

Good morning!

CS 744: MESOS

Shivaram Venkataraman

Fall 2020

ADMINISTRIVIA

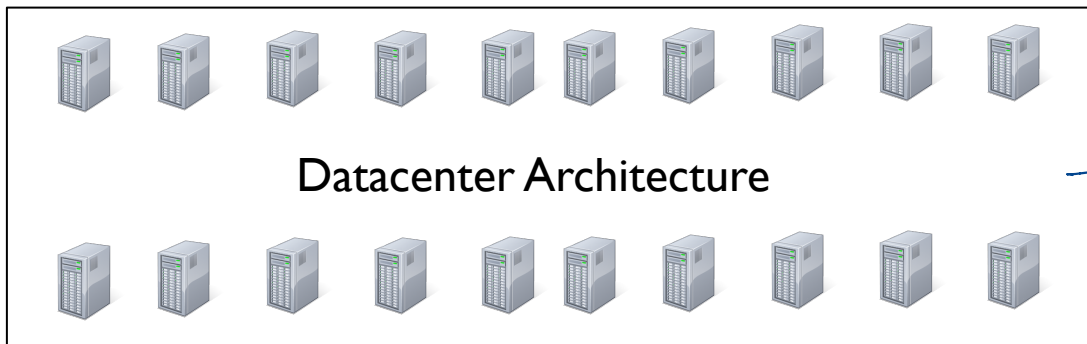
- Assignment 1: How did it go? → Fill out the poll!
- Assignment 2 out tonight → ML distributed
- Project details
 - Create project groups → ~ 3 students
 - Bid for projects/Propose your own → next week
 - Work on Introduction
 - ↳ 1-2 page
- Check-in
- Final report and poster session

COURSE FORMAT

Paper reviews

“Compare, contrast and evaluate research papers”

Discussion



Assignment

MR, Spark

GFS





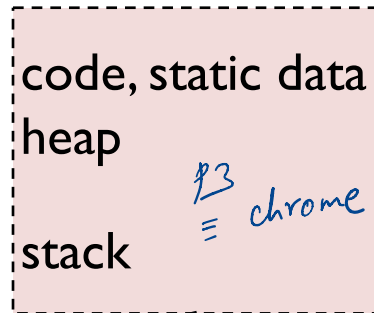
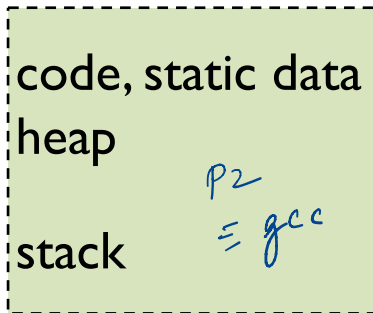
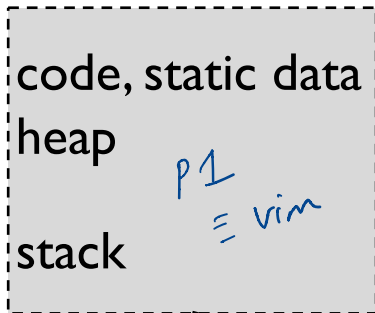
MapReduce

GFS

Spark

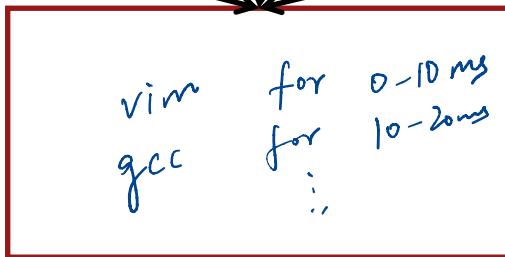


BACKGROUND: OS SCHEDULING

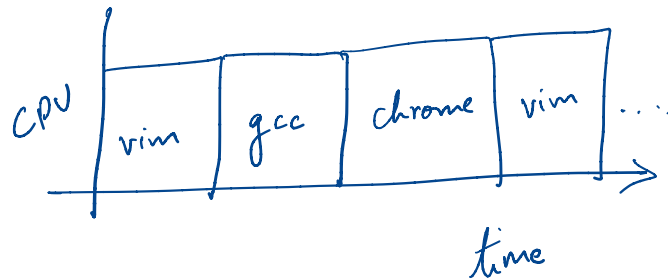


time sharing

How do we
share CPU
between
processes ?



CPU



CLUSTER SCHEDULING

Fairness → similarity between OS & cluster

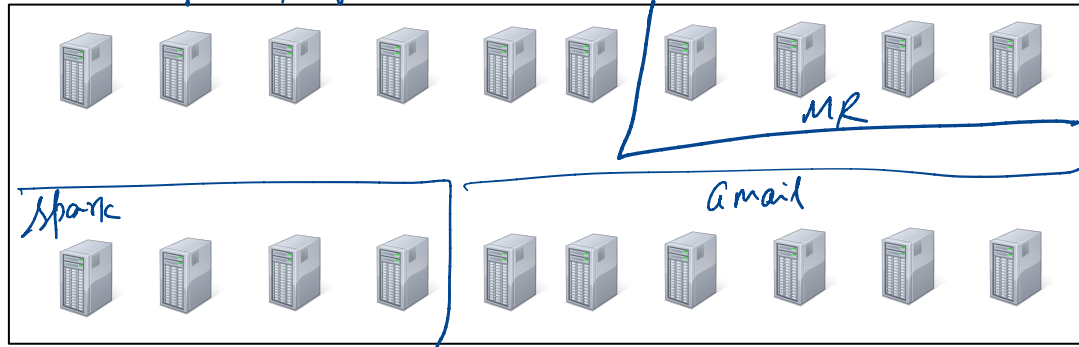
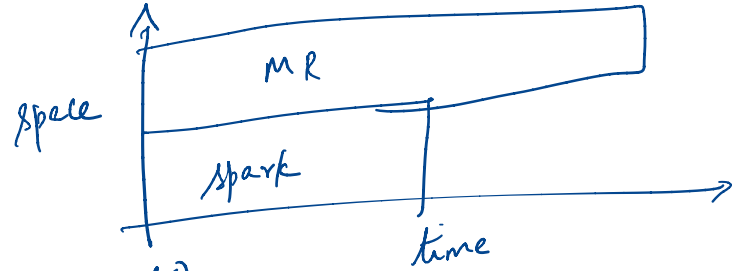
Multi-core or NVMA-aware scheduling

Scale → large number of machines
↳ one scheduler?

space sharing

fault tolerance

Preferences (placement, constraint)



TARGET ENVIRONMENT

Utilization
↳ Not all resources are used

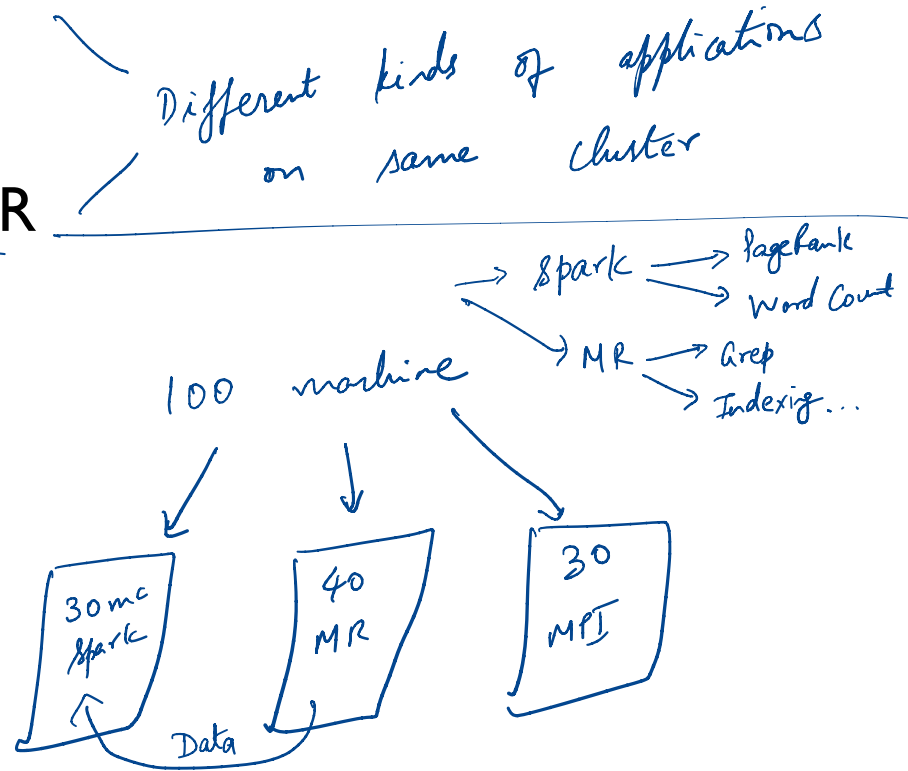
Sharing →

Multiple MapReduce versions

Mix of frameworks: MPI, Spark, MR

Data sharing across frameworks

Avoid per-framework clusters



DESIGN

Two-level scheduling

odd numbers : 3 or 5
is common

ZooKeeper
quorum

Mesos

↳ Scheduling
across
framework

Per-frame
work
scheduler

Simplicity at a cluster
wide scheduler

↳ Add new frameworks
in the future

Scalability, Flexibility

Spark
Driver

Hadoop
scheduler

MPI
scheduler

Mesos
master

Standby
master

Standby
master

Mesos slave
Hadoop
executor
task

Mesos slave
MPI
executor
task

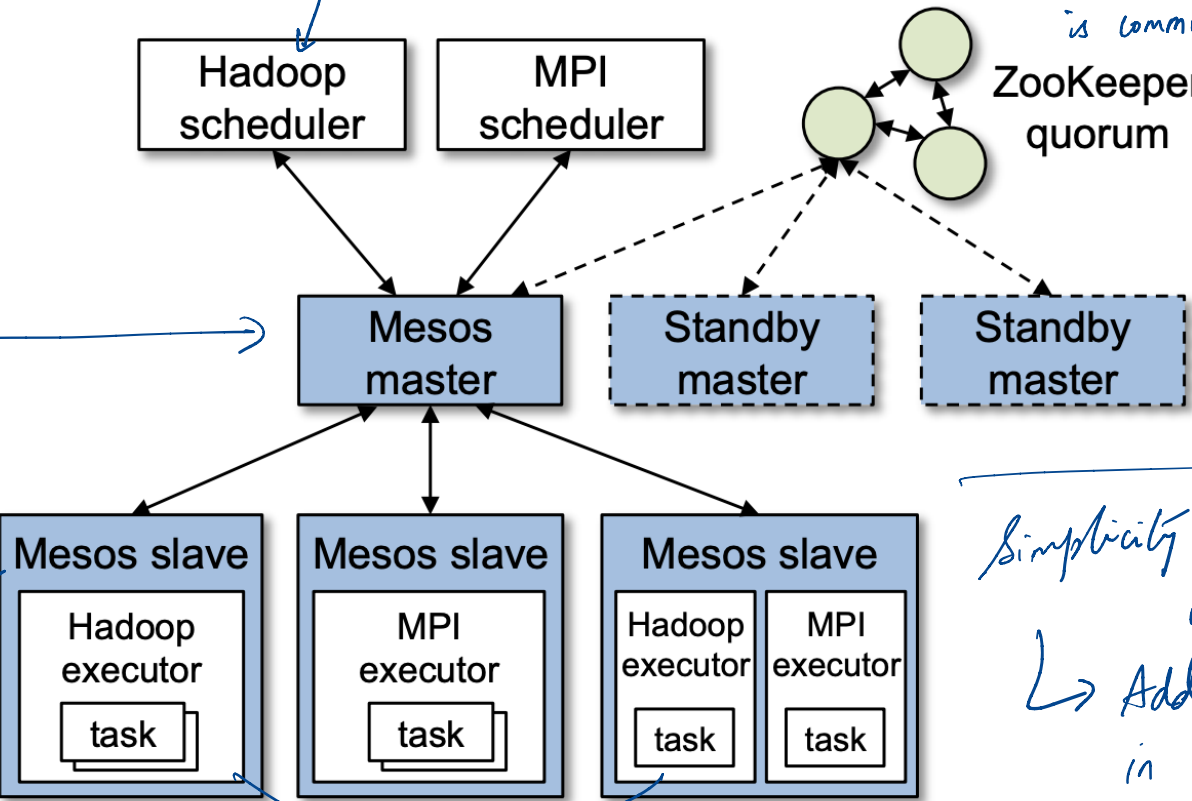
Mesos slave
Hadoop
executor
task
MPI
executor
task

GFS: single
namenode
Chunkserver

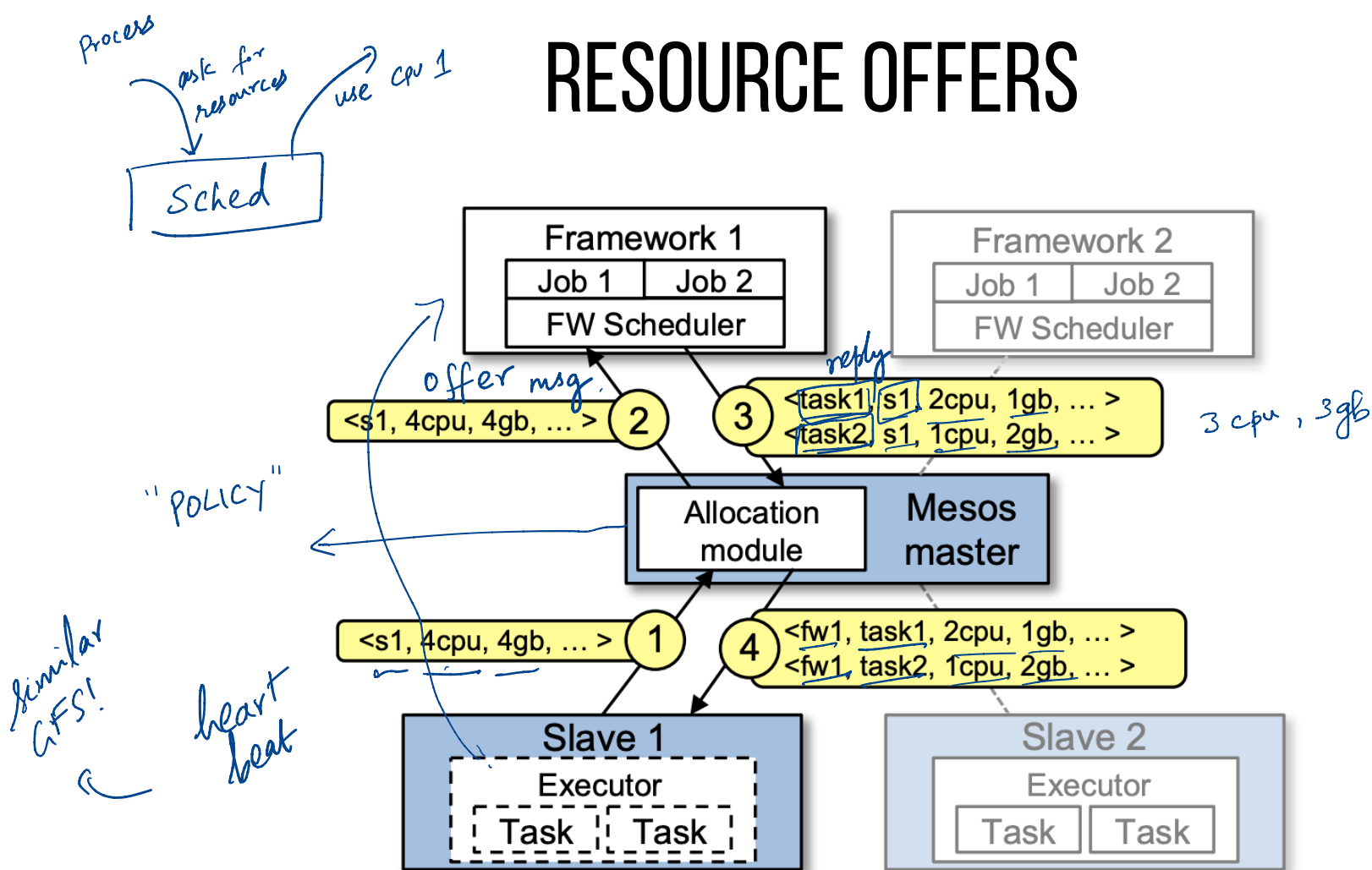
single
master

These run
on every
machine

Shuffle



RESOURCE OFFERS



CONSTRAINTS

Examples of constraints :

Data locality → soft
GPU machines → hard

Constraints in Mesos:

↳ frameworks can reject offer

↳ "Filters" → Boolean functions

DESIGN DETAILS

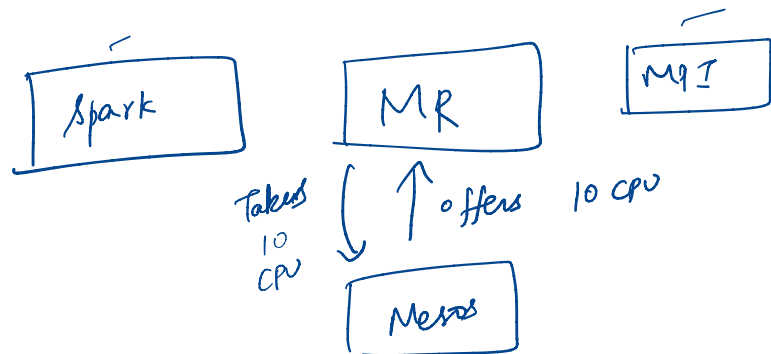
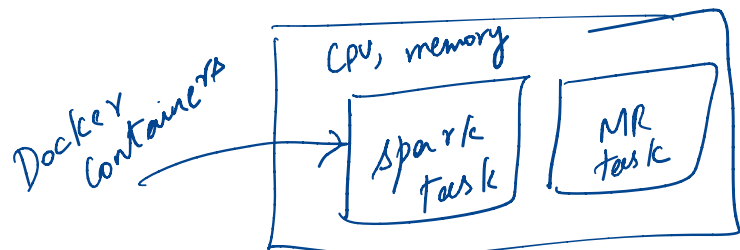
Allocation:

Guaranteed allocation, revocation

↳ short-lived tasks!

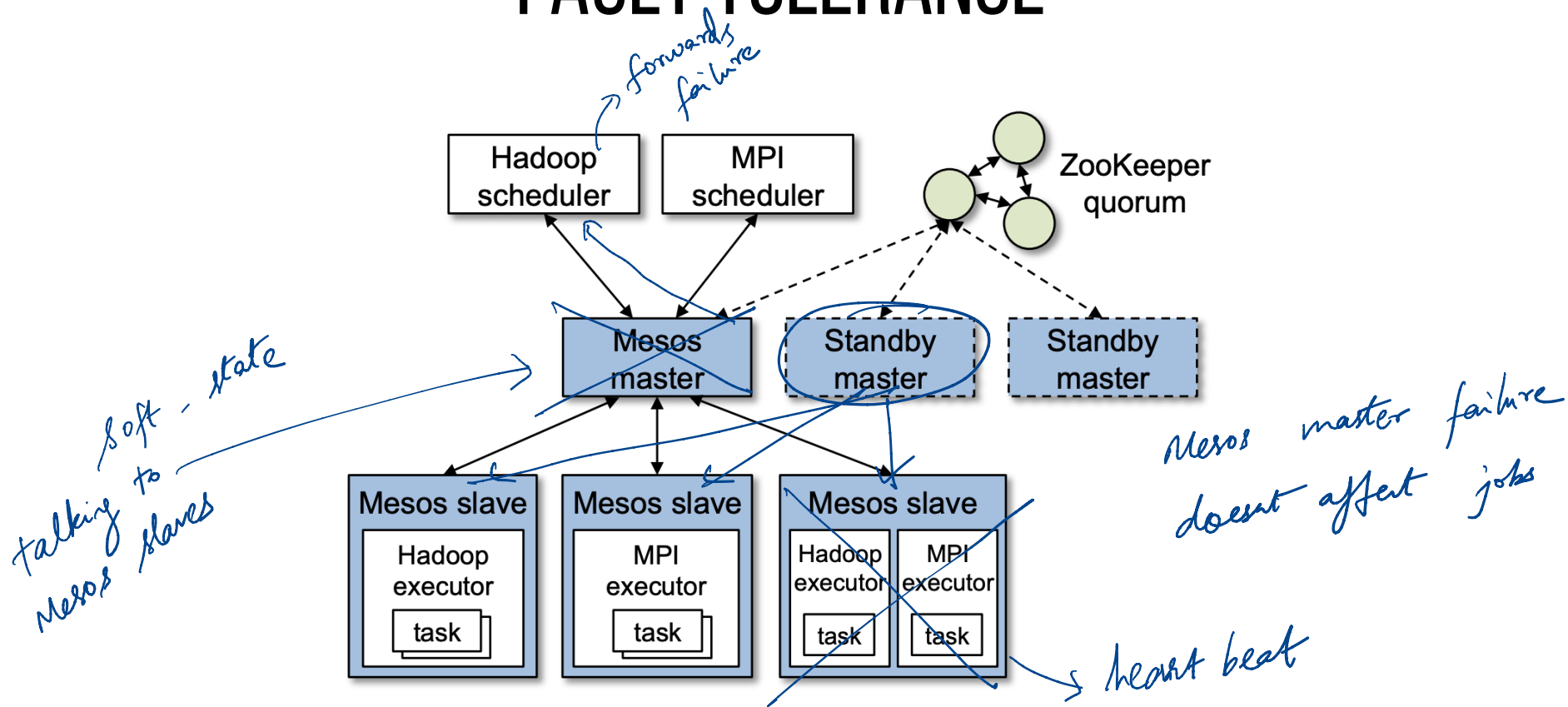
Isolation

Containers (Docker)



↳ long running tasks can be pre-empted, when other frameworks express interest

FAULT TOLERANCE



PLACEMENT PREFERENCES

What is the problem?

↳ If you have more frameworks with prefs than machines available in the cluster

How do we do allocations?

↳ weighted lottery scheme
make offers proportional in size to the overall resources that a framework needs

CENTRALIZED VS DECENTRALIZED

Centralized

→ Scalability

~100s of frameworks
each ~100s of apps

Optimal solution

Decentralized

handle new frameworks
↓
Complexity for framework
developer

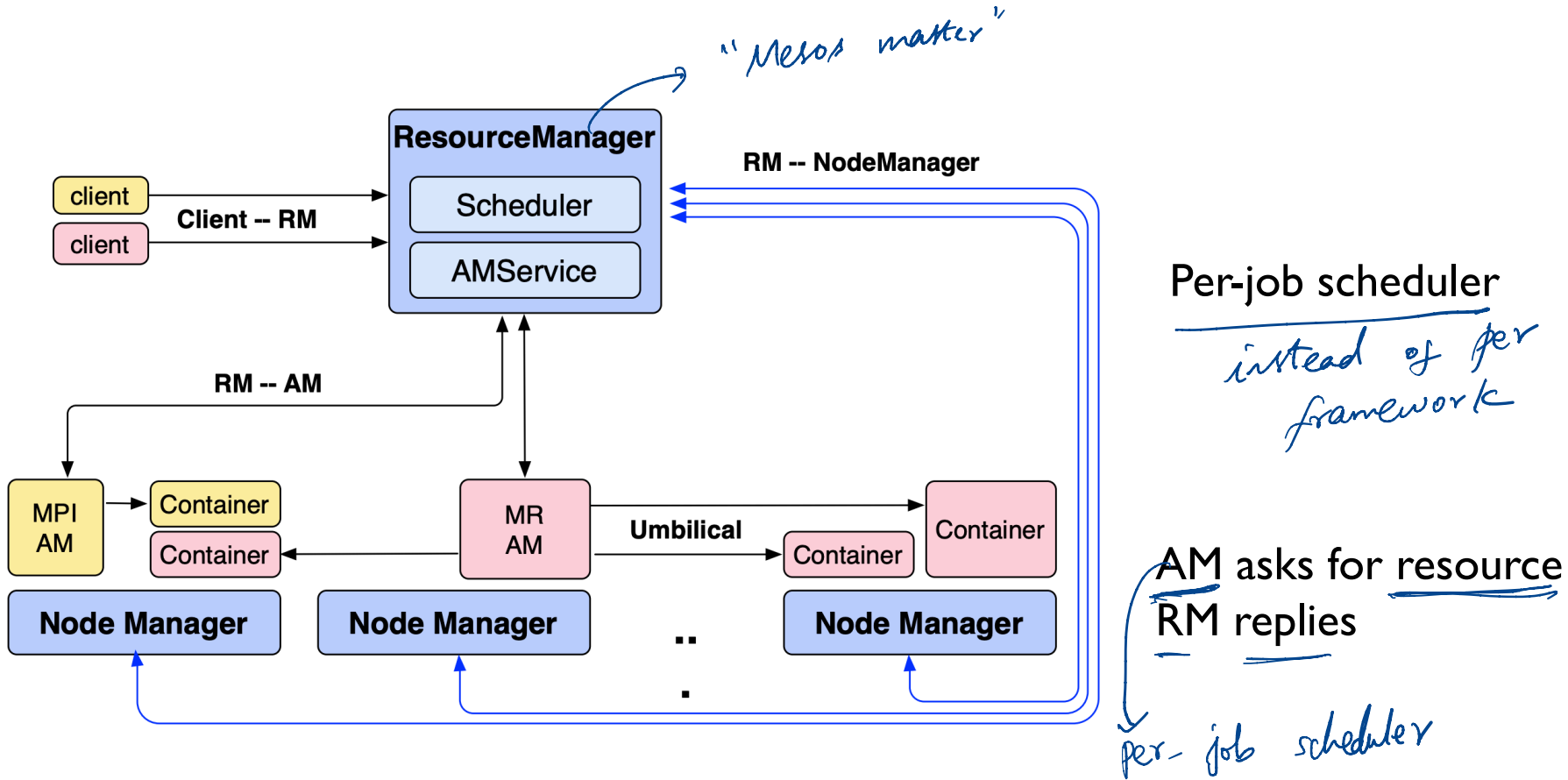
CENTRALIZED VS DECENTRALIZED

Framework complexity ✓

Fragmentation, Starvation → If resource offers
are too small

Inter-dependent framework

COMPARISON: YARN → Apache Hadoop



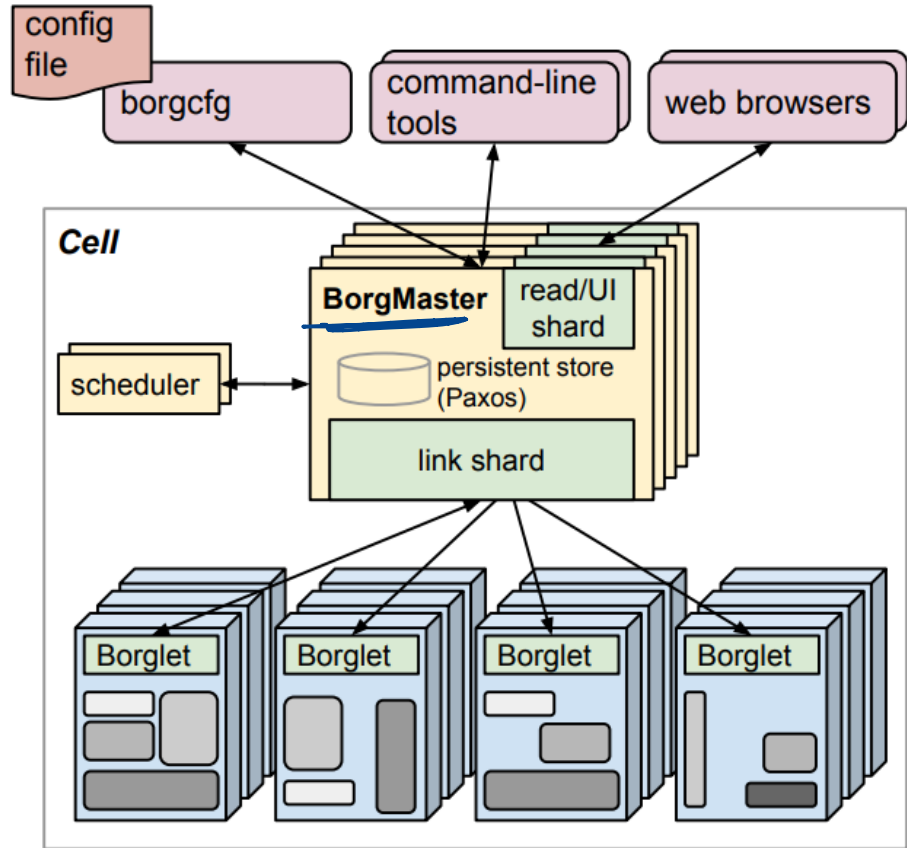
COMPARISON: BORG *→ Google*

Single centralized scheduler

Requests mem, cpu in cfg
Priority per user / service

→ Better packing

Support for quotas / reservations



SUMMARY

- Mesos: Scheduler to share cluster between Spark, MR, etc.
- Two-level scheduling with ^{framework}~~app~~-specific schedulers
- Provides scalable, decentralized scheduling
- Pluggable Policy ? Next class!

DISCUSSION

<https://forms.gle/urHSeukfyipCKjue6>

What are some problems that could come up if we scale from 10 frameworks to 1000 frameworks in Mesos?

→ Fragmentation / starvation odds go up

→ Master bottleneck?

↳ it takes time to wait for frameworks to reply

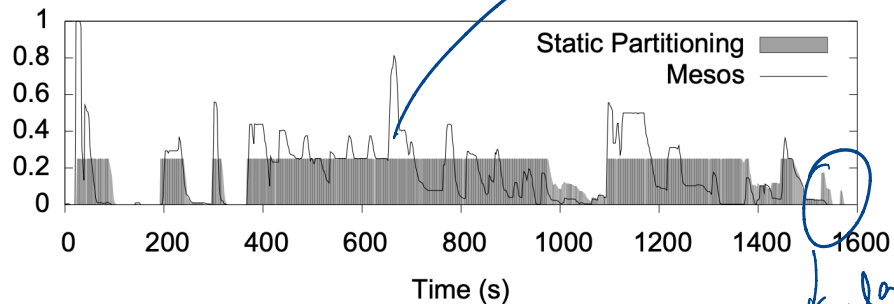
→ Pre-emption? Yes.

→ Failure recovery takes longer? why?

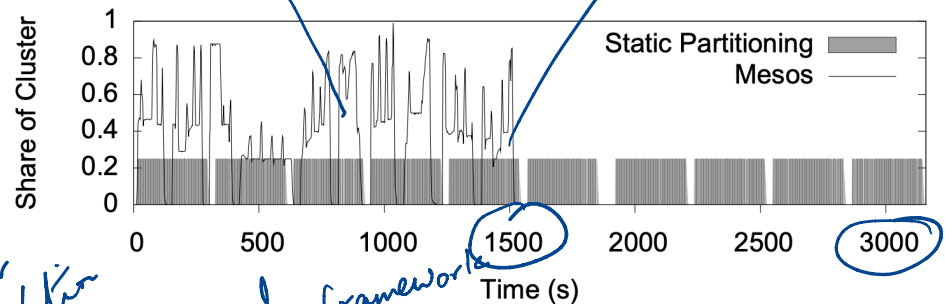
Mesos master has soft state ~ unclear?

Share of Cluster

(a) Facebook Hadoop Mix

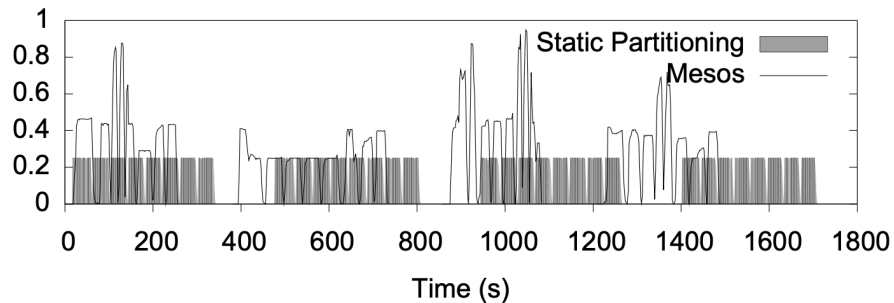


(b) Large Hadoop Mix

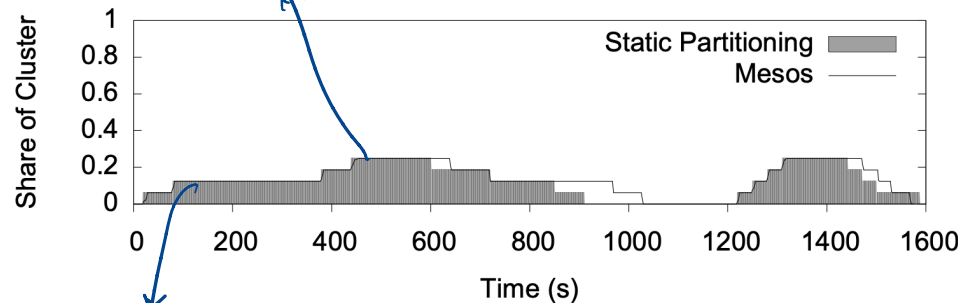


(c) Spark

Share of Cluster



(d) Torque / MPI



Elastic
can use
framework

~ 2x
faster
completion

similar
completion
times

Rigid
framework

able to
maintain
MPI's share

List any one difference between an OS scheduler and Mesos

↳ Motivation part of the lecture

Spark on oversubscribed clusters

↳ long running jobs

Executor → cache RDDs
→ "shuffle files" on local disk

↓
→ gets pre-empted → cache is blown away
↓
→ shuffle files long lived
↓
beyond "guaranteed slave"

Data locality

↳ Input (HDFS)

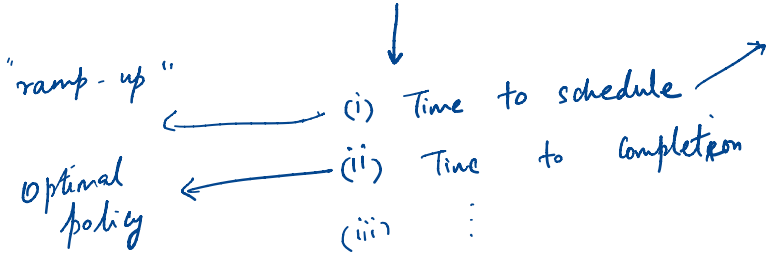
↳ Memory (Executor)

↳ shuffle (Executor)
output

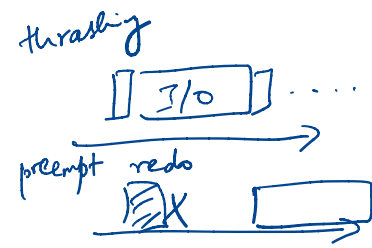
↓ Coarse Grained
Executor Backend

resource offers

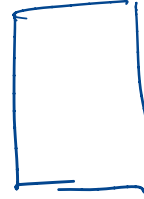
↳ how does it "perform" "better"?



Comparisons with YARN, Borg



NEXT STEPS



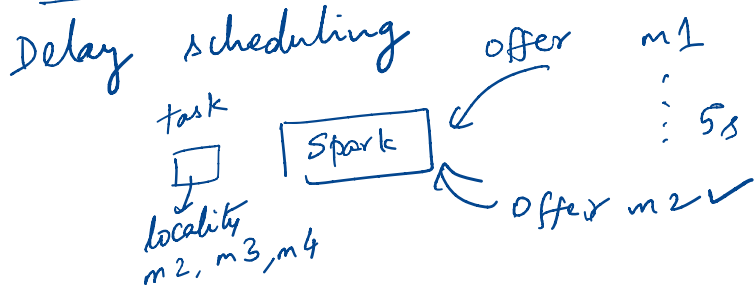
Pre-emption
↳ redo "work"

Next class: Scheduling Policy

Assignment 2 will be released

Further reading

- <https://www.umbrant.com/2015/05/27/mesos-omega-borg-a-survey/>
- <https://queue.acm.org/detail.cfm?id=3173558>



wait for 5s or so
after m1 offer is made