

Distributed computing on large data sets

1

Goals

- Do one computation faster - **latency**
- Do more computations in given time - **throughput**
- Tolerate failure of 1+ machines

2

Distributing query evaluation

- Data distribution?
- Computation distribution?
- Goals
 - Keep all machines busy
 - Be able to replace badly-behaved machines seamlessly!

3

MapReduce framework

- programming model
- implementation for large clusters
- Google introduced for index building and PageRank circa 2003
“for processing and generating large data sets”
- The Apache Hadoop project developed open-source software

4

MapReduce Programming Model

- **input set**: $\{(\text{input key}_i, \text{value}_i) \mid 0 \leq i \leq \text{input size}\}$
 - user chooses type value – e.g. whole document
- **output set**: $\{(\text{output key}_i, \text{value}_i) \mid 0 \leq i \leq \text{output size}\}$
- **Map (written by user)**:
 $(\text{input key}, \text{value}) \rightarrow \{(\text{intermed. key}_j, \text{value}_j) \mid 0 \leq j \leq \text{Map result size}\}$
- **system** groups all Map output pairs for input set by **intermediate key (shuffle phase)**
 - gathers by intermediate key value
 - supply to Reduce by iterator
- **Reduce (written by user)** process intermediate values:
 $(\text{intermed. key}, \text{list of values}) \rightarrow (\text{output key}, \text{value})$

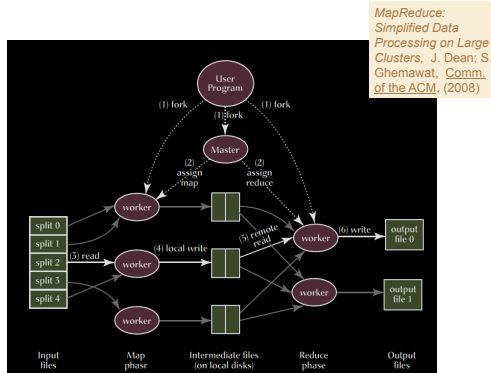
5

MapReduce for building inverted index

- Input pair: (**docID**, contents of doc)
- Map: produce $\{(\text{term}, \text{docID})\}$ for each term appearing in docID
- Input to Reduce: list of all (**term**, docID) pairs for one term
- Output of Reduce: (**term**, sorted list of docIDs containing that term)
 - postings list!

keys 6

Diagram of computation distribution



MapReduce parallelism

- Map phase and shuffle phase may overlap
- Shuffle phase and reduce phase may overlap
- Map phase must finish before reduce phase starts
 - reduce depends on all values associated with a given key

8

MapReduce Fault Tolerance

- Master fails => restart whole computation
- Worker node fails
 - Master detects failure
 - must redo all Map tasks assigned to worker
 - output of completed Map tasks on failed worker's disk
 - for failed Map worker, Master
 - reschedules each Map task
 - notifies reducer workers of change in input location
 - for failed Reduce worker, Master
 - reschedules each Reduce task
 - rescheduling occurs as live workers become available

10

Communication Cost

Rajaraman, Leskovec and Ullman

- Model: algorithm is **acyclic network of tasks**
 - cascaded MapReduce tasks
- **Communication cost task = size of input**
 - bytes versus tuples
- **Commun. cost alg. = sum** of cost of all tasks
- Captures
 - Network cost
 - Disk cost
- Why not output cost?
 - input to another algorithm
 - output “rarely large”

11

Remarks

- Google built on **large collections** of inexpensive “commodity PCs”
 - always some not functioning
- **Solve fault-tolerance problem in software**
 - redundancy & flexibility NOT special-purpose hardware
- Keep **machines** relative **generalists**
 - machine becomes free => assign to any one of set of tasks

13