

Common Algorithm Building Blocks

Diego Rivera Correa

Slides Owned by Mirek Riedewald



This work is licensed under the Creative Commons Attribution 4.0 International License. To view a copy of this license, visit <http://creativecommons.org/licenses/by/4.0/>.

Key Learning Goals

- Why is the “order inversion” design pattern called that way?
- Be able to write the MapReduce and Spark pseudo-code for all the algorithms in this module.

Introduction

- We discuss approaches and algorithms that frequently appear as building blocks of larger data processing pipelines:
 - Managing state at different granularities
 - Per-record computation
 - Content-based data splitting
 - Grouping and aggregation
 - Duplicate removal
 - Random sampling and shuffling
 - Quantiles
 - Top-k

Managing State at Different Granularities

- In previous modules, we worked with task-local state and discussed two types of **global** state.
- The first type of global state is a mechanism for broadcasting *read-only* data to all tasks of a job.

Managing State at Different Granularities

- The second are global counters and accumulators, which are initialized and read out by the driver, but updated by the tasks of a job.
- We will first review those mechanisms, then introduce the “**order inversion**” design pattern in MapReduce.

Review: Broadcasting Read-Only State

- In Hadoop MapReduce, we can use the job **Context** object to share a small number of constants with all tasks.
- The **file cache** in MapReduce allows broadcasting of larger read-only state via files. In Spark, this is supported by **broadcast variables**.

Review: Global Counters

- Hadoop MapReduce global **counters** and Spark **accumulators** offer limited functionality to avoid the complexity of shared-memory programs.
 - They are initialized and read out by the driver program, i.e., there is no parallelism.
 - They are updated in parallel by the tasks of a job. However, the only supported update operation is to add a (positive or negative) value. Addition is commutative and associative; hence it works correctly independent of the order in which the updates are applied.

Other Options for Global State?

- Should we manage larger objects in the application master, e.g., a central hash-table for counting words?

Other Options for Global State?

- Should we manage larger objects in the application master, e.g., a central hash-table for counting words? No!
 - The master's memory would become a bottleneck.
 - Updates of state require shared-memory style synchronization.

Other Options for Global State?

- How about distributing the large object over many workers?

Other Options for Global State?

- How about distributing the large object over many workers?
 - This addresses the memory bottleneck at the master.
 - If the object can be updated by multiple tasks, then there still is a need for **distributed consensus**.
 - If each task has exclusive access to a partition, then we are back to the per-task shared-nothing-style data management of MapReduce and Spark tasks.

The Order Inversion Design Pattern

- In there anything else in the spectrum between **local** data exclusively accessed by a single task and **global** counters/accumulators?
- Not really, but we will now introduce a solution that was proposed for MapReduce to give operations on fine-grained data partitions access to coarser-grained data needed by multiple small partitions. For reasons explained later, this was called “**order inversion**.”

Example

- Consider a crowd-sourcing project where citizen scientists report birds they observe. For simplicity, assume participants report a (species, color) pair every time they see a bird.
- Our goal is for each species and color to estimate the conditional probability of the color, given the species.
 - The probability of the Northern Cardinal (N.C.) being red is defined as $P(\text{color} = \text{red} \mid \text{species} = \text{N.C.})$.

Example

- We can estimate such probabilities from big data by counting the appropriate quantities. To estimate $P(\text{color} = C \mid \text{species} = S)$, we need:
 - Frequency count $f(S)$, i.e., the number of observations for species S . This is called a **marginal**.
 - Frequency count $f(S, C)$, i.e., the number of observations matching both species S and color C . This is called a **joint event**.
 - E.g., we estimate $P(\text{color} = \text{red} \mid \text{species} = \text{N.C.})$ by dividing the number of **red Northern Cardinal** observations by the **total number of Northern Cardinal** observations (including all colors), i.e., as **$f(\text{N.C., red}) / f(\text{N.C.})$**

Obvious Solution Using “Stripes”

- Both $f(S)$ and $f(S, C)$ need to count per species. Hence the species presents itself as an obvious choice for intermediate key. Starting with this observation, the MapReduce program “falls into place.” For each input record (species S , color C), Map emits the record with species S as the key and color C as the value.

```
// Note: If no combining is used, Map could emit  
// (S, C), i.e., S as the key and C as the value.  
map( observation = (species S, color C) )  
  emit( S, (C, 1) )
```

```
reduce( S , [(C1, n1), (C2, n2),...] )
```

```
// H maps a color to a count
```

```
init hashMap H
```

```
marginal = 0
```

```
for all (C, n) in input list do
```

```
  H[C] += n
```

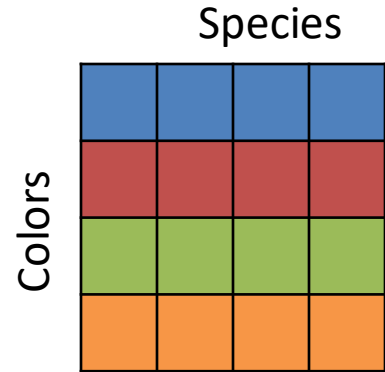
```
  marginal += n
```

```
for all C in H do
```

```
  emit( (S, C), H[C] / marginal )
```

Discussion of the Stripe-Based Approach

- Why do we call this a “stripe”-based approach? Think of a table where each row is indexed by a species and each column is indexed by a color. A table cell, indexed by the combination of species S and color C , contains count $f(S, C)$.
- By choosing the species as the key, Reduce works with an entire row of this table, which pictorially looks like a stripe.



A 4x4 grid representing a table where rows are indexed by Species and columns by Colors. The grid shows four horizontal stripes of color: blue, red, green, and orange.

	Species			
Colors	Blue	Blue	Blue	Blue
	Red	Red	Red	Red
	Green	Green	Green	Green
	Orange	Orange	Orange	Orange

Discussion of the Stripe-Based Approach

- The Stripe, i.e., table row, turns out to be a great fit for relative frequency computation. Each cell in the Stripe has the color frequency for the species; and the sum of these frequencies equals the total for the species.
- Hence all the data needed for computing $f(S, C)/f(S)$ for species S is in the corresponding stripe.

Colors	Species			
	Blue	Blue	Blue	Blue
	Red	Red	Red	Red
	Green	Green	Green	Green
	Orange	Orange	Orange	Orange

Discussion of the Stripe-Based Approach

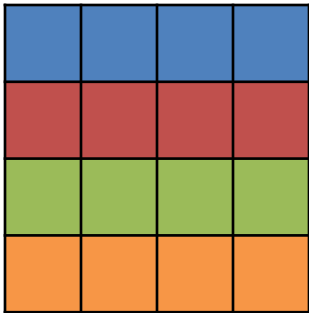
- Is there a drawback to this approach? There are indeed two major limitations:
- Do you think this approach would have memory problems? What about parallelism?

A 4x4 grid illustrating the Stripe-Based Approach. The grid is divided into four horizontal stripes of different colors: blue, red, green, and orange. The label 'Species' is at the top, and 'Colors' is on the left side.

Species			
Blue	Blue	Blue	Blue
Red	Red	Red	Red
Green	Green	Green	Green
Orange	Orange	Orange	Orange

Discussion of the Stripe-Based Approach

- Is there a drawback to this approach?
There are indeed two major limitations:
 - What if data structure H in Reduce exceeds memory size? This would not happen for the colors here, but it might for other problems with more columns.
 - The degree of parallelism of the Reducer workload is limited by the number of different species. What if we have more machines than species?



A 4x4 grid of colored squares. The columns are labeled 'Species' and the rows are labeled 'Colors'. The grid contains four distinct colors: blue, red, green, and orange, each appearing in a single row across all four columns.

	Species			
Colors	Blue	Blue	Blue	Blue
	Red	Red	Red	Red
	Green	Green	Green	Green
	Orange	Orange	Orange	Orange

New Attempt Using “Pairs”

- We could address the problems of the Stripe-based approach by splitting the Reducer work into smaller units. Unfortunately, any smaller unit would miss some of the joint events needed for computing $f(S)$.
- To see why, consider a program that uses both species and color as the intermediate key. For input record (species S , color C), Map emits $((S, C), 1)$.

New Attempt Using “Pairs”

- For key (S, C), reduce() computes $f(S, C)$, the frequency count of that species-color combination. Virtually no memory is used, and the fine granularity enables up to $\#species \times \#colors$ different Reduce function calls.
- So, does this approach have any drawbacks?

New Attempt Using “Pairs”

- So, does this approach have any drawbacks?
Yes!
 - How can it compute the marginal $f(S)$? The `Reduce()` call needs $f(S, \text{color})$ for **all** colors for species S .
 - This could be addressed by running another simple MapReduce program to pre-compute $f(S)$. However, this approach seems wasteful for Big Data as it reads the input data twice—once for computing the $f(S)$ and then again for computing the $f(S, C)$.

New Attempt Using “Pairs”

- Can we get the best of both worlds, i.e., the one-pass efficiency of Stripes and the small memory footprint and more fine-grained work partitioning of Pairs?

Fixing the Pairs-Based Approach, Attempt 1

- Notice that if keys $(S, C1)$ and $(S, C2)$ are assigned to different Reducers, then no Reducer has access to all data needed for computing $f(S)$. Hence, we must make sure that *all keys $(S, \text{any color})$ for species S end up in the same Reduce task*.
 - We already know how to achieve this by defining a custom **Partitioner** that assigns a key (S, C) to a Reducer solely based on S , ignoring the value of C .

Fixing the Pairs-Based Approach, Attempt 1

- While the custom Partitioner guarantees that all records for species *S* will be processed in the same Reduce task, there still is a separate Reduce call for each species-color combination.

Challenge Question 1

- How can we compute $f(S)$ when each Reduce call only works with a single color for species S ?

Challenge Question 1: Answers

- How can we compute $f(S)$ when each Reduce call only works with a single color for species S ?
- Possible answer 1: The individual Reduce calls for keys $(S, C1)$, $(S, C2)$, etc could be turned into a single Reduce call for species S by using a **grouping comparator**. (Review the secondary sort design pattern.)

Challenge Question 1: Answers

- Possible answer 2: State can be maintained across the different Reduce calls for keys (S, C1), (S, C2), etc to keep track of $f(S, C)$ for all colors C of species S.
- More precisely, both the marginal and a hashmap H that stores the frequency for each color could be defined at the Reducer **class level**, letting each Reduce call for (S, C) update them accordingly. (Review the in-mapper combining design pattern.)

Challenge Question 2

- Does either of these approaches (grouping comparator, in-Reducer combining) really give us a better solution than the Stripes approach?

Challenge Question 2: Answer

- Does either of these approaches (grouping comparator, in-Reducer combining) really give us a better solution than the Stripes approach?
- No. These “improved” Pairs-based approaches have the same drawbacks as the Stripes-based approach. By using a custom Partitioner that only considers the species, Reduce task **granularity** is back at the species level, like for Stripes.

Challenge Question 2: Answer

- Similarly, since $f(S)$ is computed together with the $f(S, C)$ frequencies, all separate color frequencies for species S must be kept **until the last record for species S is processed**.
- Hence the memory footprint is not smaller than for Stripes either. Essentially the attempt at improving the Pairs-based approach made it “simulate” the Stripe-based approach!

Fixing the Pairs-Based Approach, Attempt 2

- The failed attempts so far had tried to compute $f(S)$ for species S **concurrently** with the different $f(S, C)$ for that same species.
- This forced us to
 - (1) send all Map output records for species S to the same Reducer
 - (2) keep the counts for all $f(S, C)$ around until the very last record for species S was processed by the Reducer, i.e., the moment $f(S)$ would finally be known.

Fixing the Pairs-Based Approach, Attempt 2

- If the program had known $f(S)$ from the beginning, it could have been “broadcast” to the Reduce calls for keys (S, C) , for all colors C .
 - This is similar to global state, but for a smaller scope—here per individual species.
- How can we achieve this sharing of coarser-grained data across functions working on more fine-grained data?

State Sharing with Global Counters

- The program below assumes that there is a global counter for each species. There are two problems with this idea...

```
map( observation: (species S, color C) )  
  // Update global counter for species S.  
  // This counter must have been created  
  // beforehand by the driver.  
  marginal_S.add(1)  
  
emit( (S, C), 1 )
```

```
reduce( (S, C), [n1, n2,...] )  
  frequency = 0  
  
  for all n in input list do  
    frequency += n  
  
  // Here we assume that the counter can be read  
  // in a Reduce task.  
  emit( (S, C), frequency / marginal_S )
```

State Sharing with Global Counters

- The program below assumes that there is a global counter for each species. There are two problems with this idea...
 - All species must be known in advance in the driver, before any task accesses the input. We cannot dynamically create new counters in a task.
 - Tasks cannot read the counter value (during same job), therefore Reduce has no access to the marginal.

```
map( observation: (species S, color C) )  
  // Update global counter for species S.  
  // This counter must have been created  
  // beforehand by the driver.  
  marginal_S.add(1)  
  
emit( (S, C), 1 )
```

```
reduce( (S, C), [n1, n2,...] )  
  frequency = 0  
  
  for all n in input list do  
    frequency += n  
  
  // Here we assume that the counter can be read  
  // in a Reduce task.  
  emit( (S, C), frequency / marginal_S )
```

Solution Using Order Inversion

- Like Pairs, the program below uses intermediate key (species, color). In addition to ((S, C), 1), Map emits ((S, dummy), 1) for computing $f(S)$. We exploit that Reduce calls in the same Reducer are executed in key order.

```
map( ..., observation: (species S, color C) ) {  
    emit( (S, dummy), 1 )  
    emit( (S, C), 1 )  
}
```

Partitioner: partition by species

Key comparator for (species, color):

- Sort by species first, then by color
- Make sure color “dummy” comes before all real colors

```
Class Reducer {  
    marginal    // per-task state  
  
    reduce( (S, C), [n1, n2,...] ) {  
        if C = dummy {           // Compute marginal f(S)  
            marginal = 0  
            for all n in input list do marginal += n  
        } else {                 // Real color: compute f(S, C)  
            colorCnt = 0  
            for all n in input list do colorCnt += n  
            emit( (S, C), colorCnt / marginal )  
        }  
    }  
}
```

Solution Using Order Inversion

- The Reducer class needs a task-level variable to keep track of marginal $f(S)$. Since the Reduce call for the dummy color happens first, $f(S)$ will be known before the Reduce call $f(S, C)$ for any color C .

```
map( ..., observation: (species S, color C) ) {  
    emit( (S, dummy), 1 )  
    emit( (S, C), 1 )  
}
```

Partitioner: partition by species

Key comparator for (species, color):

- Sort by species first, then by color
- Make sure color “dummy” comes before all real colors

```
Class Reducer {  
    marginal    // per-task state  
  
    reduce( (S, C), [n1, n2,...] ) {  
        if C = dummy {           // Compute marginal  $f(S)$   
            marginal = 0  
            for all n in input list do marginal += n  
        } else {                 // Real color: compute  $f(S, C)$   
            colorCnt = 0  
            for all n in input list do colorCnt += n  
            emit( (S, C), colorCnt / marginal )  
        }  
    }  
}
```

Solution Using Order Inversion

- The Reduce task now has a constant (i.e., independent of the number of species and colors) memory footprint. It needs the marginal for a single species and a single counter for the currently processed species-color combination.

```
map( ..., observation: (species S, color C) ) {  
  emit( (S, dummy), 1 )  
  emit( (S, C), 1 )  
}
```

Partitioner: partition by species

Key comparator for (species, color):

- Sort by species first, then by color
- Make sure color “dummy” comes before all real colors

```
Class Reducer {  
  marginal    // per-task state  
  
  reduce( (S, C), [n1, n2,...] ) {  
    if C = dummy {           // Compute marginal f(S)  
      marginal = 0  
      for all n in input list do marginal += n  
    } else {                 // Real color: compute f(S, C)  
      colorCnt = 0  
      for all n in input list do colorCnt += n  
      emit( (S, C), colorCnt / marginal )  
    }  
  }  
}
```

Discussion

- How does this new solution compare to the previous attempts?

Discussion

- How does this new solution compare to the previous attempts?
 - Pro: It reads the input only once, like the Stripes approach.
 - Pro: It requires virtually no heap memory, like the initial Pairs approach.
 - Potential con: Map duplicates every input record, therefore two copies of the input data are transferred from Mappers to Reducers.
 - Same: Reduce granularity still is limited by the number of species.

Discussion

- Interestingly, this approach *in principle* supports a finer Reduce-task granularity at the cost of more data replication in the Map phase. The code on the next slide illustrates this idea.

Discussion

- By producing k replicas in Map, each for a different dummy color, the keys $(S, C1)$, $(S, C2)$, etc. can be distributed over k different Reduce tasks.
- The Partitioner must ensure that each of these tasks receives one of the (S, dummy_i) keys as well, so that it can compute the marginal.

Discussion

- Idea: Instead of sending all colors of species S to one reducer, split the colors across k reducers. To make that work: Each reducer also receives a dummy copy that lets it compute $f(S)$.
- Map emits: $(S, \text{dummy}_1), (S, \text{dummy}_2), \dots, (S, \text{dummy}_k), (S, \text{realColor})$
 - Partitioner: $\text{dummy}_i \rightarrow \text{reducer } i$. $\text{realColor} \rightarrow \text{one of the } k \text{ reducers, based on } \text{hash}(\text{color})$
- Reducer: First computes $f(S)$ from its dummy key. Then computes $f(S, C) / f(S)$ for its assigned colors.
- Tradeoff: More parallelism, but more replicated map output.

Discussion

- Species: Cardinal. Original order inversion:
one reducer gets everything:
f(Cardinal), red, blue, green, yellow
- Multiple partitions:
Reducer 1: dummy1 + red, yellow
Reducer 2: dummy2 + blue
Reducer 3: dummy3 + green, black

Each reducer computes the same f(Cardinal),
then handles only some colors.

Multiple Partitions (on Color)

```
map( observation: (species S, color C) ) {  
  emit( (S, dummy1), 1 )  
  emit( (S, dummy2), 1 )  
  ...  
  emit( (S, dummyk), 1 )  
  emit( (S, C), 1 )  
}
```

Partitioner: partition by species and color, distributing the colors for species S over k Reduce tasks and making sure each of these Reduce tasks receives exactly one of the dummy_i colors for that species.

Key comparator for (species, color):

- Sort by species first
- Make sure color “dummy” comes before all real colors

```
Class Reducer {  
  marginal  
  
  reduce( (S, C), [n1, n2,...] ) {  
  
    if C = dummy {  
      // Compute marginal f(S)  
      marginal = 0  
      for all n in input list do  
        marginal += n  
    } else {  
      // Real color, hence compute f(S, C)  
      colorCnt = 0  
      for all n in input list do  
        colorCnt += n  
  
      emit( (S, C), colorCnt / marginal )  
    }  
  }  
}
```

Order Inversion Design Pattern Summary

- This design pattern for computing state at different granularities and for controlling the order in which this state is computed, is called “order inversion.”
- In what sense does it invert order? We can compute $f(S)$ as the sum of the $f(S, C)$, summing over all colors C . Hence the “natural” order of computation would be to obtain all the $f(S, C)$ for species S first and then add them up to obtain $f(S)$.

Order Inversion Design Pattern Summary

- Without order inversion, the programmer would need either (1) larger and more complex data structures to bring the right data together (e.g., hashmap H for a Stripe) or (2) more MapReduce jobs to compute the intermediate results.
- Order inversion turns a synchronization problem into an ordering problem.

Order Inversion Design Pattern Summary

- Tradeoff: This approach enables the use of simpler data structures and less Reducer memory, without requiring additional MapReduce phases.
- It can achieve more fine-grained partitioning at the cost of data replication. A Combiner or in-Mapper combining may soften the negative performance impact of this data replication.

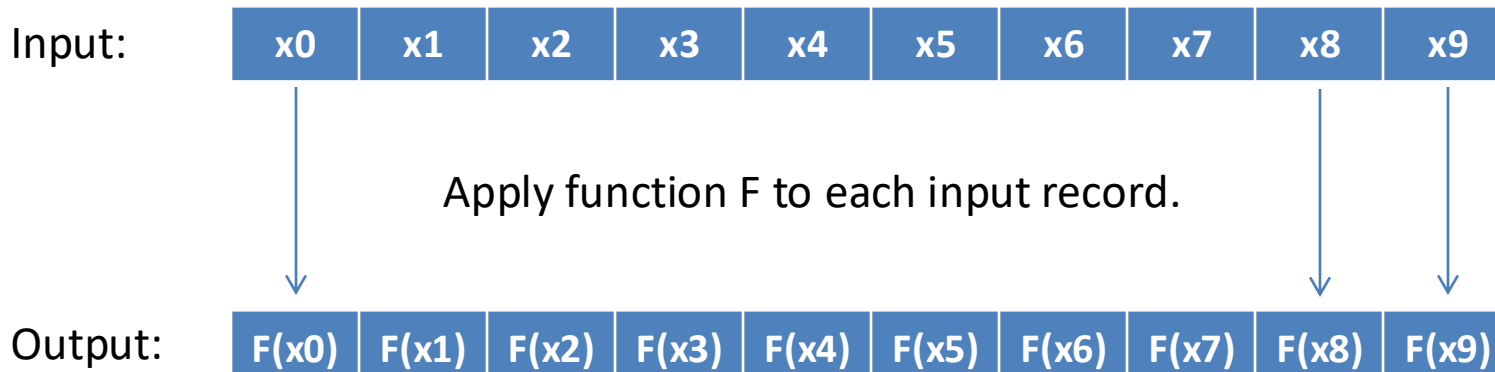
Next, we will explore a variety of useful “small” algorithms.

Per-Record Computation

- Per-record computation is easy to parallelize and fairly common:
 - The relational selection operator filters out records using a predicate. For example, from a dataset of flights, it can return all flights into Boston.
 - The relational projection operator removes fields from a record. For example, it could remove the flight arrival time field if it is not needed for the analysis.

Per-Record Computation

- **Task:** Given function $F()$, transform each input record x to $F(x)$.
- **Solution:** Each task applies $F()$ to its input records one-by-one. No shuffling is needed. For operators that filter records, e.g., selection, $F(x)$ conceptually returns NULL if x is filtered out.



Implementations

DBMS (SQL): Assume the input relation R has schema (A1, A2, A3)

Generic solution: SELECT F(A1, A2, A3) FROM R

Selection: SELECT * FROM R WHERE F(A1, A2, A3)

MapReduce:

// Generic and projection

```
map( ..., x )  
  emit( F(x) )
```

// Selection

```
map( ..., x )  
  if F(x) then emit( x )
```

Spark:

```
myRDD.map(x => F(x))           // generic  
myRDD.filter( F(x) )         // selection
```

```
myDS.map(x => F(x))           // generic  
myDS.filter( F(x) )         // selection  
myDS.select("A1")           // projection
```

Discussion

- The MapReduce implementation is a Map-only job. To specify this, the number of Reducers must be set to zero.
 - Look at the grep program from the Miner/Shook book on MapReduce design patterns:
<http://khoury.northeastern.edu/home/mirek/code/DistributedGrep.java>
- The MapReduce, Spark, and generic SQL solution read the **entire input**. If every input record must be processed anyway, then this is the best one can do.

Content-Based Data Splitting

- Consider product reviews by users of an online shopping site, stored as records with schema (userID, isPreferred, productID, productCategory, review). An analyst might be interested in exploring reviews by preferred users separately from those by regular users.
- Similarly, a product-centric exploration might require training of separate data mining models for different product categories.

Content-Based Data Splitting

- **Task:** Partition the input into p separate subsets based on properties of the input records.
- **Solution:** Like per-record computation, each task goes through its input records one-by-one, sending the record to the desired output channel. No data shuffling is needed.
- This approach can also be used for problems where some input records are assigned to zero or more than one partitions.

Input:



Output:



Implementation in MapReduce

- A simple Map-only job can solve this task. It uses the `MultipleOutputs` class to write to different output files. Each of them can store records of a different type, e.g., one might store text data, another pairs of integer numbers.
- The Map function determines the partition the input record belongs to, then emits it to the appropriate output file using `MultipleOutputs.write()`.

Implementation in MapReduce

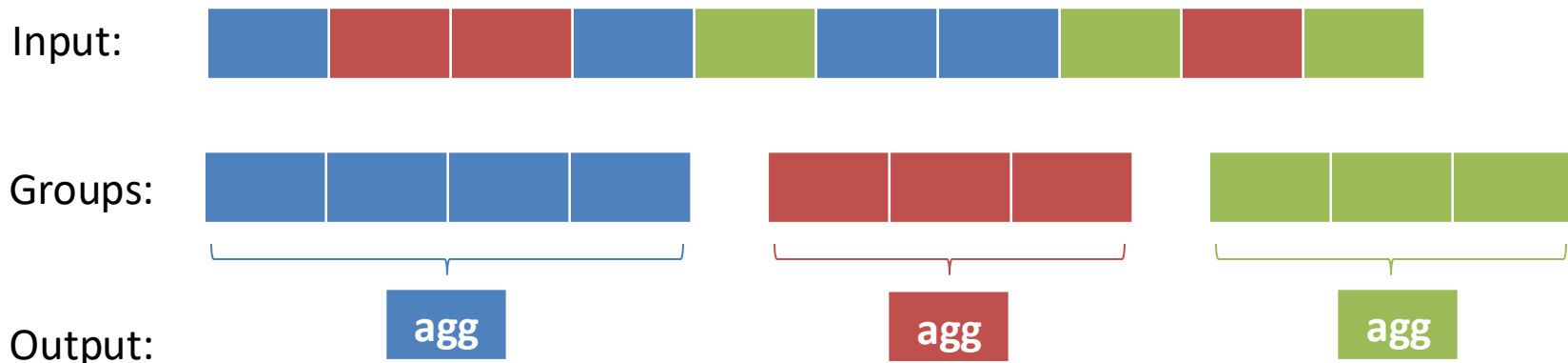
- For m Map tasks and p output partitions, this program will generate $m \times p$ output files. (Each Map task writes to its own set of p output files!)
- If these files are too small, they can be concatenated using HDFS file-system commands.

Grouping and Aggregation

- Data can be partitioned without shuffling only if the possible partitions are known in advance.
- For aggregation and when the partitions are not known, shuffling is needed. We have seen this for Word Count.

Grouping and Aggregation

- **Task:** Partition the input on a key and compute an aggregate of the values in each partition.
- **Solution:** The first round filters records and removes irrelevant attributes. Then the shuffle phase ensures that all records with the same key end up together, so that the next stage can perform the per-group aggregation.



Implementations

DBMS (SQL): Assume the input relation R has schema (Key, Val)

```
SELECT Key, myAGG( Val ) FROM R GROUP BY Key
```

MapReduce:

```
map( key k, val v )  
  emit( k, v )
```

```
reduce( key, [val1, val2,...] )  
  compute agg = myAGG(val1, val2,...)  
  emit( key, agg)
```

Spark:

```
myPairRDD.aggregateByKey( aggFunction )  
  
myDS.groupBy("Key").agg( aggFunction )
```

Grouping RDDs in Spark

- `groupByKey` transformation: for a pair RDD of type (K, V) , it returns an RDD of type $(K, \text{Iterable}[V])$ —one group per key.
- When computing aggregates, it usually is more memory-efficient to use `aggregateByKey`, `reduceByKey`, or `foldByKey` instead.
 - These combine data in the same partition!

Global Aggregation: Standard Approach

- Sometimes one would like to compute a single “global” aggregate for the entire input, e.g., the average delay over all flights. This can be done like grouping-and-aggregation.
- In MapReduce, each Mapper computes the per-split aggregate, using in-Mapper combining or a Combiner. Then these values are associated with a “dummy” key and aggregated by a single Reduce call.

Global Aggregation: Standard Approach

- In MapReduce, each Mapper computes the per-split aggregate, using in-Mapper combining or a Combiner. Then these values are associated with a “dummy” key and aggregated by a single Reduce call.

DBMS (SQL): Assume the input relation R has attribute A
SELECT myAGG(A) FROM R

MapReduce:

```
map( key k, val v )  
  emit( dummy, v )  
  
reduce( dummy, [val1, val2,...] )  
  emit( myAGG(val1, val2,...) )
```

Spark:

```
myRDD.aggregate( ..., aggFunction,... )  
  
myDS.agg( aggFunction )
```

Global Aggregation with Global Counters/Accumulators

- The standard approach we just discussed requires shuffling. Can we avoid this, i.e., can we run the equivalent of a Map-only job?
- Based on what you have seen so far, this seems impossible, because a task aggregating a data partition cannot compute the global aggregate.

Global Aggregation with Global Counters/Accumulators

- Simple aggregates like sum and count can be computed with a Map-only job by exploiting Hadoop's **global counter** feature. It is available through the Counter class.
- Global counter variables can be defined in MapReduce user code and any task can increment them or set their value. No matter which task updates a counter, in the end it will reflect the updates performed by all completed tasks. In Spark, the corresponding feature are **Accumulators**.

Duplicate Removal

- Duplicate removal eliminates repeat occurrences of records in an input file. It is a special case of grouping-and-aggregation: identical records form groups and from each group exactly one representative is output.

Duplicate Removal

- **Task:** Eliminate all duplicates from the input.
- **Solution:** Group by the record content, then emit a single representative for each group.

Implementation

DBMS (SQL):

```
SELECT DISTINCT * FROM R
```

```
SELECT DISTINCT column1, column2,... FROM R
```

MapReduce:

```
map( key k, val v )  
  emit( (k, v), NULL )
```

```
reduce( (k, v), [NULL, NULL,...] )  
  emit( k, v )
```

Spark:

```
myRDD.distinct()
```

```
myDS.dropDuplicates()
```

```
myDS.dropDuplicates( column1, column2,...)
```

Real Code

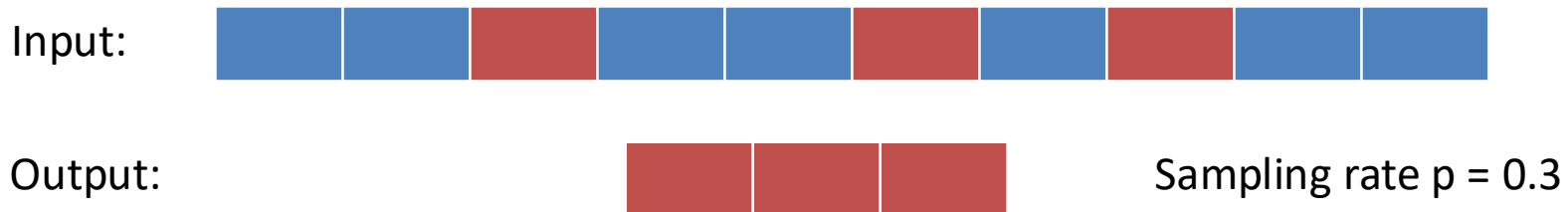
- Check out the example code from the Miner/Shook book at <http://khoury.northeastern.edu/home/mirek/code/DistinctUserDriver.java>
- The program finds all distinct user IDs. Since the input records are in XML format, they are parsed to access the user ID element.

Random Sampling

- Random sampling is crucially important for big-data analysis. Working with a random sample lowers computational cost while often still producing a good approximation of the desired result.
- **Task:** Sample a fraction of approximately p , $0.0 < p \leq 1.0$, of the input records *uniformly at random*.
 - This means each input record should have the same probability p of being selected for the output.

Random Sampling

- **Solution:** Each data partition can be sampled independently with sampling rate p . A pseudorandom-number generator determines if the input record will be emitted: If the generator produces a floating-point number rnd in the range $0.0 \leq \text{rnd} < 1.0$, then the input record is emitted if and only if $\text{rnd} < p$. No shuffling is needed.



Implementation

DBMS (SQL):

```
SELECT * FROM R WHERE random() <= p
```

MapReduce:

```
map( key k, val v )  
  // Assume random() produces a pseudorandom  
  // number n in the range  $0.0 \leq n < 1.0$   
  if random() < p then emit( k, v )
```

Spark:

```
// Sample without replacement: set first  
// argument to FALSE  
myRDD.sample( FALSE, p )  
  
myDS.sample( p )  
myDS.sample( FALSE, p )
```

Real Code

- Check out the example code from the Miner/Shook book at <http://khoury.northeastern.edu/home/mirek/code/SimpleRandomSampling.java>
- Notice how this program uses a single pseudorandom-number generator per Map task.

Random Shuffling

- Random shuffling, i.e., arranging records in a file in a random order, can be useful in many situations. After a file is randomly shuffled, each block will contain a random sample of records.
- Hence one can obtain a random sample of exactly k records by simply reading the first k records from the shuffled file.

Random Shuffling

- **Task:** Randomly shuffle a given input file. Each input record should have the same probability of ending up in any position in the shuffled file.
- **Solution:** We sort the input by a (pseudo-) random key, which is assigned by tasks in the first round of computation.

Implementation

DBMS (SQL):

```
SELECT * FROM R ORDER BY random()
```

MapReduce:

```
map( key k, val v )
```

```
// To avoid duplicate keys, chose the random  
// number from a large domain, e.g., floating  
// point numbers between 0.0 and 1.0.  
emit( random(), (k, v) )
```

```
// If none of the random keys are identical, then  
// the input list contains only a single record
```

```
reduce( rnd, [(k1, v1), (k2, v2),...] )
```

```
for each (k, v) in input list  
  emit( k, v )
```

Spark:

```
myRDD.map( x => (rand(), x) )  
      .sortByKey()
```

```
myRDD.map( x => (rand(), x) )  
      .groupByKey
```

```
myDS.sort( rand() )
```

Real Code

- Check out the example code from the Miner/Shook book at <http://khoury.northeastern.edu/home/mirek/code/AnonymizeDriver.java>
- This program assigns a random integer as the key.

Approximate Quantiles

- Quantiles such as the **median** provide important information about a data distribution. For instance, the median housing price is a measure of wealth for a neighborhoods.
- Furthermore, range-partitioning based on quantiles results in balanced load distribution because the number of records between consecutive quantiles is the same (unless there are many duplicates).

Approximate Quantiles

- **Task:** Find approximate quantiles or percentiles.
- **Solution:** The main idea is to randomly sample input records in the first round. Then all sampled records are sorted by a single task in round 2. The approximate quantiles are collected from the sorted sample.
 - A good choice is to create a sample that is large, but still fits into the memory of the round-2 task.

Records sampled:



Median selected
from sample:



Challenge Question

- What are the drawbacks of a sample that is too small or too large?

Challenge Question: Answer

- What are the drawbacks of a sample that is too small or too large?
 - The smaller the sample, the less data is available for determining the quantiles. This leads to poorer approximation quality.
 - On the other hand, if sample size exceeds the amount of memory, sorting it would be significantly more expensive.

Implementation in MapReduce

- The algorithm for approximate quantiles is shown below. Would **secondary sort** be helpful here? It eliminates the need for sorting in user code.

```
map( key k, val v )  
  // Assume random() produces a  
  // pseudorandom number n in the  
  // range  $0.0 \leq n < 1.0$   
  if random() < p then emit( dummy, (k, v) )
```

```
reduce( dummy, [(k1, v1), (k2, v2),...] ) {  
  copy all records from the input list into array A  
  sort array A  
  
  for each quantile position i in sorted array A  
    emit( A[i] )  
}
```

Top-K Records

- In addition to quantiles, one can gain valuable information about a big dataset from the K most **important** records.
- This could be the top-K largest (or smallest) records based on some field or attribute, e.g., the most active users, people who spend the most money or time, or users with the most friendship links.

Top-K Records

- **Task:** Find the K records that have the largest value of some attribute of interest.
- **Solutions:**
 - For large K, sort the input and then select the first K records.
 - For small K, sorting can be avoided. The main idea is to scan the input only once. In the first round, each task finds the **local** top-K in its partition. All local top-K are sent to a **single task** that merges them into the final result. Correctness is guaranteed, because if a record is in the global top-K, it must be in one of the local top-K.

MapReduce Implementation

- The algorithm below could also use the secondary sort design pattern to simplify the reduce function. With secondary sort, the reduce function simply picks up the first K records in the input list.

```
Class Mapper {  
    localTopK  
  
    setup() { init localTopK }  
  
    map( v )  
        if (v ranks above the last value in localTopK)  
            // Adding v must also evict the now  
            // (K+1)-st value from localTopK  
            localTopK.add( v )  
  
    cleanup()  
        for each v in localTopK emit( dummy, v )  
}
```

```
reduce( dummy, [v1, v2,...] )  
    init globalTopK  
  
    for each value v in input list  
        if (v ranks above the last value in globalTopK)  
            // Adding v must also evict the now  
            // (K+1)-st record from globalTopK  
            globalTopK.add( v )  
  
    for each value v in globalTopK  
        emit( v )  
}
```

Real Code

- Check out the example code from the Miner/Shook book at <http://khoury.northeastern.edu/home/mirek/code/TopTenDriver.java>
- This program finds the top-10 users with the greatest reputation. Since the input records are in XML format, they are parsed to access the reputation element.

Implementation in Spark and DBMS

DBMS (SQL): Assume the input relation R has attribute A

```
SELECT * FROM R  
ORDER BY A  
FETCH FIRST K ROWS ONLY
```

Spark:

```
// The RDD needs to have an implicit Ordering for the type of objects stored  
myRDD.top( K )  
myRDD.takeOrdered( K )  
  
myDS.sort( A ).head( K )
```

More Info about Top-K in Spark

- `takeOrdered(k)` and `top(k)` actions return the k first and last elements of the RDD, respectively. Ordering is determined by an implicit Ordering object defined in scope.
- When applied to pair RDDs of type (K, V), order would not be based on key, but on (K, V) tuples. To change this, one needs to have an implicitly defined Ordering[(K, V)] object in scope. This is the case for simple key and value types by default.
- Implementation in Spark corresponds to the top-K MapReduce algorithm: it takes the top-K elements from each partition, then merges them into a single top-K list.
 - Note that the final result will end up in the driver's memory.

Summary

- Big data analysis often involves partitioning, sampling, and summarizing of data. Parallel solutions in MapReduce and Spark are comparably straightforward, sometimes not requiring any data shuffling.
- In most cases, design patterns introduced previously, e.g., in-Mapper combining and secondary sort, can significantly improve performance or reduce coding effort.

References

- M. Greenwald and S. Khanna. Space-efficient Online Computation of Quantile Summaries. In Proc. ACM SIGMOD Int. Conf. on Managment of Data, pages 58-66, 2001
 - https://scholar.google.com/scholar?cluster=6184540005789557130&hl=en&as_sdt=0,22