

Distributed Programming with MapReduce

Mike Hicks
University of Maryland

Recall: Kinds of parallelism

- Data parallelism
 - The same task run on different data in parallel
- Task parallelism
 - Different tasks running on the same data
- Hybrid data/task parallelism
 - A parallel pipeline of tasks, each of which might be data parallel

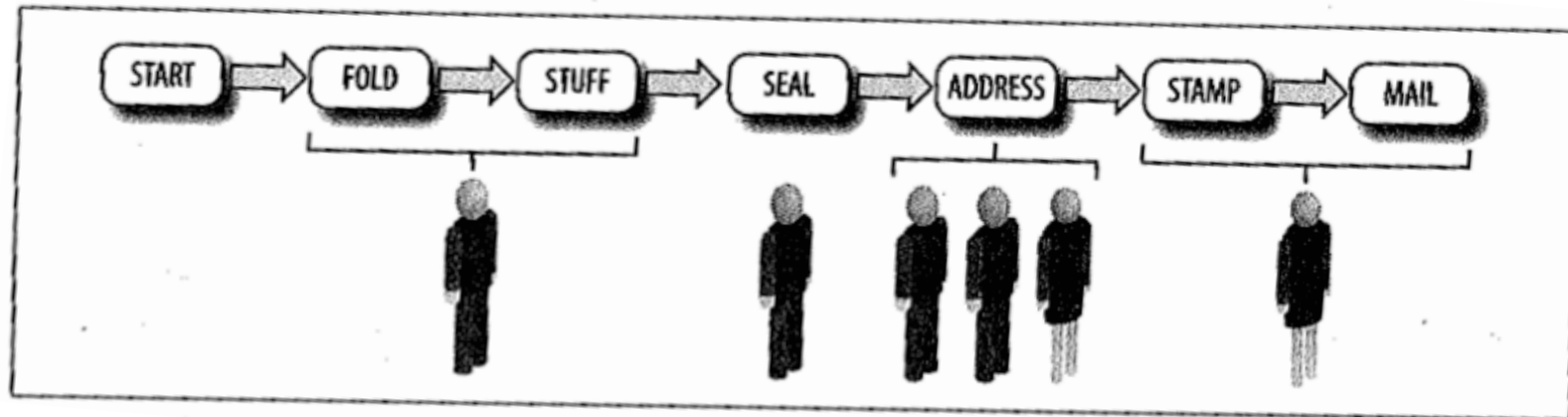
Pipeline parallelism

- Output of one task is the input to the next
 - Each task can run in parallel
 - Throughput impacted by the longest-latency element in the pipeline



Pipeline load balancing

- Assign more than one computational process to each task
 - Combines data- and pipeline- parallelism



MapReduce: Programming the Pipeline

- Pattern inspired by Lisp, ML, etc.
 - Many problems can be phrased this way
- Results in clean code
 - Easy to program/debug/maintain
 - Simple programming model
 - Nice retry/failure semantics
 - Efficient and portable
 - Easy to distribute across nodes

Map and Reduce in Lisp (Scheme)

- (map *f list*)
- (map square '(1 2 3 4))
- (1 4 9 16)
- (reduce + '(1 4 9 16) 0)
- (+ 1 (+ 4 (+ 9 (+ 16 0))))
- 30
- (reduce + (map square '(1 2 3 4)) 0)

Unary operator

Binary operator

MapReduce a la Google (and Yahoo! Hadoop)

- `map(key, val)` is run on each item in set
 - emits new-key / new-val pairs
 - type is $k1*v1 \rightarrow (k2*v2)$ list
- `reduce(key, vals)` is run for each unique key emitted by `map()`
 - emits final output
 - type is $k2*(v2 \text{ list}) \rightarrow v2 \text{ list}$

Count words in docs

- Input consists of (url, contents) pairs
- map(key=url, val=contents):
 - For each word w in contents, emit (w , “1”)
- reduce(key=word, values=uniq_counts):
 - Sum all “1”s in values list
 - Emit result “(word, sum)”

Count, Illustrated

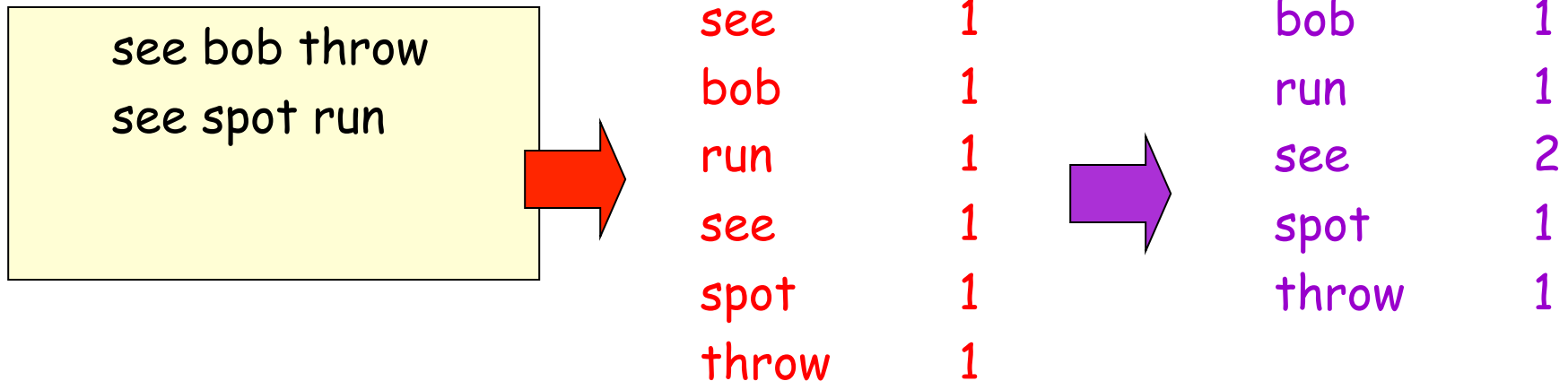
map(key=url, val=contents):

For each word w in contents, emit (w , "1")

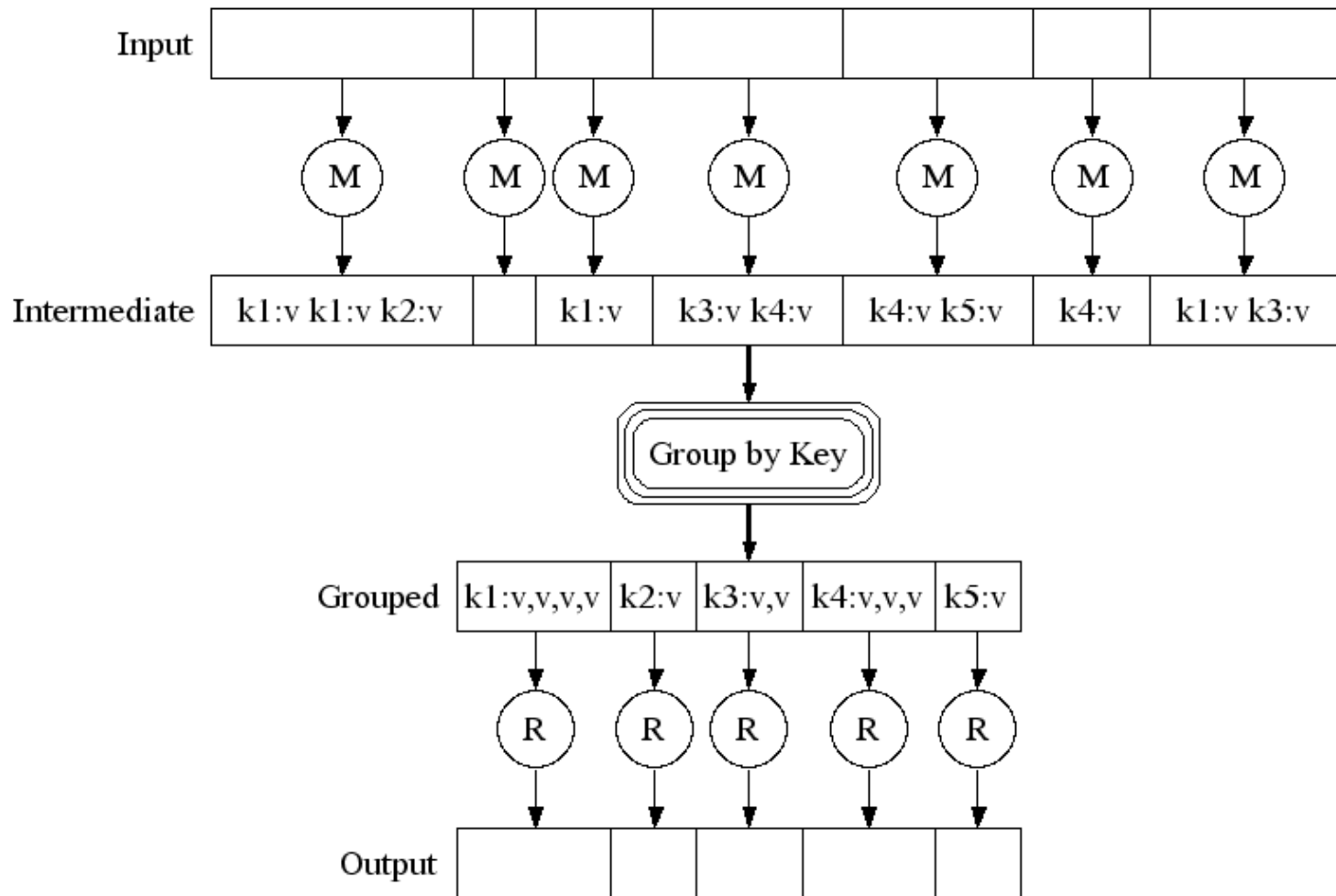
reduce(key=word, values=uniq_counts):

Sum all "1"s in values list

Emit result "(word, sum)"



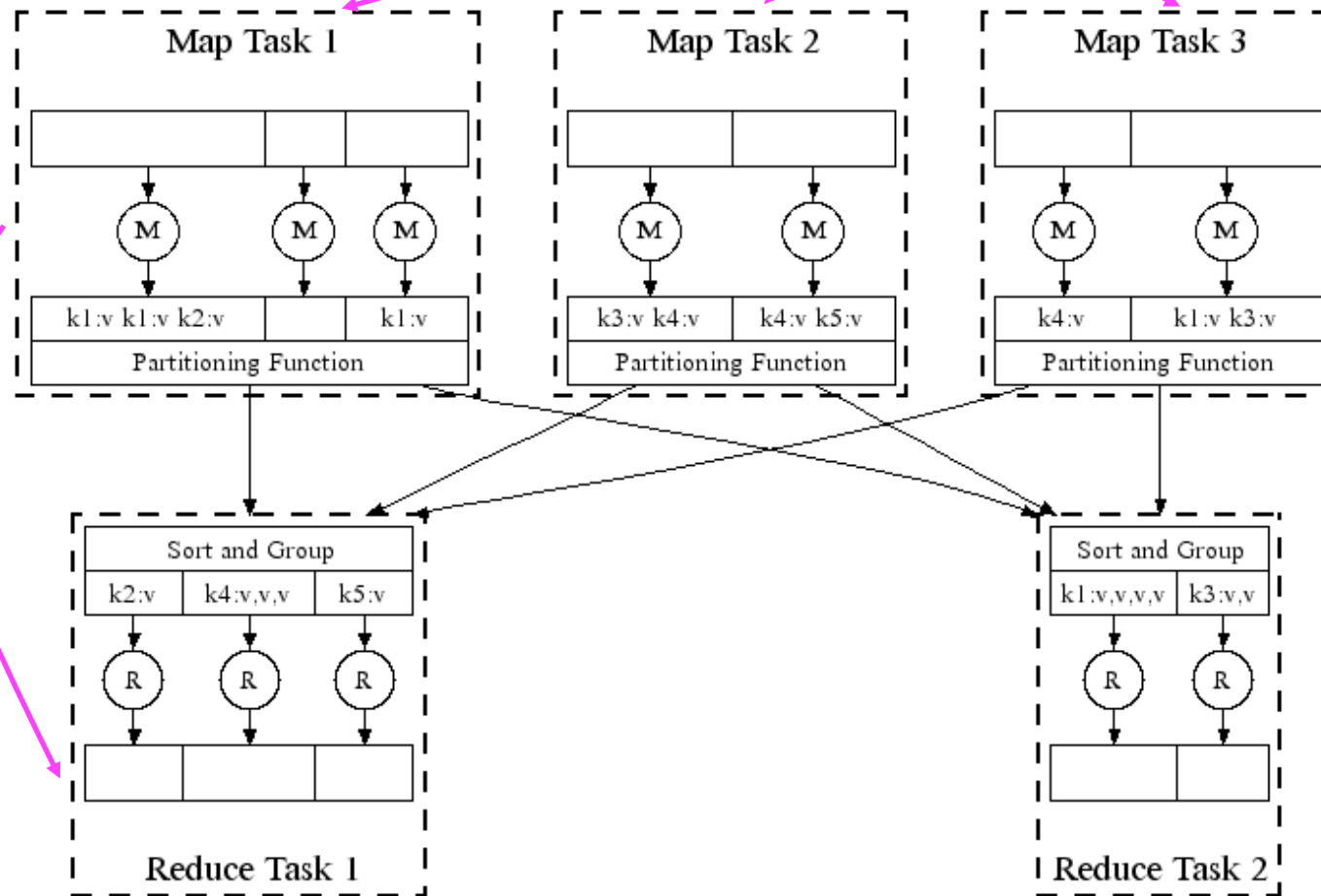
Execution



Parallel Execution

data
parallelism

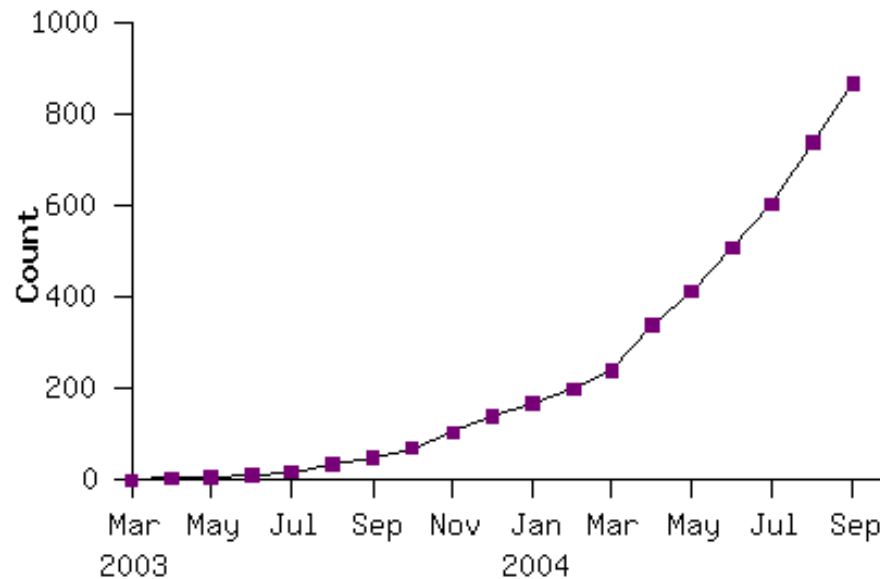
pipeline
parallelism



Key: no implicit dependencies between map or reduce tasks

Model is Widely Applicable

MapReduce Programs In Google Source Tree at end of 2004



Example uses:

distributed grep
term-vector / host
document clustering

...

distributed sort
web access log stats
machine learning

...

web link-graph reversal
inverted index construction
statistical machine translation

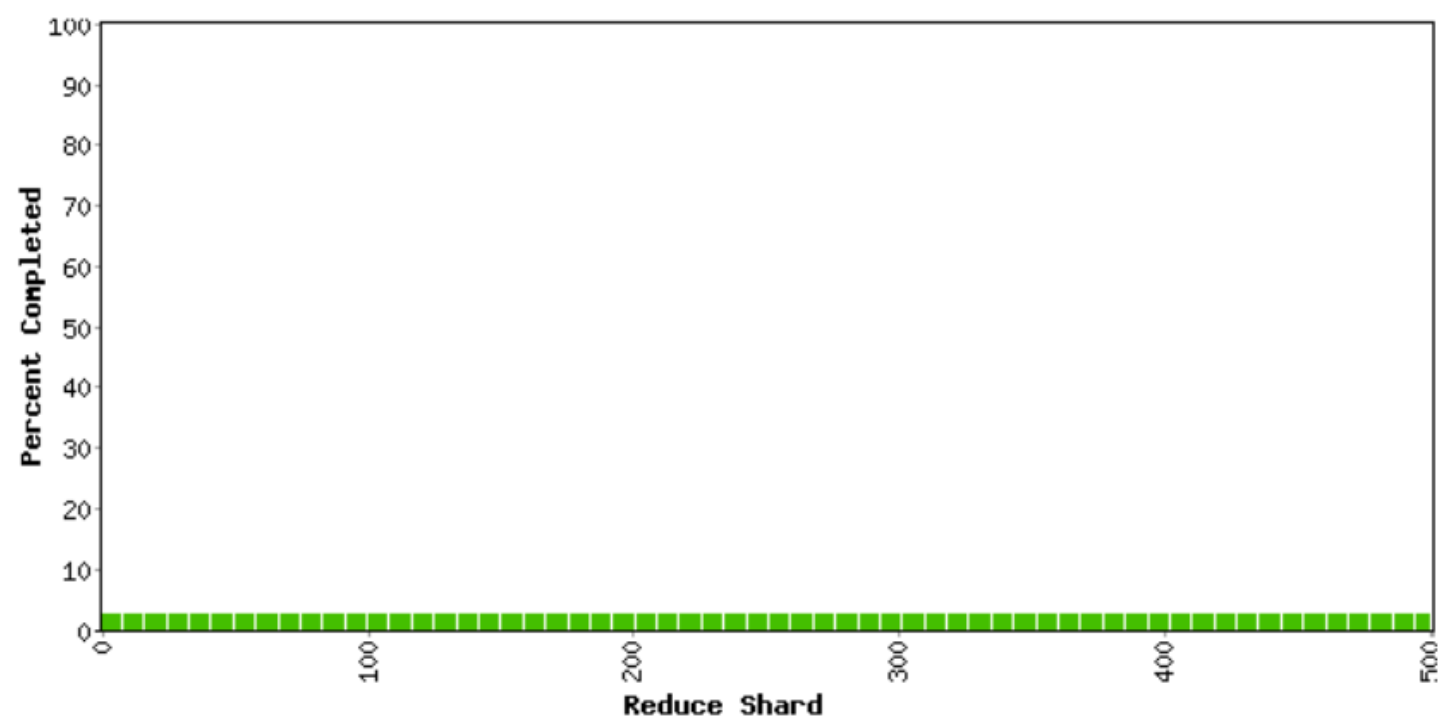
...

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 00 min 18 sec

323 workers; 0 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	0	323	878934.6	1314.4	717.0
Shuffle	500	0	323	717.0	0.0	0.0
Reduce	500	0	0	0.0	0.0	0.0



Counters

Variable	Minute
Mapped (MB/s)	72.5
Shuffle (MB/s)	0.0
Output (MB/s)	0.0
doc-index-hits	145825686
docs-indexed	506631
dups-in-index-merge	0
mr-operator-calls	508192
mr-operator-	506631

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

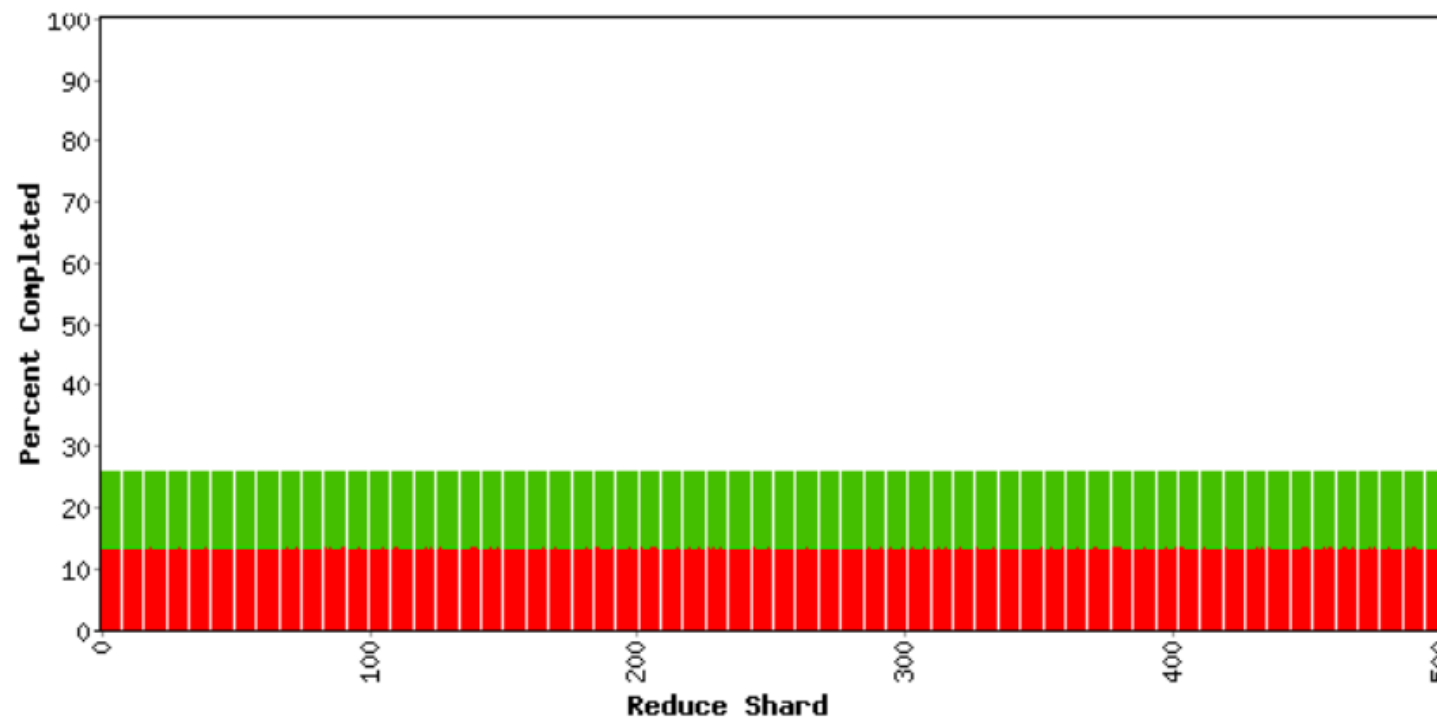
Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 05 min 07 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	1857	1707	878934.6	191995.8	113936.6
Shuffle	500	0	500	113936.6	57113.7	57113.7
Reduce	500	0	0	57113.7	0.0	0.0

Counters

Variable	Minute
Mapped (MB/s)	699.1
Shuffle (MB/s)	349.5
Output (MB/s)	0.0
doc-index-hits	5004411944
docs-indexed	17290135
dups-in-index-merge	0
mr-operator-calls	17331371
mr-operator-outputs	17290135

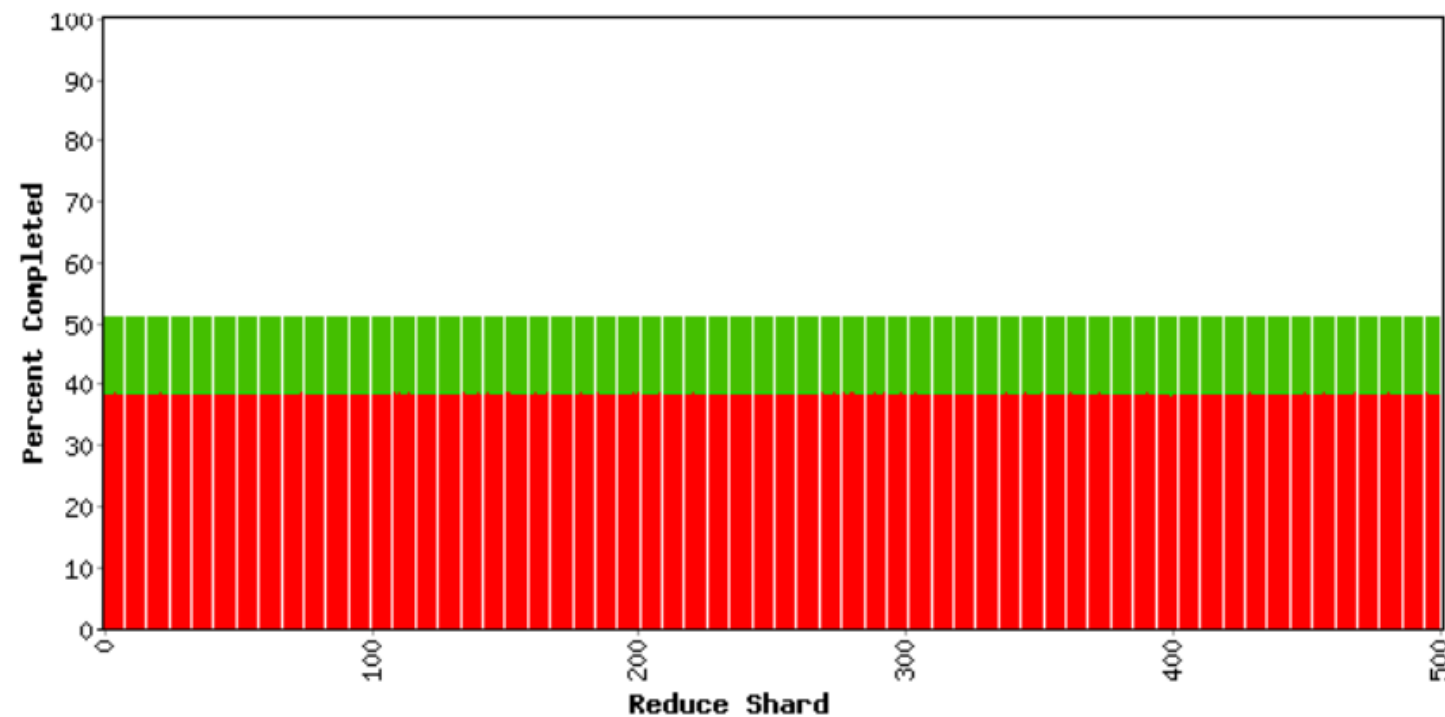


MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 10 min 18 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	5354	1707	878934.6	406020.1	241058.2
Shuffle	500	0	500	241058.2	196362.5	196362.5
Reduce	500	0	0	196362.5	0.0	0.0



Counters

Variable	Minute
Mapped (MB/s)	704.4
Shuffle (MB/s)	371.9
Output (MB/s)	0.0
doc-index-hits	5000364228
docs-indexed	17300709
dups-in-index-merge	0
mr-operator-calls	17342493
mr-operator-outputs	17300709

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 15 min 31 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	8841	1707	878934.6	621608.5	369459.8
Shuffle	500	0	500	369459.8	326986.8	326986.8
Reduce	500	0	0	326986.8	0.0	0.0

Counters

Variable	Minute
Mapped (MB/s)	706.5
Shuffle (MB/s)	419.2
Output (MB/s)	0.0
doc-index-hits	4982870667
docs-indexed	17229926
dups-in-index-merge	0
mr-operator-calls	17272056
mr-operator-outputs	17229926

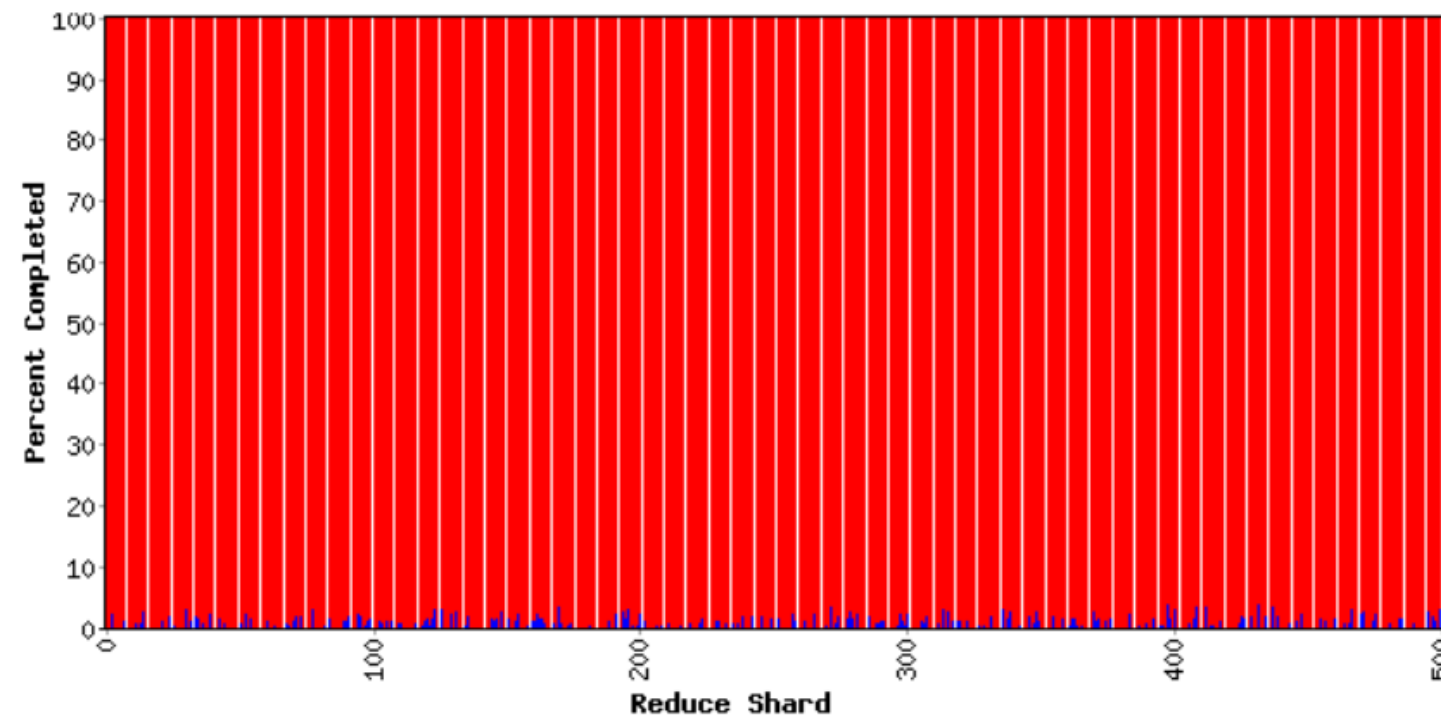


MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 29 min 45 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	13853	0	878934.6	878934.6	523499.2
Shuffle	500	195	305	523499.2	523389.6	523389.6
Reduce	500	0	195	523389.6	2685.2	2742.6



Counters

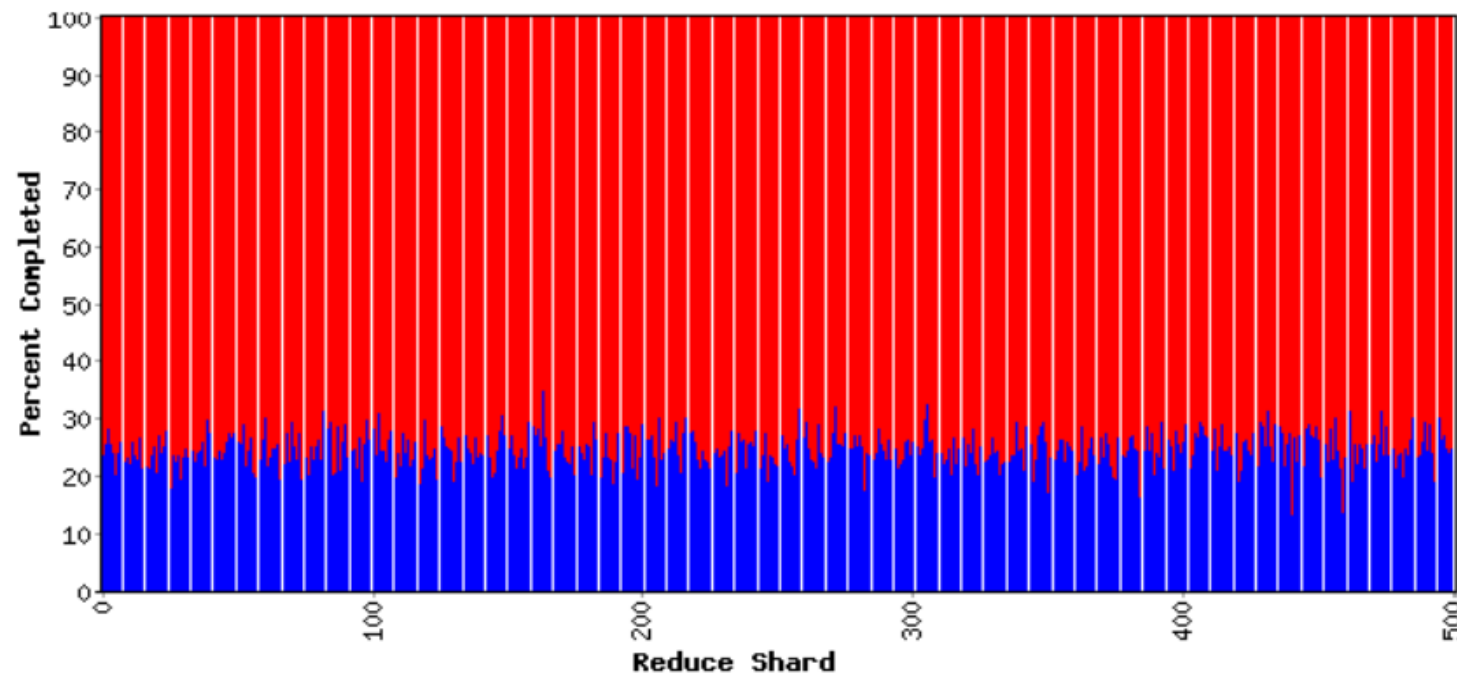
Variable	Minute	
Mapped (MB/s)	0.3	
Shuffle (MB/s)	0.5	
Output (MB/s)	45.7	
doc-index-hits	2313178	105
docs-indexed	7936	
dups-in-index-merge	0	
mr-merge-calls	1954105	
mr-merge-outputs	1954105	

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 31 min 34 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	13853	0	878934.6	878934.6	523499.2
Shuffle	500	500	0	523499.2	523499.5	523499.5
Reduce	500	0	500	523499.5	133837.8	136929.6



Counters

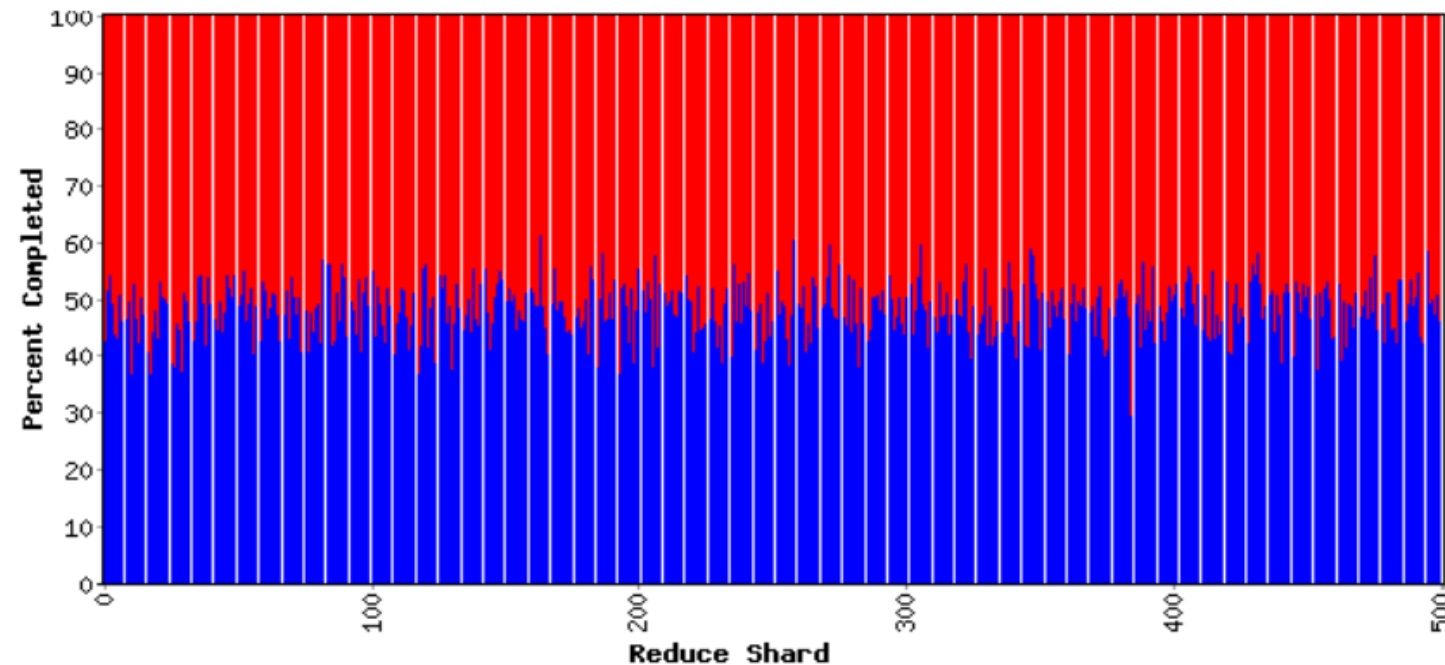
Variable	Minute	
Mapped (MB/s)	0.0	
Shuffle (MB/s)	0.1	
Output (MB/s)	1238.8	
doc-index-hits	0	10
docs-indexed	0	
dups-in-index-merge	0	
mr-merge-calls	51738599	
mr-merge-outputs	51738599	

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 33 min 22 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	13853	0	878934.6	878934.6	523499.2
Shuffle	500	500	0	523499.2	523499.5	523499.5
Reduce	500	0	500	523499.5	263283.3	269351.2



Counters

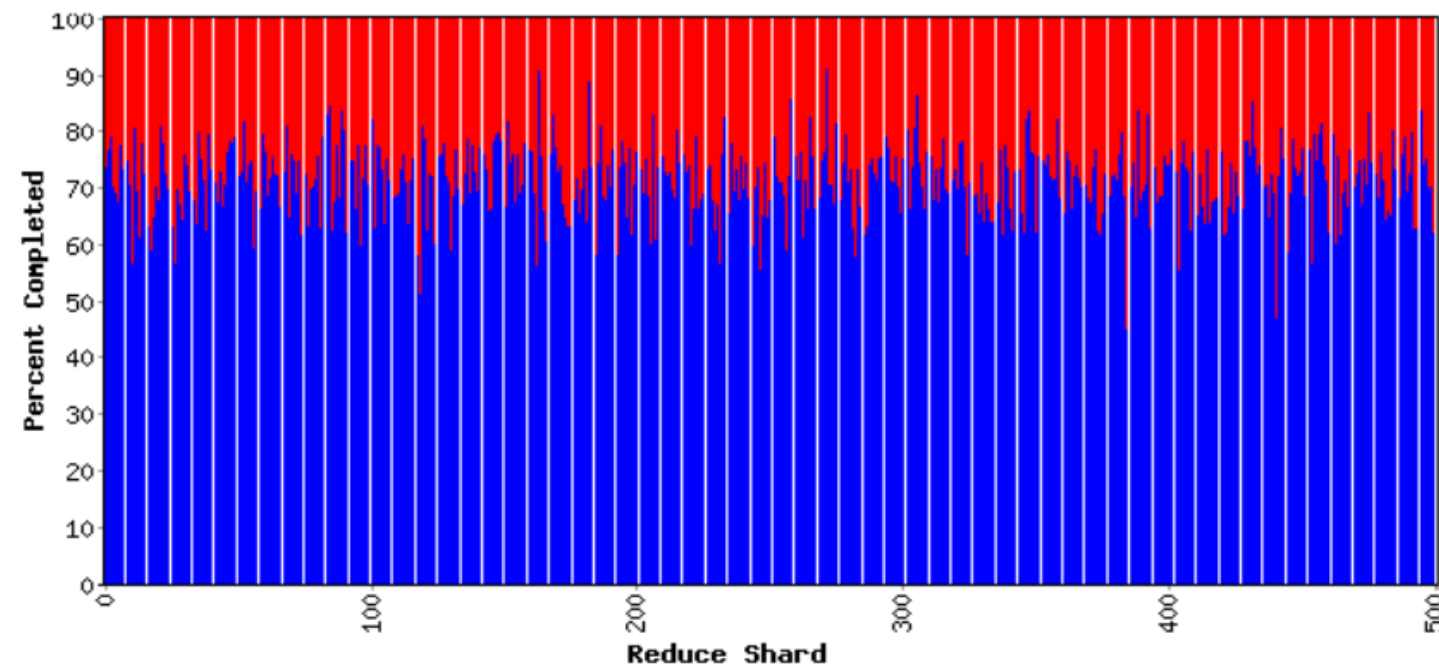
Variable	Minute	
Mapped (MB/s)	0.0	
Shuffle (MB/s)	0.0	
Output (MB/s)	1225.1	
doc-index-hits	0	10
docs-indexed	0	
dups-in-index-merge	0	
mr-merge-calls	51842100	
mr-merge-outputs	51842100	

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 35 min 08 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	13853	0	878934.6	878934.6	523499.2
Shuffle	500	500	0	523499.2	523499.5	523499.5
Reduce	500	0	500	523499.5	390447.6	399457.2



Counters

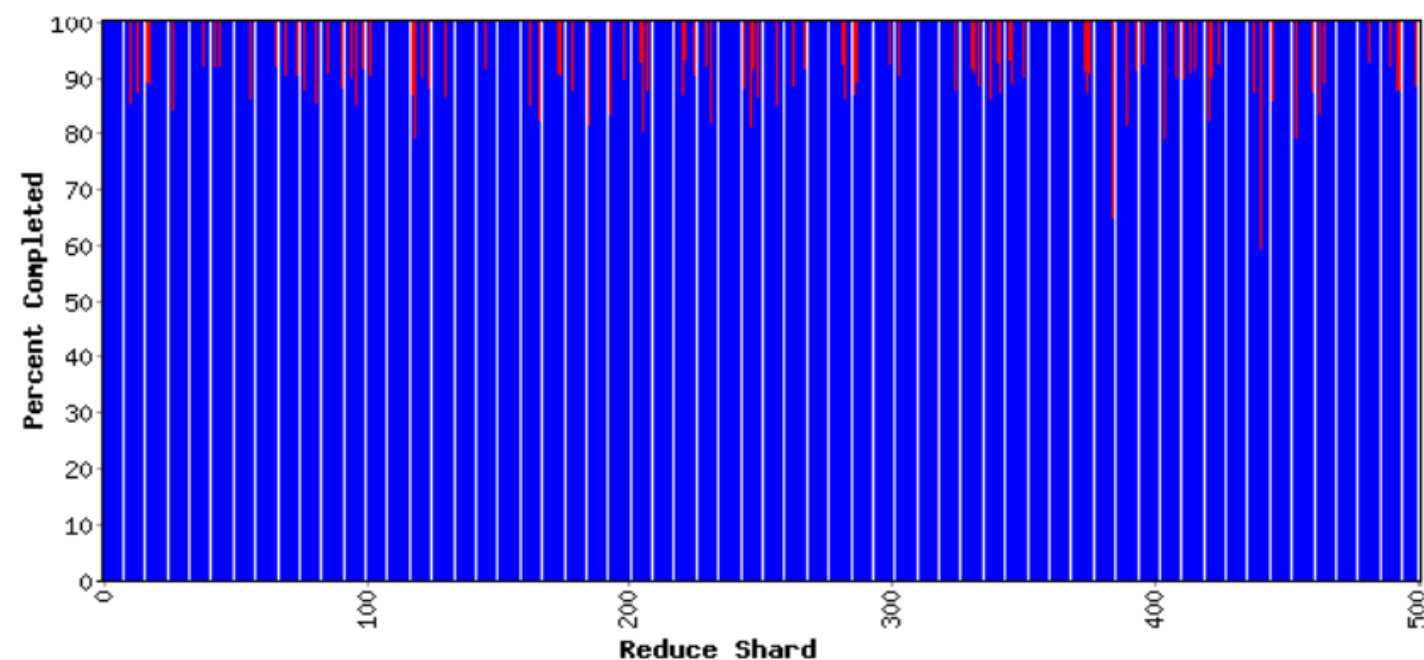
Variable	Minute	
Mapped (MB/s)	0.0	
Shuffle (MB/s)	0.0	
Output (MB/s)	1222.0	
doc-index-hits	0	10
docs-indexed	0	
dups-in-index-merge	0	
mr-merge-calls	51640600	
mr-merge-outputs	51640600	

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 37 min 01 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	13853	0	878934.6	878934.6	523499.2
Shuffle	500	500	0	523499.2	520468.6	520468.6
Reduce	500	406	94	520468.6	512265.2	514373.3



Counters

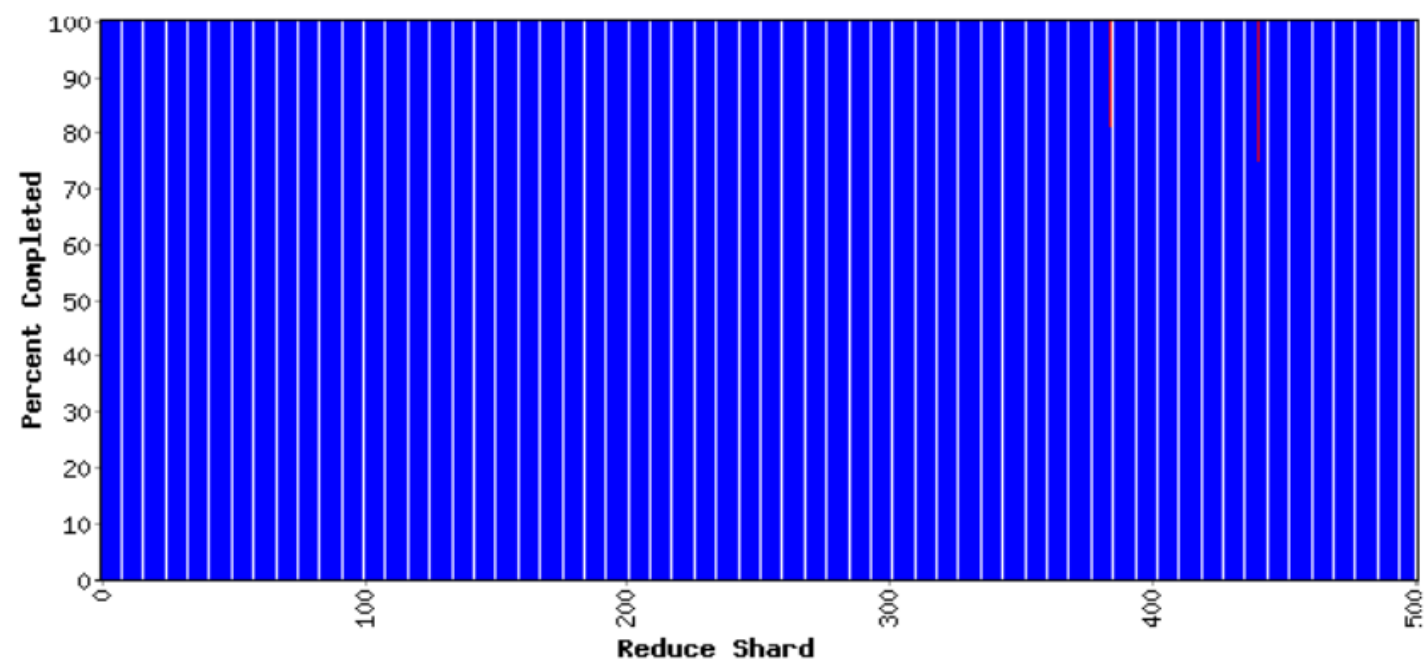
Variable	Minute
Mapped (MB/s)	0.0
Shuffle (MB/s)	0.0
Output (MB/s)	849.5
doc-index-hits	0 10
docs-indexed	0
dups-in-index-merge	0
mr-merge-calls	35083350
mr-merge-outputs	35083350

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 38 min 56 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	13853	0	878934.6	878934.6	523499.2
Shuffle	500	500	0	523499.2	519781.8	519781.8
Reduce	500	498	2	519781.8	519394.7	519440.7



Counters

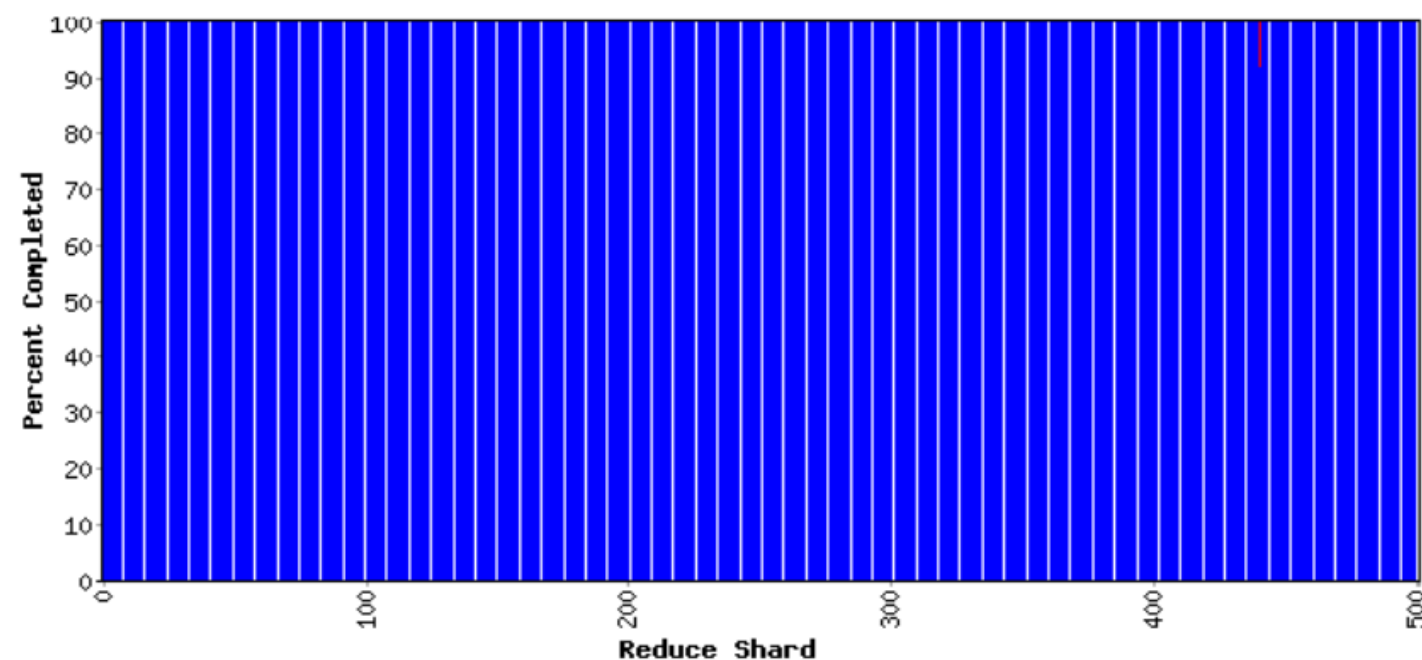
Variable	Minute	
Mapped (MB/s)	0.0	
Shuffle (MB/s)	0.0	
Output (MB/s)	9.4	
doc-index-hits	0	1056
docs-indexed	0	3
dups-in-index-merge	0	
mr-merge-calls	394792	3
mr-merge-outputs	394792	3

MapReduce status: MR_Indexer-beta6-large-2003_10_28_00_03

Started: Fri Nov 7 09:51:07 2003 -- up 0 hr 40 min 43 sec

1707 workers; 1 deaths

Type	Shards	Done	Active	Input(MB)	Done(MB)	Output(MB)
Map	13853	13853	0	878934.6	878934.6	523499.2
Shuffle	500	500	0	523499.2	519774.3	519774.3
Reduce	500	499	1	519774.3	519735.2	519764.0



Counters

Variable	Minute	
Mapped (MB/s)	0.0	
Shuffle (MB/s)	0.0	
Output (MB/s)	1.9	
doc-index-hits	0	1050
docs-indexed	0	:
dups-in-index-merge	0	
mr-merge-calls	73442	:
mr-merge-outputs	73442	:

Communication by files

- Communication between workers happens by writing files to the Google Filesystem
 - This is a distributed, fault-tolerant file system that uses replication
- Mapper files are “intermediate” and determined on the fly
 - The location of these files is given to the master, who passes along the locations to the reducers
- Reducer files are final results, given to the application

Fault Tolerance

- Handle worker failures via re-execution
 - Detect failure via periodic heartbeats
- Re-execute in-progress **reduce** tasks and in-progress and completed **map** tasks
 - Can do this easily since inputs are stored on the file system
- Key: Map and Reduce are *functional*
 - So they can always be reexecuted to produce the same answer

Optimizations

- Perform redundant computations on idle resources
 - Whichever one finishes "wins"
- Exploit locality: send tasks to the data
- Exploit **reduce** func. properties to further parallelize, decrease net. bandwidth
 - Requires this function is commutative and associative. Why?

The programming model is the key

- Simple control makes dependencies evident
 - Can automate scheduling of tasks and optimization
 - Map, reduce for different keys, embarrassingly parallel
 - Pipeline between mappers, reducers evident
- **map** and **reduce** are pure functions
 - Can rerun them to get the same answer
 - in the case of failure, or
 - to use idle resources toward faster completion
 - No worry about data races, deadlocks, etc. since there is no shared state

Hadoop

- Apache Hadoop is an open source implementation of map reduce
 - Uses Java as the programming model
 - <http://wiki.apache.org/hadoop/HadoopMapReduce>
- Uses HDFS instead of GFS for implemented distributed, fault-tolerant file system
 - But you don't have to run it in this mode

MapReduce in Hadoop

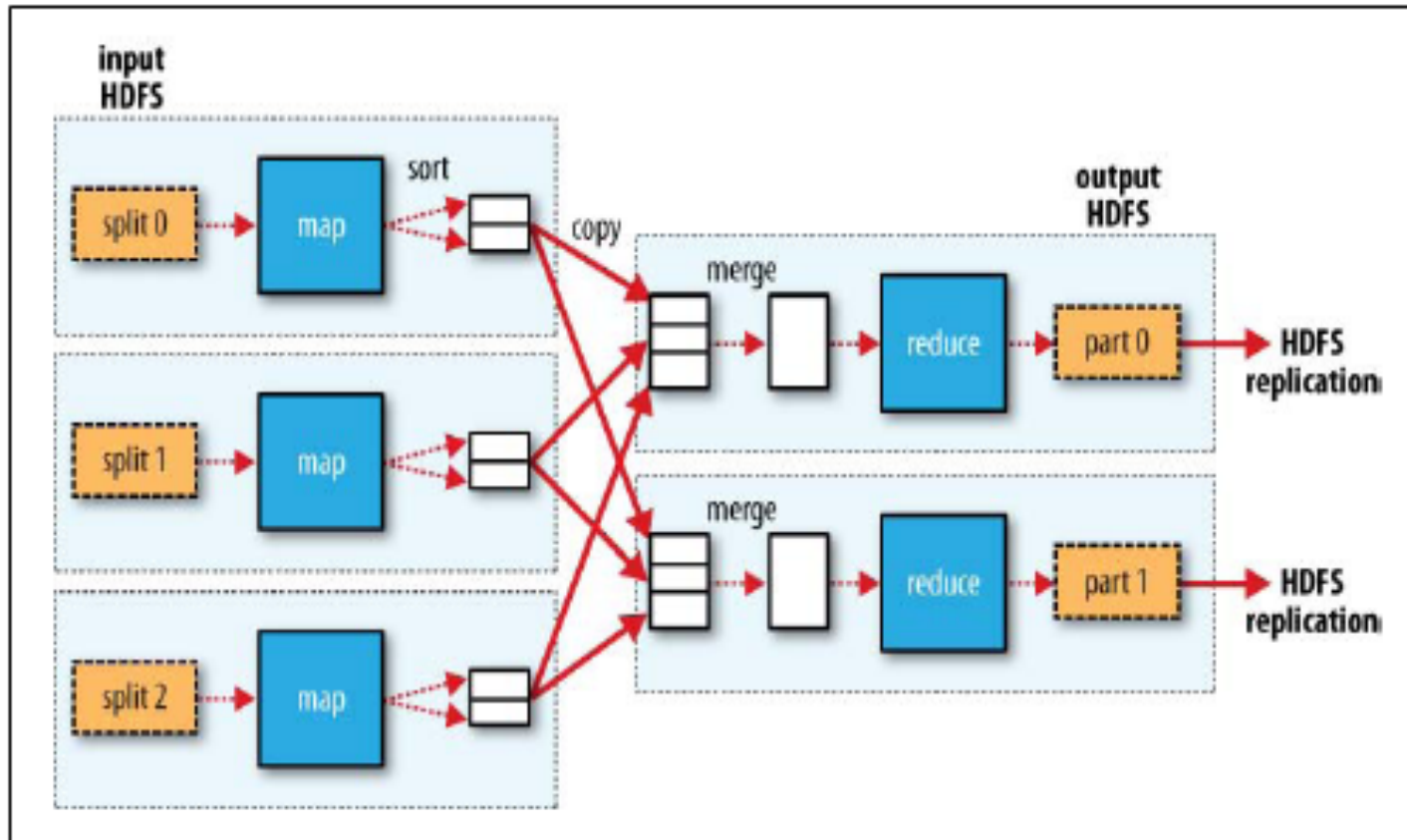
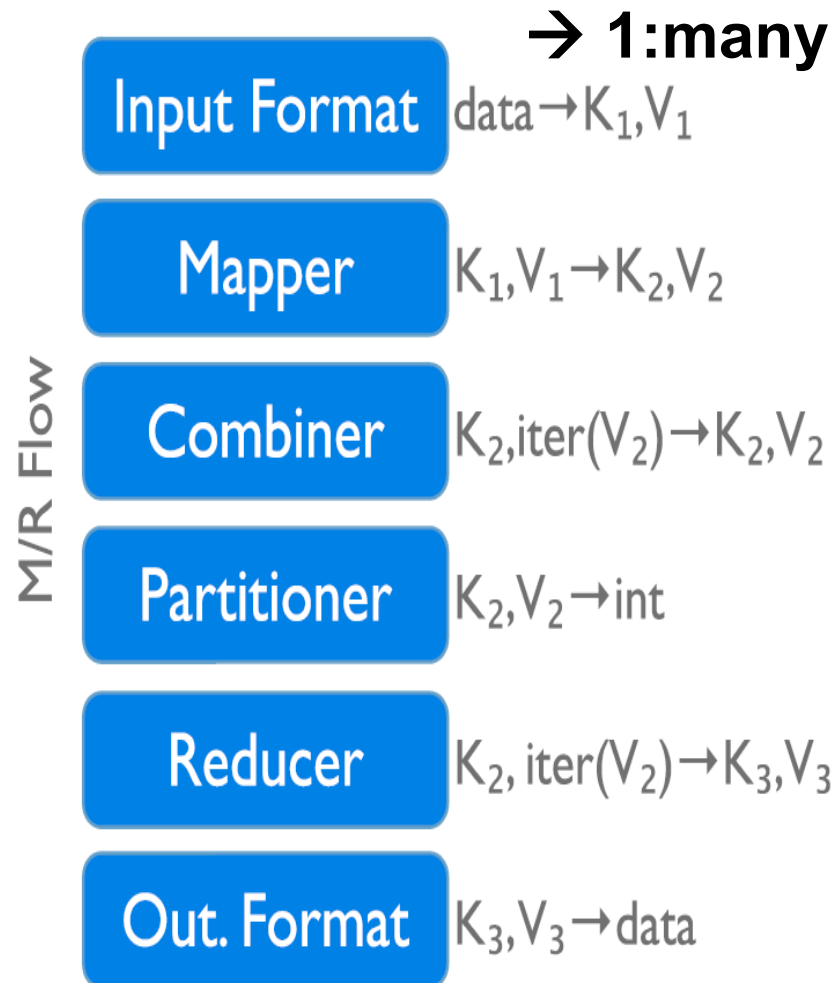


Figure 2-3. MapReduce data flow with multiple reduce tasks

Data Flow in a MapReduce Program in Hadoop

- InputFormat
- Map function
- Partitioner
- Sorting & Merging
- Combiner
- Shuffling
- Merging
- Reduce function
- OutputFormat



Lifecycle of a MapReduce Job

```
File Edit Options Buffers Tools Java Help
public class WordCount {

    public static class Map extends MapReduceBase implements
        Mapper<LongWritable, Text, Text, IntWritable> {
        private final static IntWritable one = new IntWritable(1);
        private Text word = new Text();

        public void map(LongWritable key, Text value, OutputCollector<Text, IntWritable>
            output, Reporter reporter) throws IOException {
            String line = value.toString();
            StringTokenizer tokenizer = new StringTokenizer(line);
            while (tokenizer.hasMoreTokens()) {
                word.set(tokenizer.nextToken());
                output.collect(word, one);
            }
        }

    public static class Reduce extends MapReduceBase implements
        Reducer<Text, IntWritable, Text, IntWritable> {
        public void reduce(Text key, Iterator<IntWritable> values, OutputCollector<Text,
            IntWritable> output, Reporter reporter) throws IOException {
            int sum = 0;
            while (values.hasNext()) { sum += values.next().get(); }
            output.collect(key, new IntWritable(sum));
        }

    public static void main(String[] args) throws Exception {
        JobConf conf = new JobConf(WordCount.class);
        conf.setJobName("wordcount");
        conf.setOutputKeyClass(Text.class);
        conf.setOutputValueClass(IntWritable.class);
        conf.setMapperClass(Map.class);
        conf.setCombinerClass(Reduce.class);
        conf.setReducerClass(Reduce.class);
        conf.setInputFormat(TextInputFormat.class);
        conf.setOutputFormat(TextOutputFormat.class);
        FileInputFormat.setInputPaths(conf, new Path(args[0]));
        FileOutputFormat.setOutputPath(conf, new Path(args[1]));

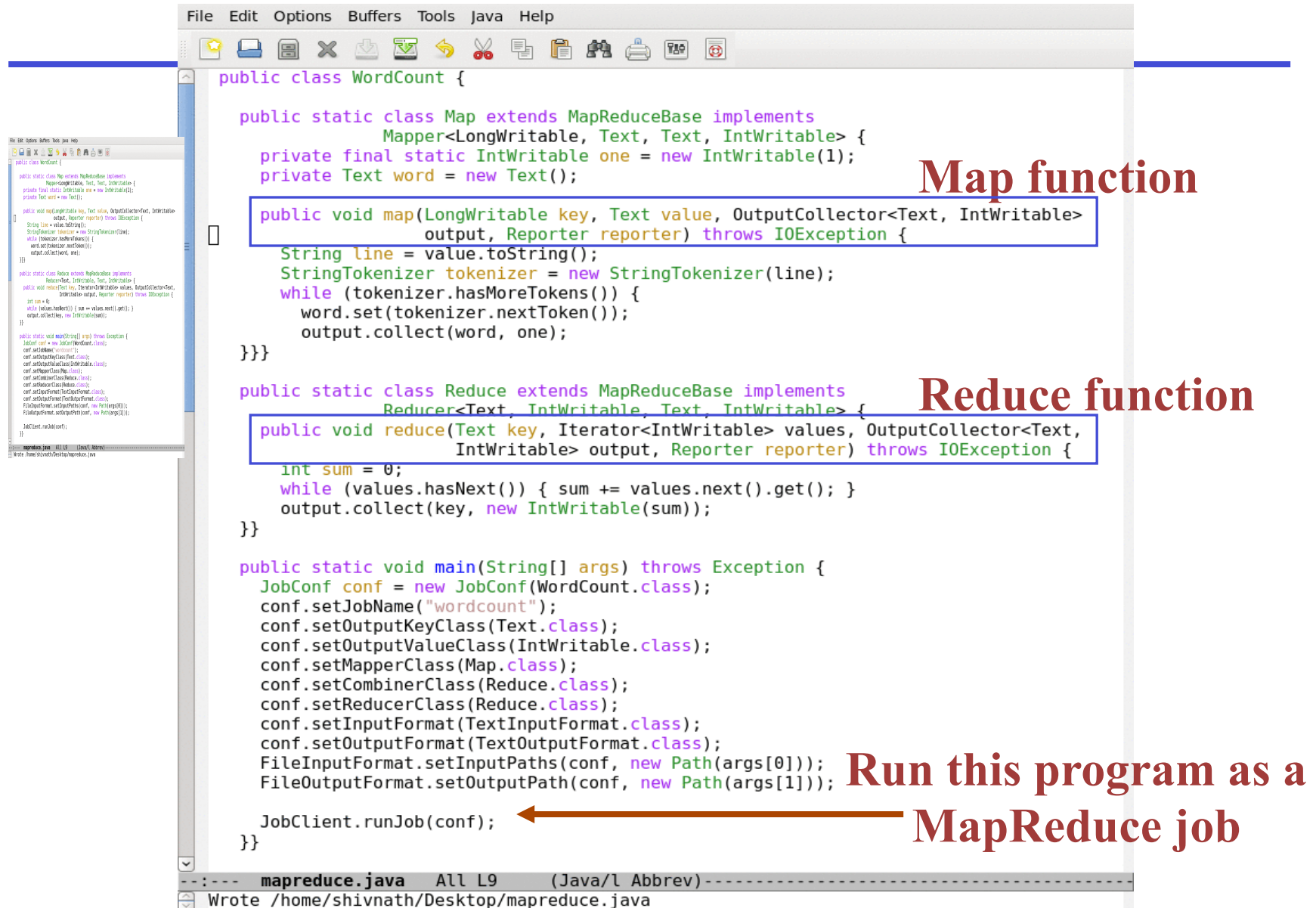
        JobClient.runJob(conf);
    }
}
```

Map function

Reduce function

Run this program as a
MapReduce job

Lifecycle of a MapReduce Job



The screenshot shows a Java IDE with the following code for `WordCount.java`:

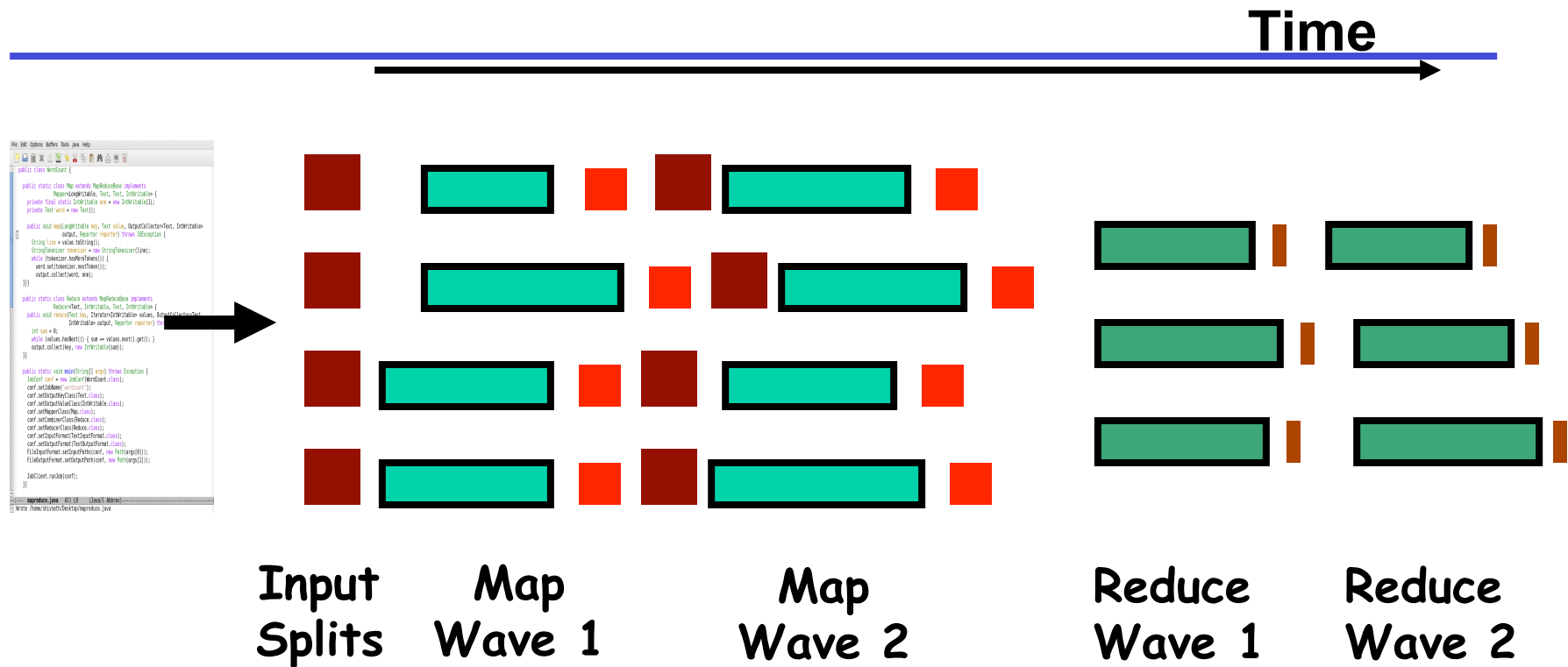
```
public class WordCount {  
  
    public static class Map extends MapReduceBase implements  
        Mapper<LongWritable, Text, Text, IntWritable> {  
        private final static IntWritable one = new IntWritable(1);  
        private Text word = new Text();  
  
        public void map(LongWritable key, Text value, OutputCollector<Text, IntWritable>  
            output, Reporter reporter) throws IOException {  
            String line = value.toString();  
            StringTokenizer tokenizer = new StringTokenizer(line);  
            while (tokenizer.hasMoreTokens()) {  
                word.set(tokenizer.nextToken());  
                output.collect(word, one);  
            }  
        }  
    }  
  
    public static class Reduce extends MapReduceBase implements  
        Reducer<Text, IntWritable, Text, IntWritable> {  
        public void reduce(Text key, Iterator<IntWritable> values, OutputCollector<Text,  
            IntWritable> output, Reporter reporter) throws IOException {  
            int sum = 0;  
            while (values.hasNext()) { sum += values.next().get(); }  
            output.collect(key, new IntWritable(sum));  
        }  
    }  
  
    public static void main(String[] args) throws Exception {  
        JobConf conf = new JobConf(WordCount.class);  
        conf.setJobName("wordcount");  
        conf.setOutputKeyClass(Text.class);  
        conf.setOutputValueClass(IntWritable.class);  
        conf.setMapperClass(Map.class);  
        conf.setCombinerClass(Reduce.class);  
        conf.setReducerClass(Reduce.class);  
        conf.setInputFormat(TextInputFormat.class);  
        conf.setOutputFormat(TextOutputFormat.class);  
        FileInputFormat.setInputPaths(conf, new Path(args[0]));  
        FileOutputFormat.setOutputPath(conf, new Path(args[1]));  
  
        JobClient.runJob(conf);  
    }  
}
```

Annotations on the code:

- Map function**: Points to the `map` method in the `Map` class.
- Reduce function**: Points to the `reduce` method in the `Reduce` class.
- Run this program as a MapReduce job**: Points to the `JobClient.runJob(conf);` line.

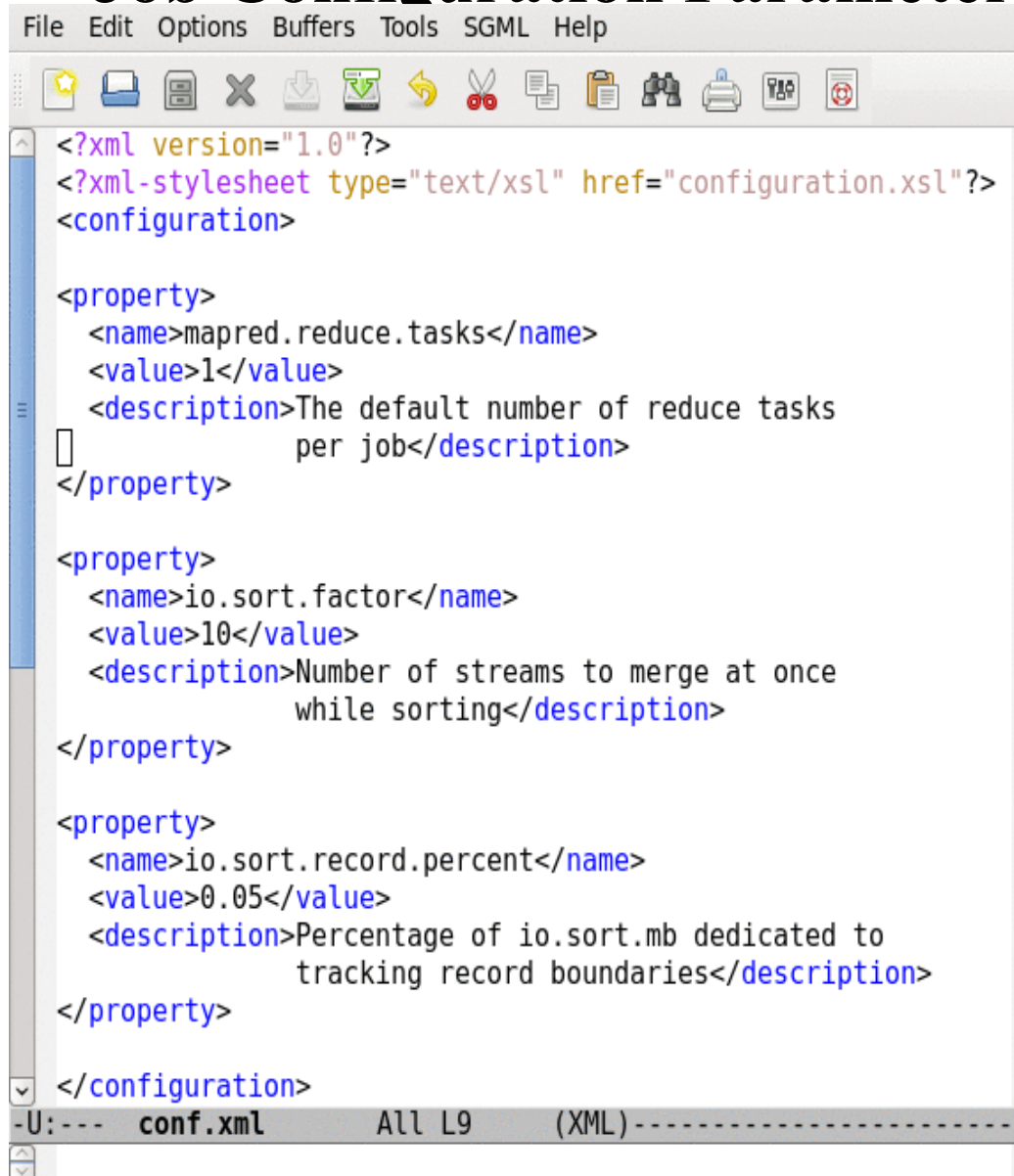
The IDE status bar at the bottom shows: `mapreduce.java All L9 (Java/l Abbrev)` and `Wrote /home/shivnath/Desktop/mapreduce.java`.

Lifecycle of a MapReduce Job



How are the number of splits, number of map and reduce tasks, memory allocation to tasks, etc., determined?

Job Configuration Parameters

A screenshot of an XML editor window. The title bar shows 'File Edit Options Buffers Tools SGML Help'. The toolbar contains icons for opening, saving, undo, redo, cut, copy, paste, and other editing functions. The main text area displays an XML configuration file named 'conf.xml'. The XML structure is as follows: <?xml version="1.0"?> <?xml-stylesheet type="text/xsl" href="configuration.xsl"?> <configuration> <property> <name>mapred.reduce.tasks</name> <value>1</value> <description>The default number of reduce tasks per job</description> </property> <property> <name>io.sort.factor</name> <value>10</value> <description>Number of streams to merge at once while sorting</description> </property> <property> <name>io.sort.record.percent</name> <value>0.05</value> <description>Percentage of io.sort.mb dedicated to tracking record boundaries</description> </property> </configuration> The status bar at the bottom shows '-U:--- conf.xml All L9 (XML)-----'.

```
<?xml version="1.0"?>
<?xml-stylesheet type="text/xsl" href="configuration.xsl"?>
<configuration>

  <property>
    <name>mapred.reduce.tasks</name>
    <value>1</value>
    <description>The default number of reduce tasks
per job</description>
  </property>

  <property>
    <name>io.sort.factor</name>
    <value>10</value>
    <description>Number of streams to merge at once
while sorting</description>
  </property>

  <property>
    <name>io.sort.record.percent</name>
    <value>0.05</value>
    <description>Percentage of io.sort.mb dedicated to
tracking record boundaries</description>
  </property>

</configuration>
```

- 190+ parameters in Hadoop
- Set manually or defaults are used

Project 5

- Use (old version of) Hadoop in standalone mode to implement word count example, and a variant of it
- Will run these on your machine and upload some statistics about the results to a Google spreadsheet
 - You submit your code to `submit.cs`

Distributed comp. compared to “big iron”

- According to Wikipedia, Google uses
 - 450,000 servers from 533 MHz Celeron to dual 1.4GHz Pentium III
 - 80GB drive per server, at least
 - 2-4GB memory per machine
 - Jobs processing 100 terabytes of distributed data
- Compare the computing power here to even the most powerful supercomputer
 - Not even close ...