ray

Reinforcement learning applications
Actors and tasks, Local and global scheduler

Inference

Clipper

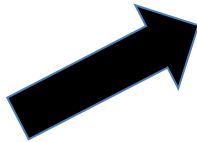


MACHINE LEARNING WORKFLOW?

 PyTorch

 TensorFlow

 mxnet



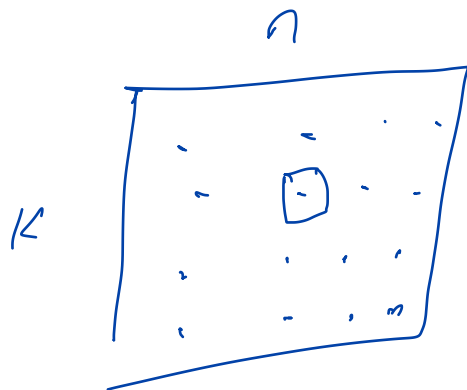
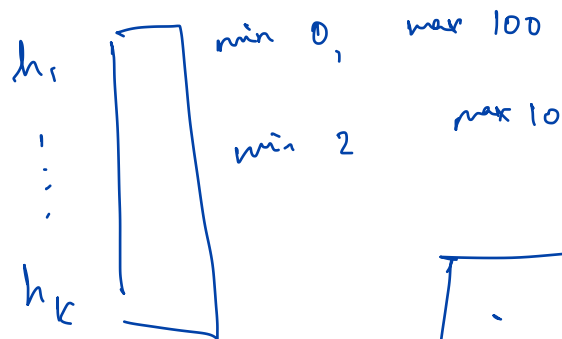
```
class Net(nn.Module):  
    def __init__(self):  
        super(Net, self).__init__()  
        self.conv1 = nn.Conv2d(1, 10, kernel_size=5)  
        self.conv2 = nn.Conv2d(10, 20, kernel_size=5)  
        self.conv2_drop = nn.Dropout2d()  
        self.fc1 = nn.Linear(320, 50)  
        self.fc2 = nn.Linear(50, 10)
```



WORKLOAD

Feedback-driven exploration

Hyper parameter search



$K \times N$ jobs

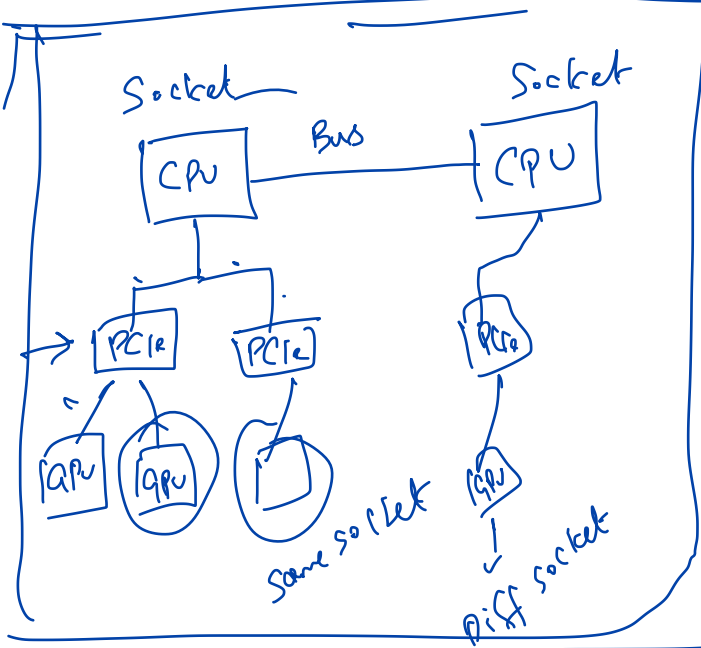
$$\min f(x) + \lambda \|x\|_2$$

time

value? regularization

→ learning rate

Early stopping
truncate execution
of suboptimal hyperparameters



~~AFFINITY~~

Locality

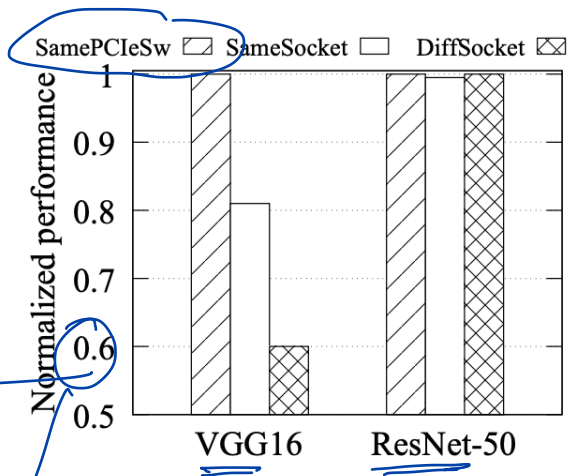


Figure 1: Intra-server locality.

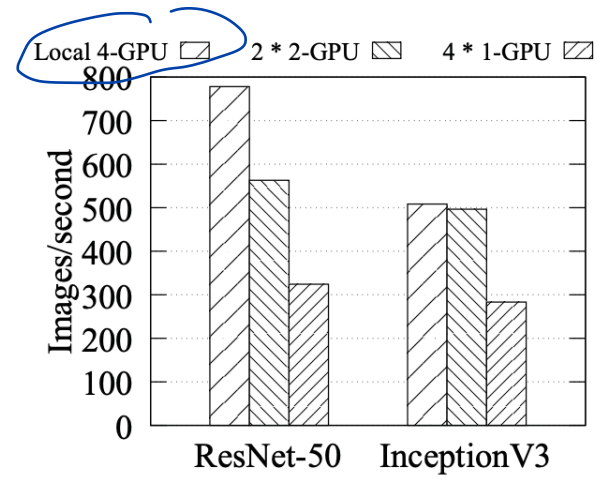
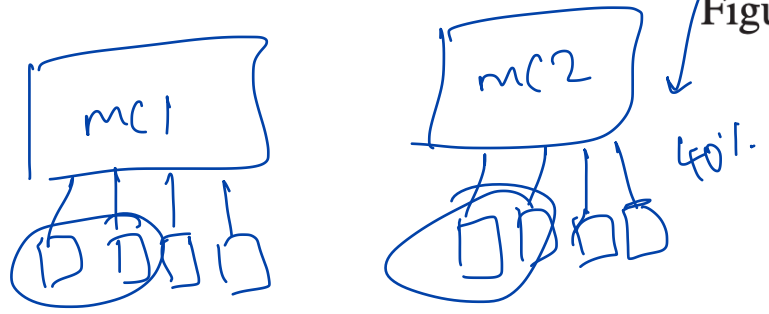


Figure 2: Inter-server locality.



ML models are sensitive to locality of placement

INTRA JOB PREDICTABILITY

while

Sample



Compute ∇f

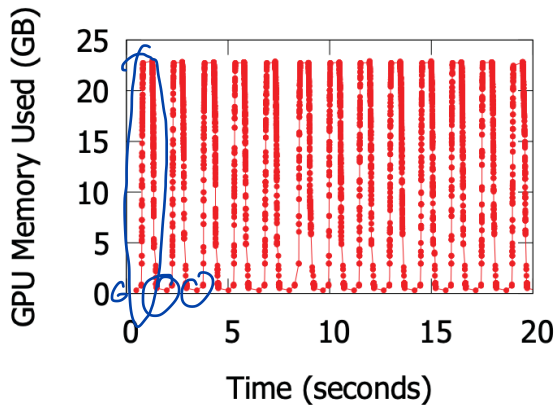
Aggregate Sync/async



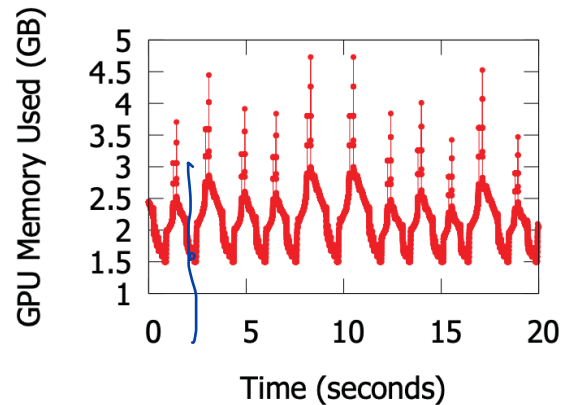
end

Predictability $\rightarrow$ useful

Heterogeneity

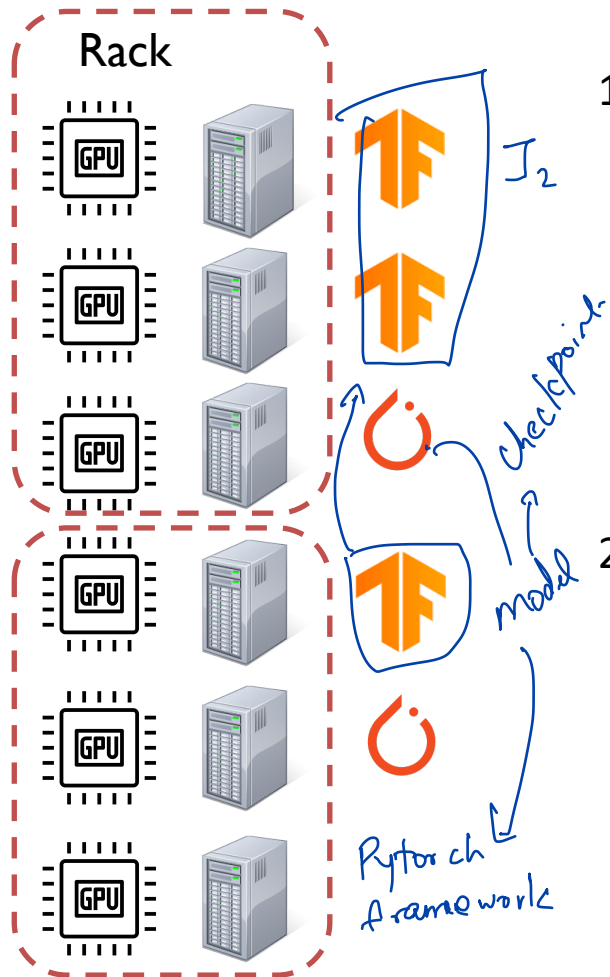


(a) ResNet50/Imagenet



(b) GNMT/WMT'14 En-De

MECHANISMS (1)



1. Suspend-Resume

→ suspend job 1, place job 2 on GPUs for time quantum → 1 min

→ suspend job 2, resume Job 1

→ Use Intra job predictability to suspend MB finishes

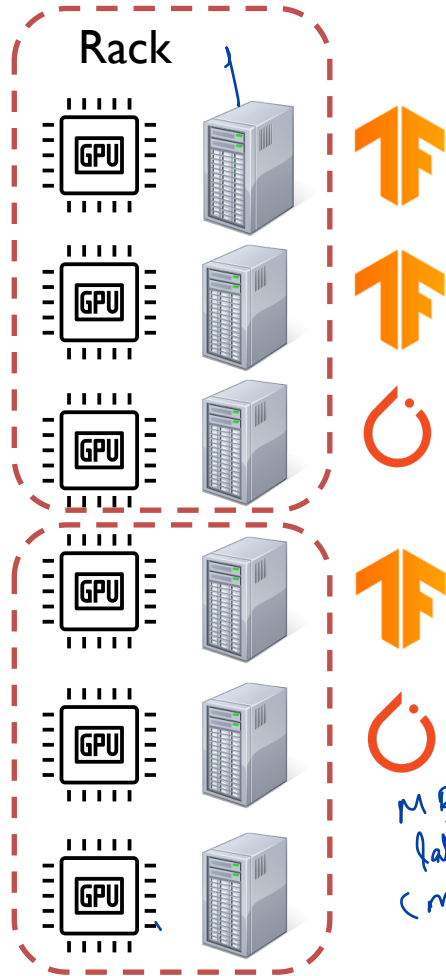
2. Migration

→ More tasks to improve locality

or improve utilization

→ Copy state old mc and new mc
↳ 1GB ≈ longer time period

MECHANISMS (2)



3. Grow-shrink

→ if GPUs idle, can we give them to existing tasks

$b = 256$, 2 GPUs

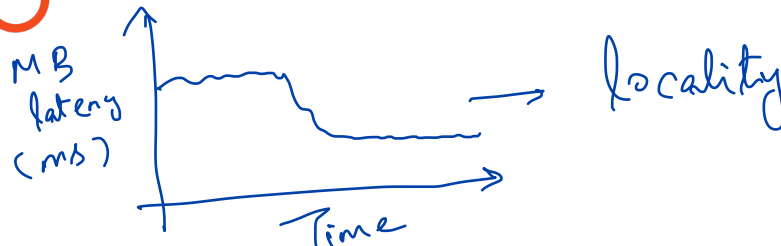
+ 2 GPUs when they become free

4. Profiling

→ Mem usage

mini batch time taken →

lookup table of profiles



SCHEDULING POLICY

Goals

early feedback

cluster efficiency

cluster-level fairness? → Non-goal

Two modes

Reactive

Introspective

"job similarity"

REACTIVE MODE

React to events

Job arrivals, departures, failures

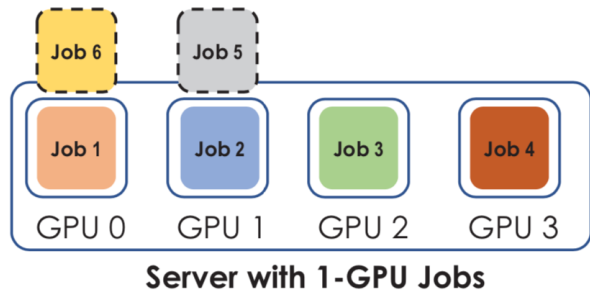
Hierarchical Preference

Nodes with same "affinity"

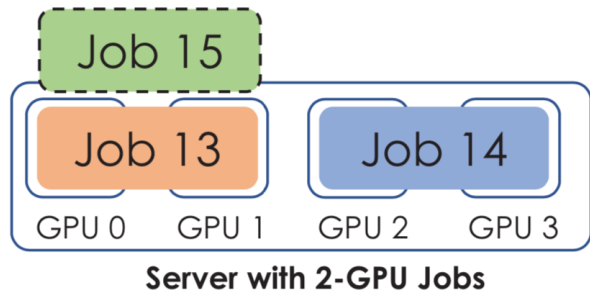
Nodes with "different affinity"

Nodes with "no affinity"

Suspend-resume ... *same affinity*



4 1-GPU
Jobs
J1 - J4



2 2-GPU
Jobs
Job 13, Job 14

INTROSPECTIVE MODE

bins: // tf-job-3243:8888
- worker-1

Monitor and optimize placement of jobs periodically

Actions

Packing



Over subscribe resources
by using predictability



greedy → profiling
undo mis

Migration

Grow-shrink

Priority = High 2 time quantum
low 1 time quantum

DISCUSSION

<https://forms.gle/aHYbNcTFdGjtXefj9>

What are some guarantees provided by Mesos that are not provided by Gandiva? Explain with an example.

↳ DRF as its scheduling policy

Mesos - "resource guarantee"

- strategy proof, while Gandiva. Jobs can inflate mini batch size
- Fairness → sharing incentive

Are mechanisms in Gandiva also useful in a cluster running Apache Spark jobs? Provide one example either for or against

↳ Migration is not that useful because of backup tasks?

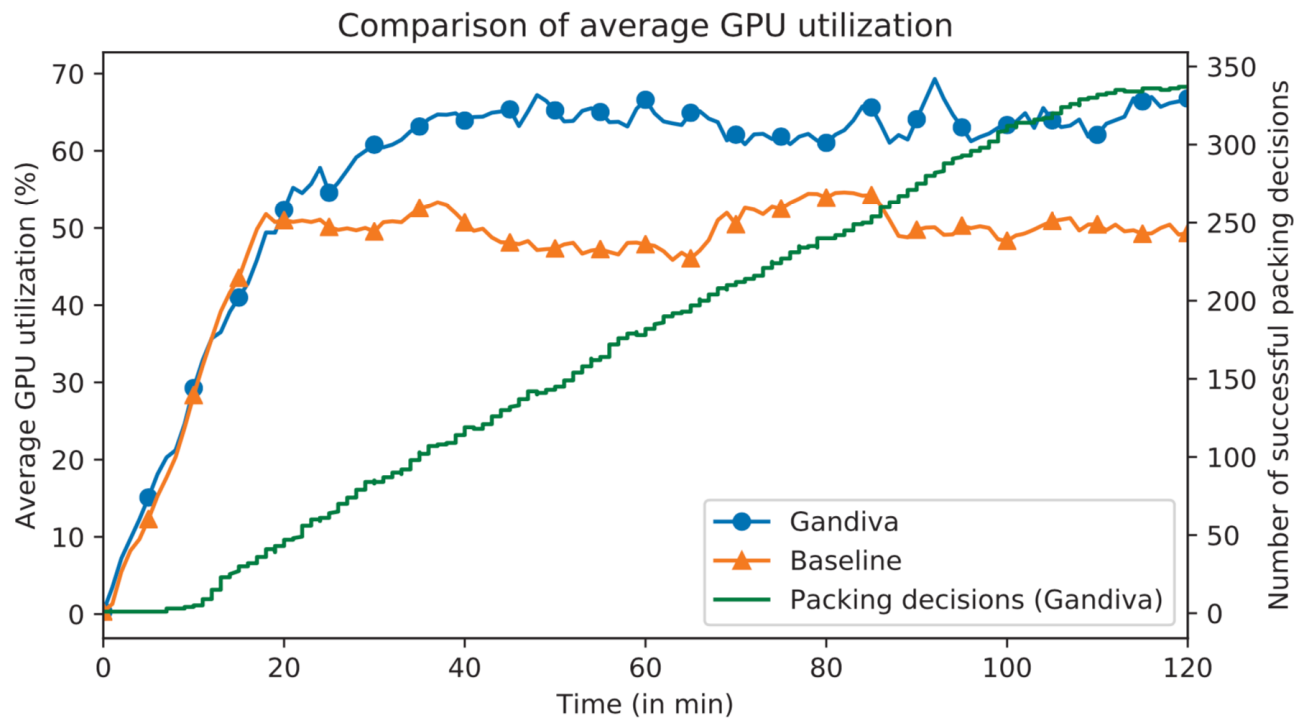
↳ need to transfer data, RDD or shuffle data

Can be
useful
for
large
shuffles

↳ Suspend / Resume may not be useful because of lack
of predictability

↳ PageRank, Logistic Regression do have predictability
within the DAG

↳ Naturally happening task boundaries



NEXT STEPS

New module on SQL!

Course project introductions

Midterm