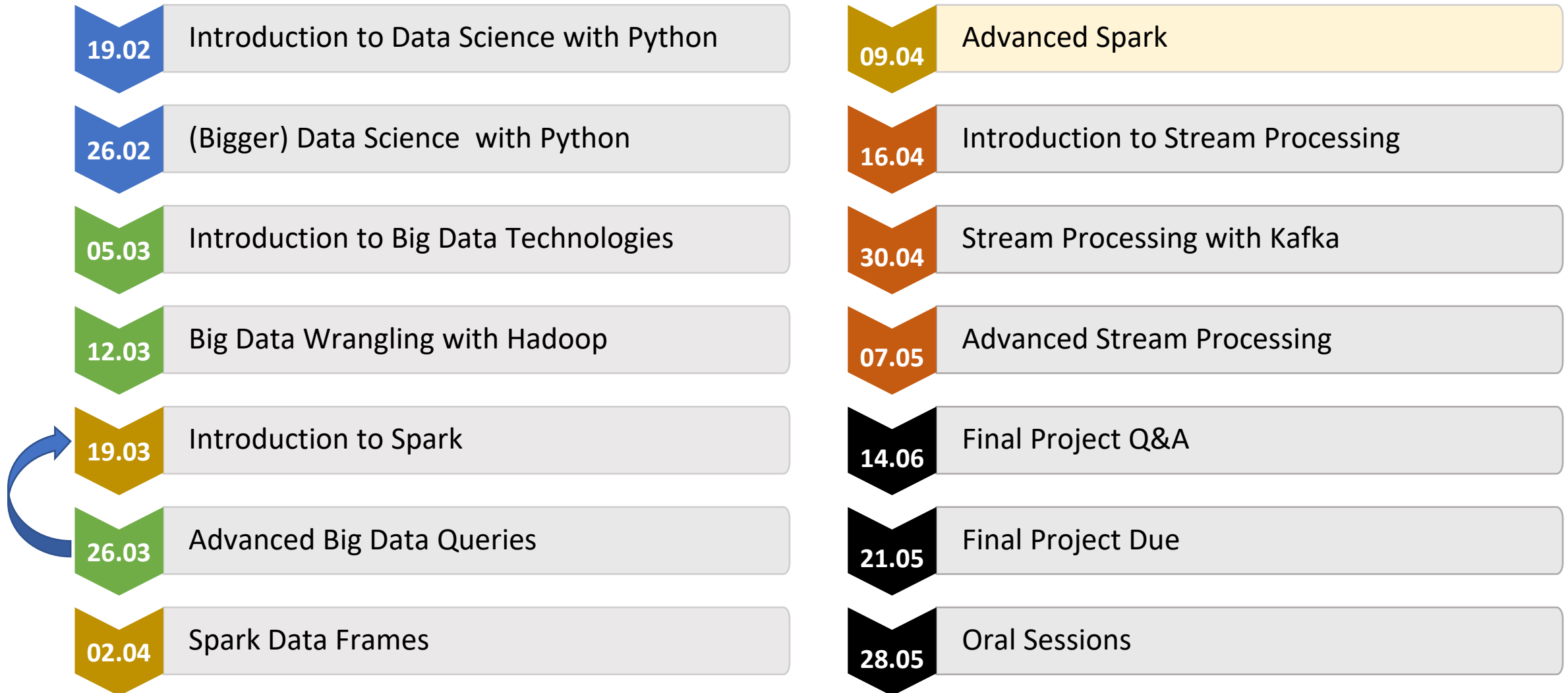


THE DATA SCIENCE LAB

Advanced Spark Optimization

COM 490

Agenda 2025 – module 3b



Week 7 Recap

- Revisit RDDs
- DataFrames
 - Structured data
 - Higher level API, declarative programming
- RDDs vs DataFrames
- DataFrames under the hood
 - PySpark internals (Py4J)
 - Catalyst optimization (query plan optimization)
- Exercises: DataFrames, Covid, Twitter

Week 7 – Questions?

Today's Agenda

- Spark parallelization
- Spark and Big Data
- The cost of shuffle operations and memory
- Spark partitioning
- Partitioning demo
- Exercise:
 - Spark DataFrames with partitioning
- Homework 3

Introduction to Spark Parallelization

RDD - Revisited

- RDD - Resilient Distributed Dataset
 - Immutable: once an RDD is created it cannot be modified, actions create new RDDs
 - Lineage: RDD points to its parents
 - Resilient: if a node dies data is easily regenerated
 - Partitioned → pieces of RDD ***distributed*** over the cluster

RDD – parallelize API

- Creation of RDD

```
sc.parallelize(data)
```

- Reference external or local data
- Each piece of data becomes a partition
- Each partition is processed in a task
- Tasks are executed in parallel

- Why parallelize?

- Take advantage of distributed resources

Spark units of work

- **Job**

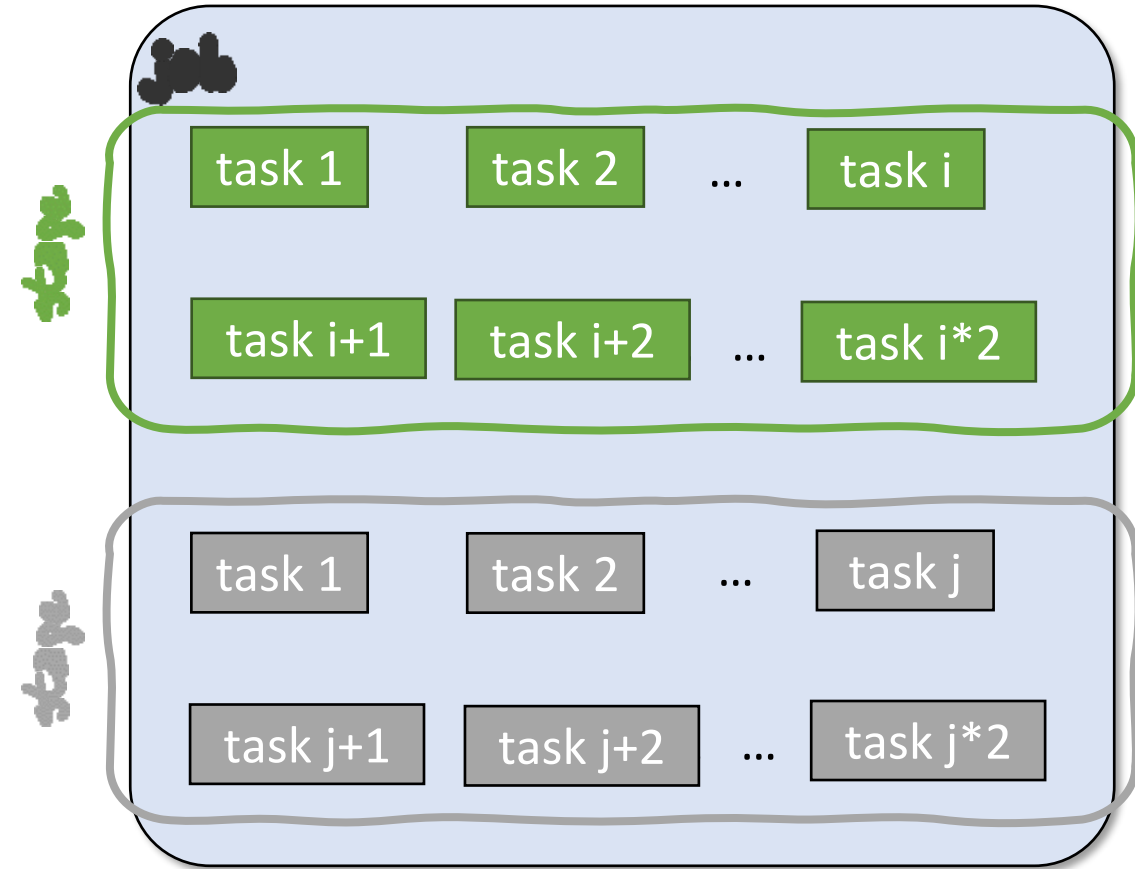
- Sequence of **Stages**
- Triggered by an **action**
- `collect()`, `read()`, `write()`

- **Stage**

- Sequence of **Tasks**
- Sequences run in parallel without a **shuffle**
- Output of an execution plan

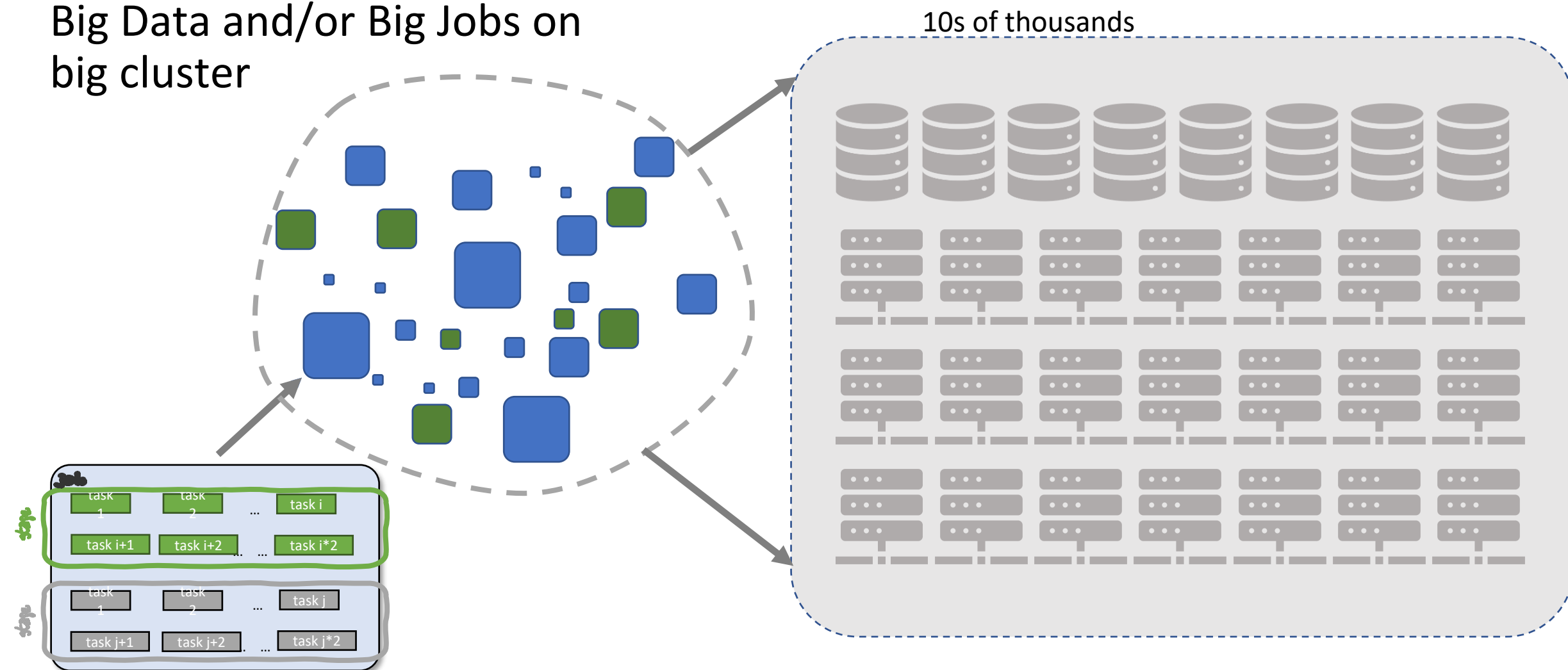
- **Task**

- Single operation applied to a partition
- Triggered by a **transformation**
- `map()`, `filter()`

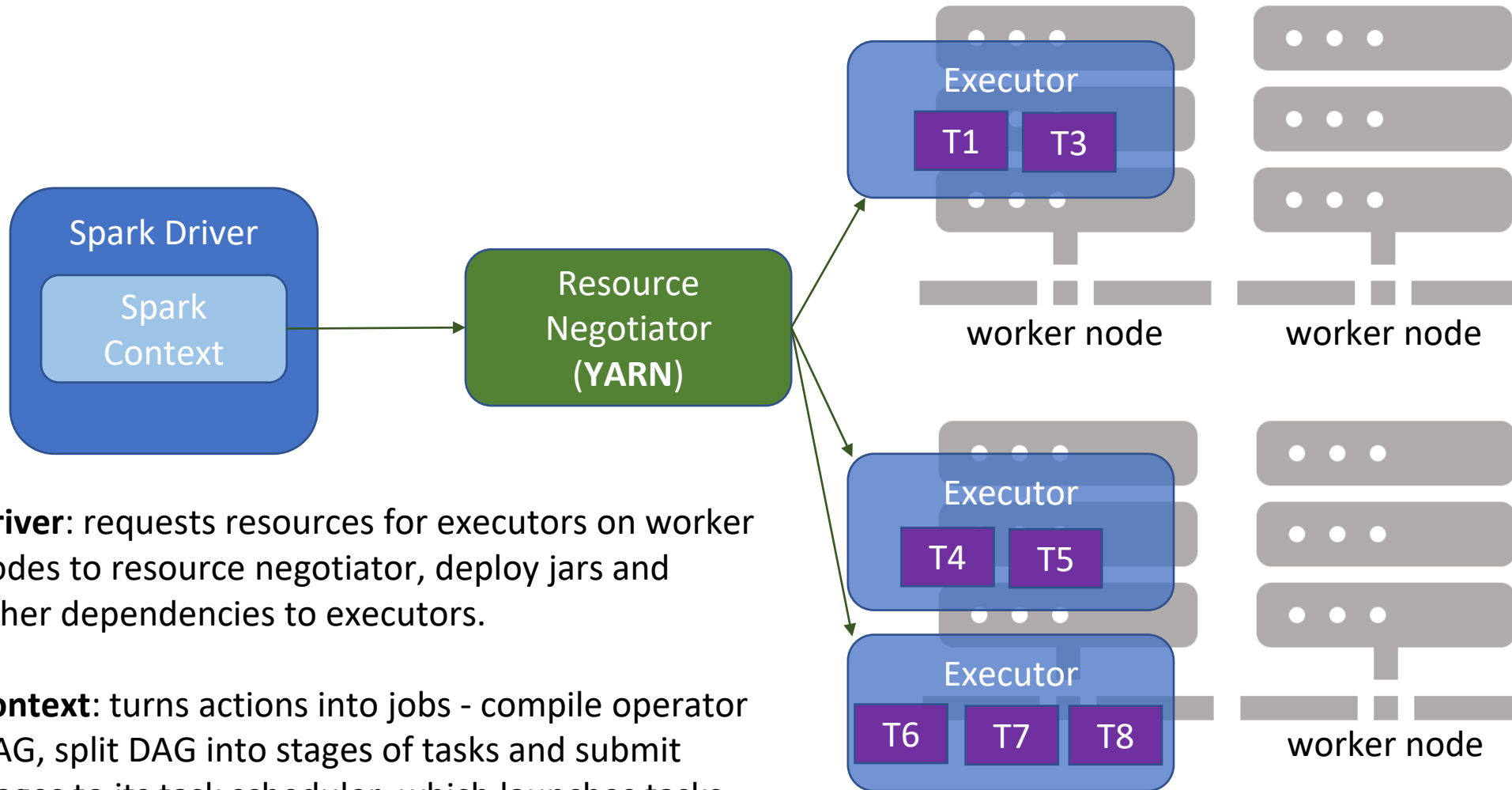


Spark and Big Data

Big Data and/or Big Jobs on big cluster



Spark Architecture - Revisited

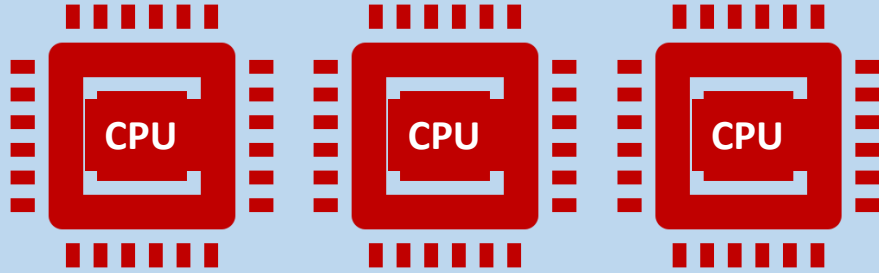


Driver: requests resources for executors on worker nodes to resource negotiator, deploy jars and other dependencies to executors.

Context: turns actions into jobs - compile operator DAG, split DAG into stages of tasks and submit stages to its task scheduler, which launches tasks.

Spark Executor – for heavy lifting

Executor



Number of cores



Memory

Execution 50%

In-memory storage 50%

“Local storage”



Handling Big Data

Time Efficiency

- Minimize data transfer
- locality principle:
 - Run with 'local' data
 - Avoid shuffle operations
- Minimize process time¹
 - Spark/PySpark optimizations
- Tune partitions³



Space Efficiency

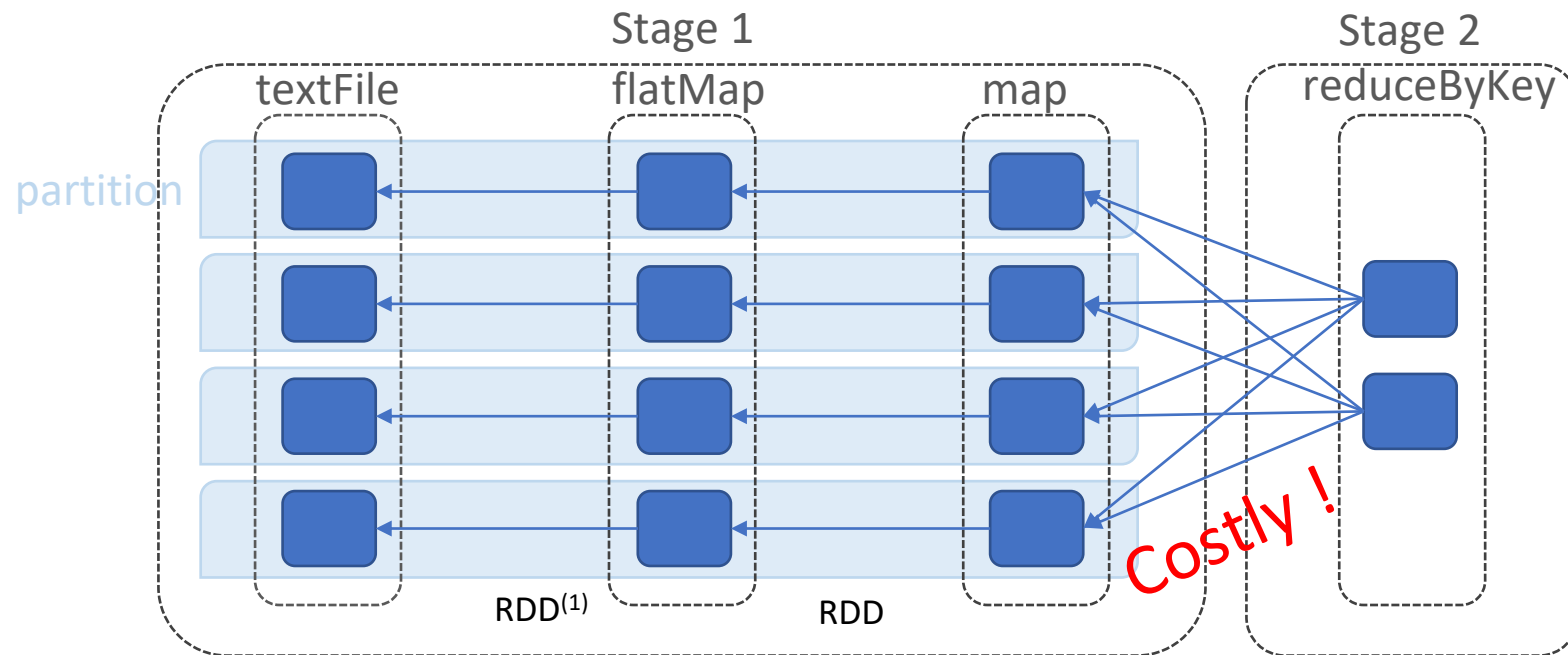
- Be aware of memory usage²
- Minimize amount of data processed
- Tune partitions³



Managing data shuffle and memory – best practices

Shuffle – what is it and how to avoid it

- Shuffle: all-to-all **Costly!**
- When: transfers data among executors
- e.g. join, <reduce/sort/...> ByKey



⁽¹⁾RDDs point to parent

Why is Shuffling costly

- All-to-all data transfers - N^2
- Every transfer involves
 - Serialization data objects $\rightarrow$ bytes $\rightarrow$ data objects
 - Network I/O
 - Disk I/O
 - Operations run on a JVM
 - (often) Heap memory $\rightarrow$ spill to disk $\rightarrow$ more disk I/O

Shuffle Optimizations

- User optimization - avoid shuffle whenever possible
 - Thin the data early
 - Order the data or filter it in the partitions
 - Reorder transformations
 - When order of operations allow it, use operations that reduce the data the most
 - Partial aggregation
- Spark optimization
 - Write intermediate files to disk, so that the data does not get copied to all of the nodes

Memory

- Spark is written in Scala
 - JVM
 - (infamous) Garbage Collector
- When to suspect out of memory issues
 - Poor performance
 - Slow / dead executors (in unpredictable ways)

Memory - diagnostic (expert level)

- How to diagnose
 - set in `spark.executor.extraJavaOptions`
 - `XX:+PrintGCDetails`
 - `XX:+HeapDumpOnOutOf`
 - In linux, use `dmesg` for oom-killer logs
 - Spilling to disk
- If using YARN: take into account the memory constraints - executor could just be killed because they exceed their allocated capacity

Memory optimization

- Avoid spilling to disk and OOM deaths
- Aim for key-value that fit in memory
 - Number of keys – if too large, it may kill the driver (in reduce stage)
 - Number of values per keys – if too large it may kill executors
- Adjust partition size
 - Increase number of partitions
 - Decrease size of partitions
- Minimize shuffle, which is also memory intensive

Partitioning

How is partitioning done? Number of partitions

- Automatically
 - HDFS - partition by HDFS block
 - Spark default partitions

```
sc.defaultParallelism()  
rdd.getNumPartitions()
```

- Programmatically
 - Create RDD with custom number of partitions

```
sc.parallelize(data, num_partitions)
```

- Change the number of partitions in next RDD

```
rdd.repartition(num_partitions)  
rdd.repartitionAndSortWithinPartitions(num_partitions,...)  
rdd.coalesce(num_partitions,shuffle=false)
```

Preferred

Configure partitions

Input - control size

- `spark.default.parallelism()`
- `spark.sql.files.maxPartitionBytes(b)`

Shuffle - control count

- `spark.sql.shuffle.partitions(n)`

Output - control size

- `coalesce(n)` to shrink in an efficient way
- `repartition(n), partitionBy(n, func)` to increase and/or balance
- `df.write.option("maxRecordsPerFile", N)`

Optimum number of partitions

Too few partitions

- less concurrency
- more susceptible to data skew
- increased memory pressure

Too many partitions

- less concurrency
- too much overhead

Right balance

- 2 – 4x number of cores
- Aim for 100ms of executions at least

Repartitioning

- Understand data access pattern: e.g. write once, read many
- Coalesce partitions to size-down file set
- Repartition = another stage = expensive
 - `df.repartition(n)` HashPartitioner
 - `df.repartition(n, [colA, ...])` RangePartitioner
- `df.localCheckpoint + repartition or coalesce`
 - stage barrier

Spark Partitioners

- HashPartitioner
 - Hash code of an object
- RangePartitioner
 - Sortable records
- CustomPartitioner
 - Extend Partitioner
 - Overwrite methods
 - And define the comparison method

```
def numPartitions: Int  
def getPartition(key: Any): Int
```

```
def equals(other: Any): Boolean
```

Spark Partitioners - RDD

- HashPartitioner: [pyspark.RDD.repartition](#)

```
newRdd = oldRdd.repartition(numPartitions)
```

- Custom partitioner: [pyspark.RDD.partitionBy](#)

```
def partitioner(key):  
    if key == 'foo':  
        return 1  
    else:  
        return random.randint(2, numPartitions)  
  
newRdd = oldRdd.repartitionBy(numPartitions, partitioner)
```

PySpark will use: $\text{partitioner}(\text{key}) \bmod \text{numPartitions}$

Spark Partitioners - DataFrame

- HashPartitioner: [pyspark.sql.DataFrame.repartition](#)

- Hash code of an object

```
newDf = oldDf.repartition(numPartitions, "colA", "colB")
```

- RangePartitioner: [pyspark.sql.DataFrame.repartitionByRange](#)

```
newDf = oldDf.repartitionByRange(numPartitions, "colA", "colB")
```

Optimization check list

- Ensure you have enough partitions, but that partitions are not underused
- Optimize placement of functions
 - Use the right functions at the right place
- Choose partitioning functions wisely
 - Aim for well balanced partitions, avoid partition skew
 - Minimize shuffling and amount of data shuffled
- Minimize memory consumption, i.e. memory hungry functions (sorting, groupBy)
 - Can cause OOM

Spark & YARN Troubleshooting

Spark UI

- Use Spark UI to tune tasks/partitions
- See stages, see tasks (#of tasks, shuffle read sizes)
- See failed tasks, tasks that are taking longer
- Collect and read standard output and error logs
- Useful for shared clusters if you are not admin

Spark UI – COM490 Spring 2024

http://iccluster08x.iccluster.epfl.ch:xyz/proxy/application_xyz

History Server

Event log directory: hdfs://iccluster067.iccluster.epfl.ch:8020/user/spark/spark3ApplicationHistory

Last updated: 2024-04-15 14:58:31

Client local time zone: Europe/Zurich

Show entries

Search:

Version	App ID	App Name	Started	Completed	Duration	Spark User	Last Updated	Event Log
3.3.2.3.3.7190.2-1	application_1709909721189_0974	livy-session-142	2024-04-15 14:03:51	2024-04-15 14:51:57	48 min	livy	2024-04-15 14:51:57	Download
3.3.2.3.3.7190.2-1	application_1709909721189_0973	livy-session-141	2024-04-15 13:48:36	2024-04-15 14:51:57	1.1 h	livy	2024-04-15 14:51:57	Download
3.3.2.3.3.7190.2-1	application_1709909721189_0972	livy-session-140	2024-04-15 13:43:36	2024-04-15 13:47:58	4.4 min	livy	2024-04-15 13:47:58	Download
3.3.2.3.3.7190.2-1	application_1709909721189_0940	livy-session-132	2024-04-14 11:35:33	2024-04-14 15:11:19	3.6 h	livy	2024-04-14 15:11:19	Download
3.3.2.3.3.7190.2-1	application_1709909721189_0950	livy-session-139	2024-04-14 14:12:29	2024-04-14 14:13:24	55 s	livy	2024-04-14 14:13:24	Download
3.3.2.3.3.7190.2-1	application_1709909721189_0947	livy-session-138	2024-04-14 13:08:22	2024-04-14 14:10:19	1.0 h	livy	2024-04-14 14:10:19	Download

Spark Jobs (?)

User: livy

Total Uptime: 1.0 h

Scheduling Mode: FIFO

Completed Jobs: 12

Event Timeline

Completed Jobs (12)

Pages: 1

1 Pages. Jump to 1. Show 100 items in a page. Go

Job Id (Job Group)	Description	Submitted	Duration	Stages: Succeeded/Total	Tasks (for all stages): Succeeded/Total
11 (25)	Job group for statement 25 showString at NativeMethodAccessorImpl.java:0	2024/04/11 22:35:19	0.4 s	1/1	425/425
10 (22)	Job group for statement 22 showString at NativeMethodAccessorImpl.java:0	2024/04/11 22:35:18	31 ms	1/1	1/1
9 (21)	Job group for statement 21 showString at NativeMethodAccessorImpl.java:0	2024/04/11 22:35:17	0.4 s	1/1 (1 skipped)	1/1 (425 skipped)
8 (21)	Job group for statement 21 showString at NativeMethodAccessorImpl.java:0	2024/04/11 22:35:16	1 s	1/1	425/425

Applications

Jobs

Stages

Tasks

Details for Stage 9 (Attempt 0)

Resource Profile Id: 0

Total Time Across All Tasks: 15 s

Locality Level Summary: Process local: 425

Input Size / Records: 16.8 / 0

Shuffle Write Size / Records: 103.4 KB / 1880

Associated Job Ids: 0

Summary Metrics for 425 Completed Tasks

Metric	Min	25th percentile	Median	75th percentile	Max
Duration	0.0 ms	11.0 ms	16.0 ms	25.0 ms	0.0 s
GC Time	0.0 ms	0.0 ms	0.0 ms	0.0 ms	13.0 ms
Input Size / Records	16.8 / 0	180.1 KB / 1	720.1 KB / 1	917.6 KB / 1	1.5 MB / 1
Shuffle Write Size / Records	0.0 KB / 0	1.5 KB / 14	5.2 KB / 52	6.8 KB / 71	9.7 KB / 105

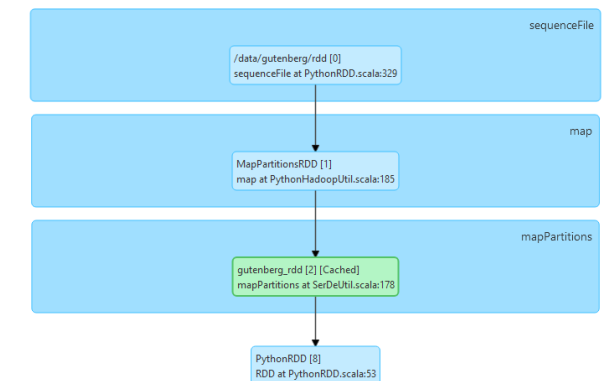
Aggregated Metrics by Executor

Tasks (425)

Show 20 entries

Task ID	Attempt	Status	Locality level	Executor ID	Host	Launch Time	Duration	GC Time	Input Size / Records	Shuffle Write Size / Records	Errors
0	001	0	SUCCESS	PROCESS_LOCAL	8	iccluster071.iccluster.epfl.ch	2024-04-11 22:35:16	0.2 s	959.9 KB / 72	7.2 KB / 78	
1	008	0	SUCCESS	PROCESS_LOCAL	5	iccluster079.iccluster.epfl.ch	2024-04-11 22:35:16	0.2 s	647 KB / 1	5.6 KB / 95	

DAG Visualization



Know the infrastructure (expert)

- Memory and cores per VM / Worker node
- Speed per core
- Network speed
- Disk capacity and disk type (SSDs, remote)
- Goal: improve utilization, aim for 70%

Start your engines

<https://com490-2024.epfl.ch>

<https://dslabgit.datascience.ch/course/2025/module-3c>

(fork and git clone)