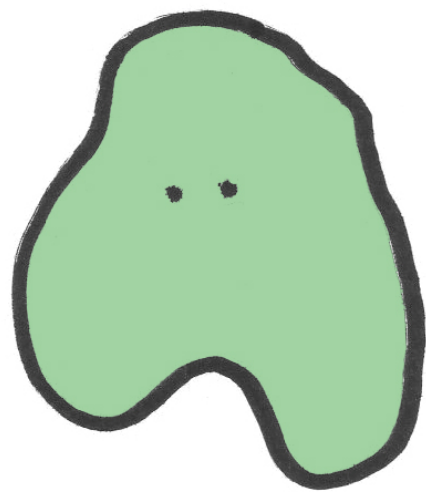




amoeba

A Robust Partitioning Scheme for Ad-Hoc Query Workloads

ANIL SHANBHAG
MIT

J/W Alekh Jindal,
Microsoft

Sam Madden,
MIT

Jorge Quiane,
QCRI

Aaron J. Elmore
Univ. Chicago

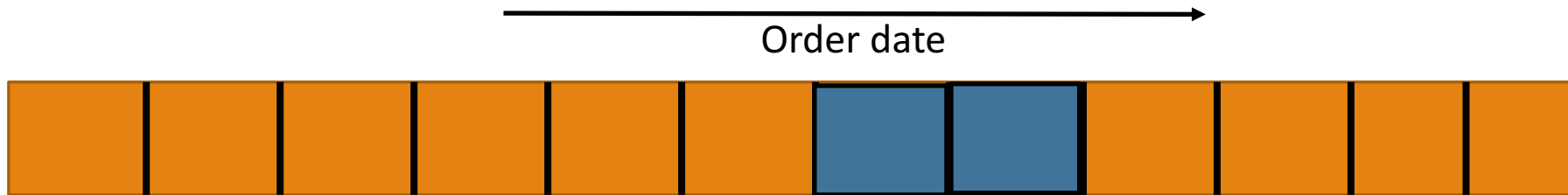
Today

Data collection is cheap => Lots of data !



Data Partitioning

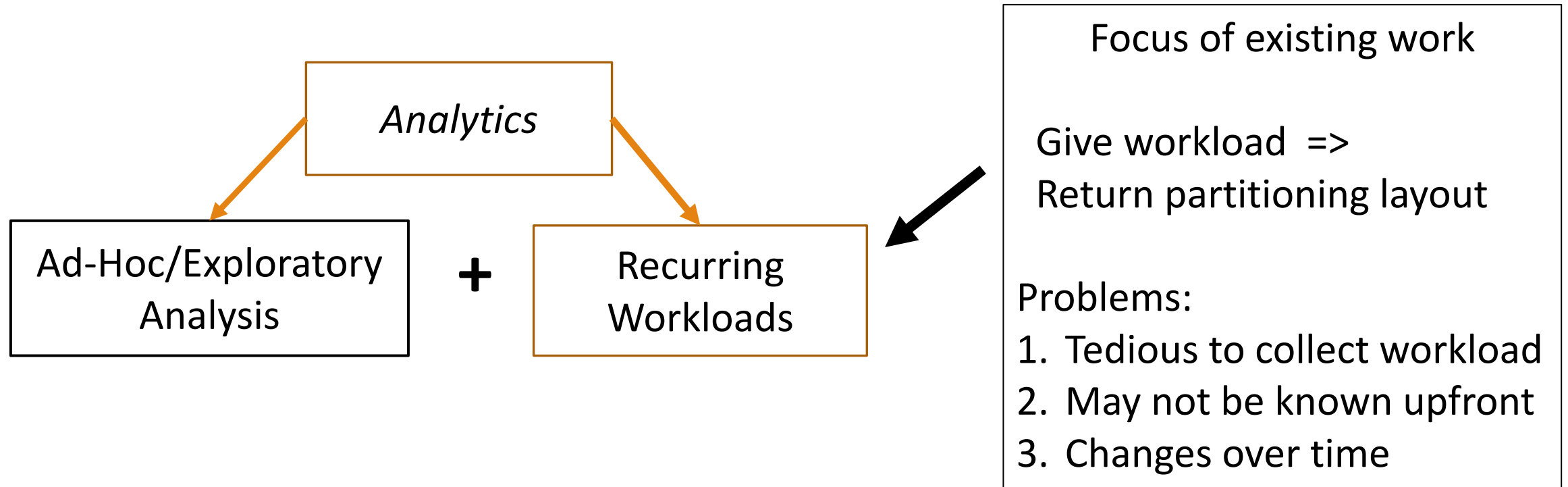
Find average order size for all orders between Sept 10 and Sept 11, 2017



Data Skipping - Skip data blocks not necessary

10% selectivity query => 10x faster if data partitioned on selection predicate

The Problem



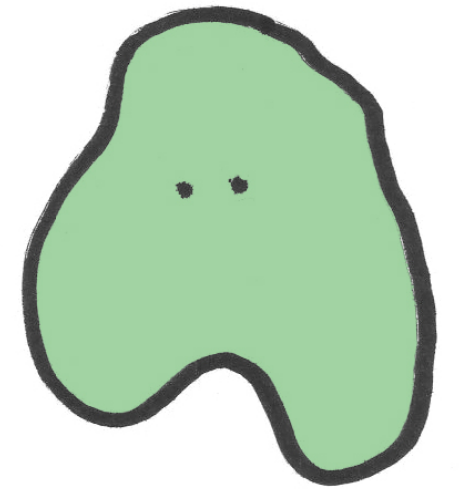
How to get benefits of partitioning in this case ?

Our Approach

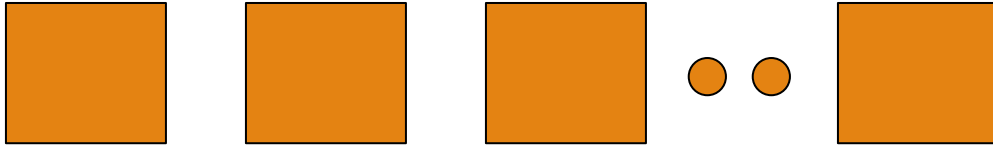
Do everything adaptively !

Two step process:

1. Upfront load the dataset partitioned
2. As users query, *incrementally* improve the partitioning of the data



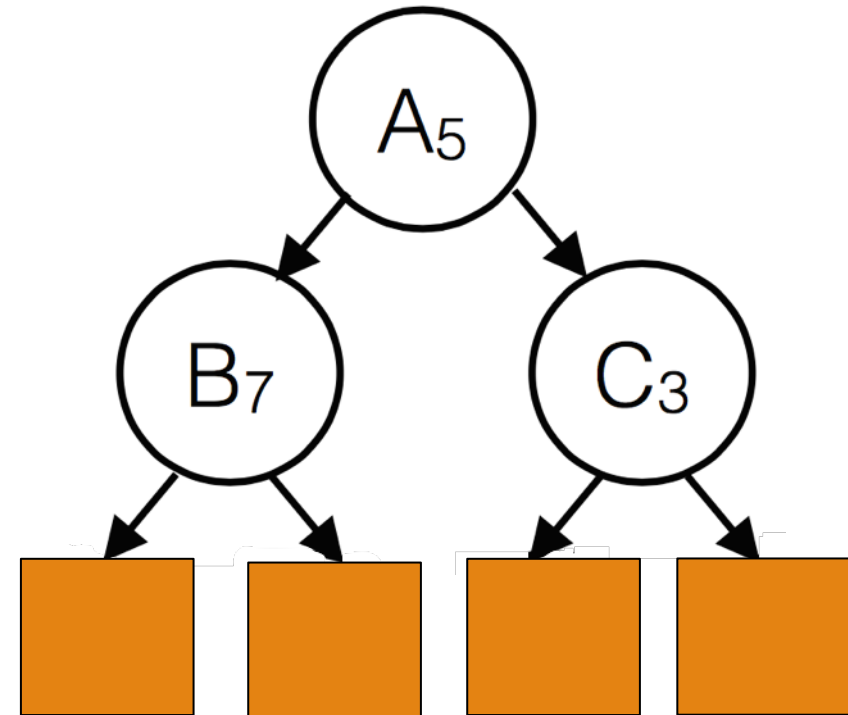
amoeba



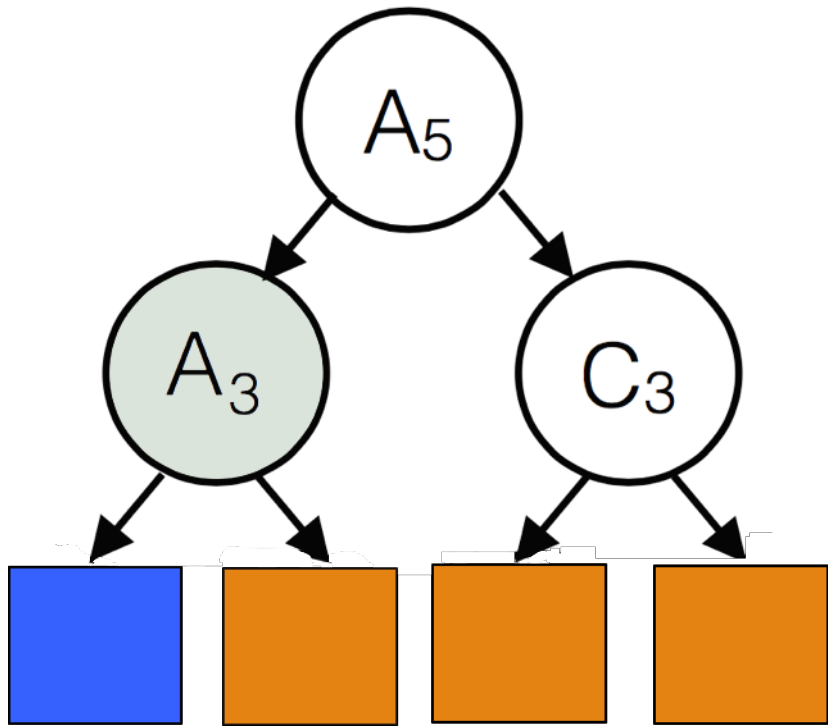
Distributed storage systems like HDFS, files broken into blocks (128 MB chunks)

Upfront Partitioning

- > Instead of partitioning by size, partition by attributes.
- > Same number of blocks created as in HDFS. Each block now has additional metadata



$A \leq 5$ and $B \leq 7$



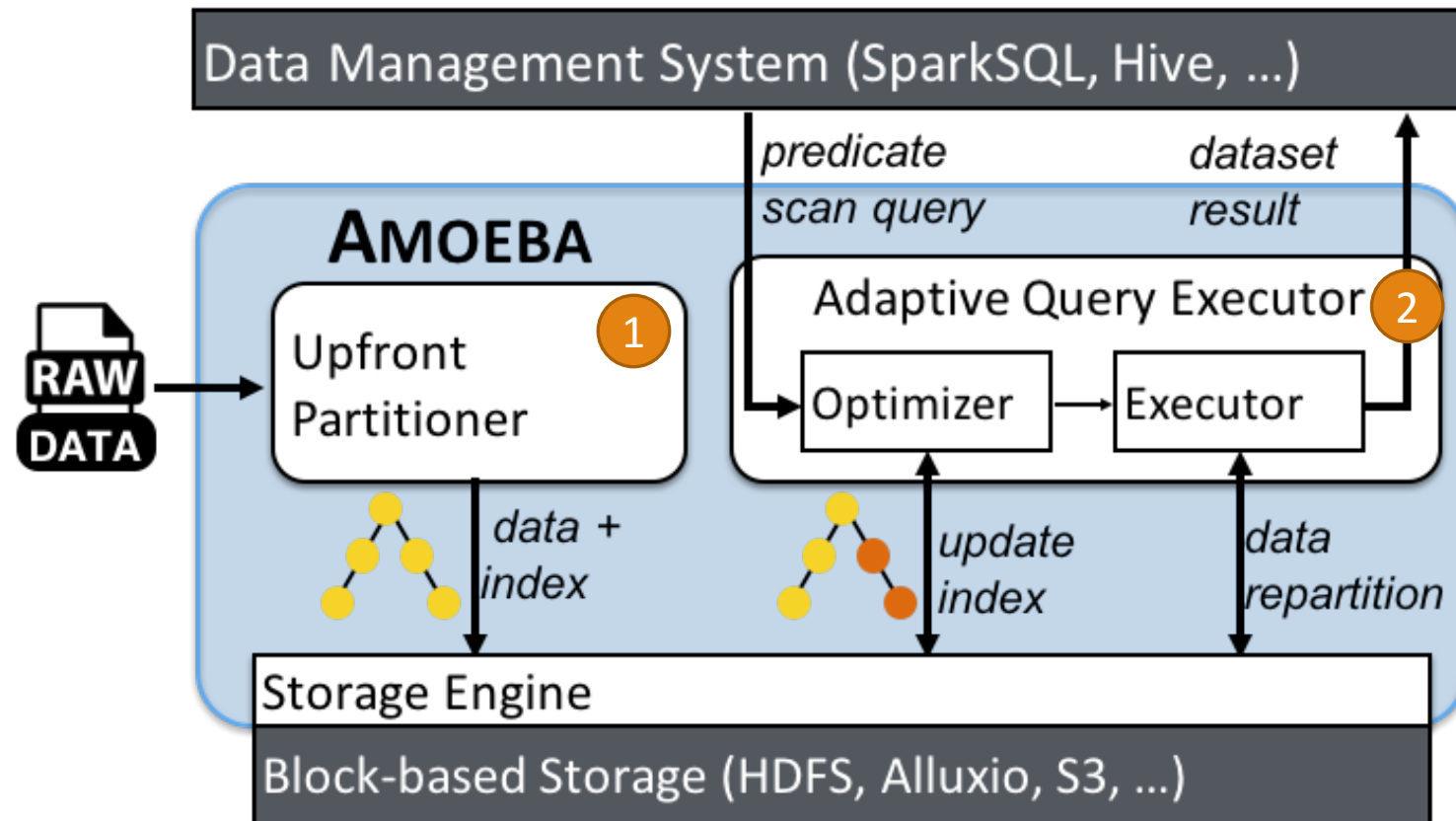
Adaptive Re-Partitioning

When user submits a query, optimizer tries to improve the partitioning by reorganizing the partitioning tree

Here if queries ask $A \leq 3$ many times, replace B_7 by A_3

Done on datasets which are O(1TB) with ~ 8000 node partition trees.

System Architecture



Predicated Scan Query
Example:

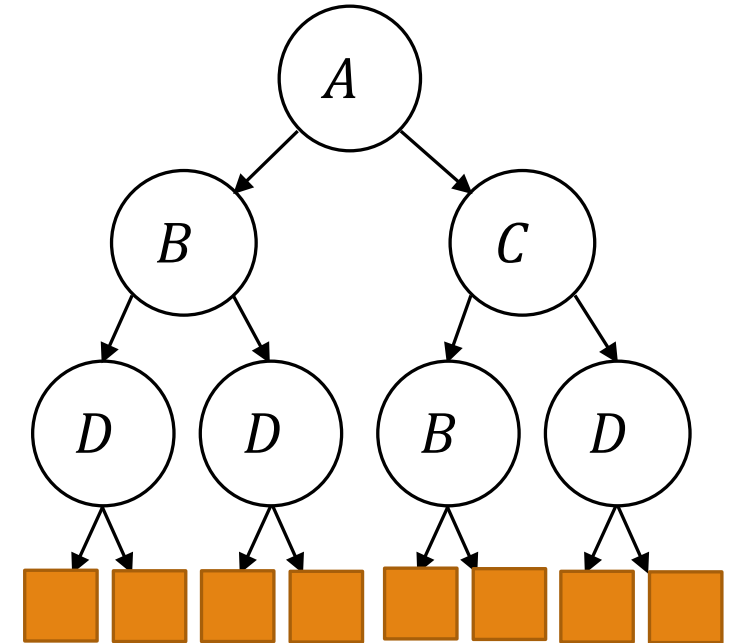
FIND employees WITH
Age < 30 AND
20k < Salary < 40k

1. Upfront Partitioner

Goal: Generate a partitioning tree

WITHOUT an upfront query workload

- > Generates a tree with heterogeneous branching
- > Balance the partitioning benefit across all attributes

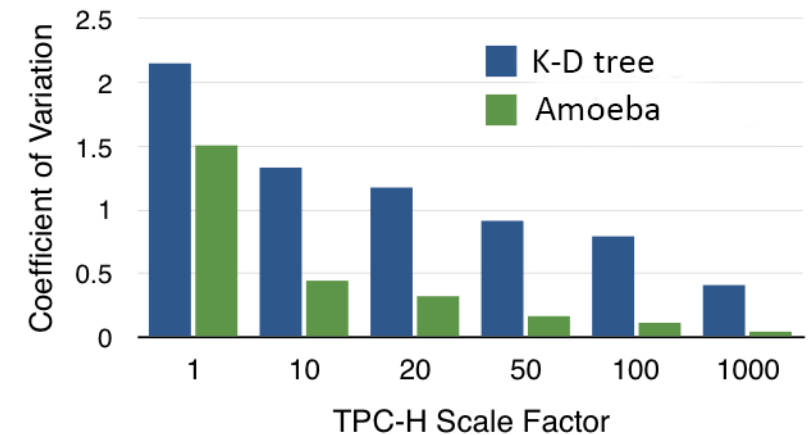
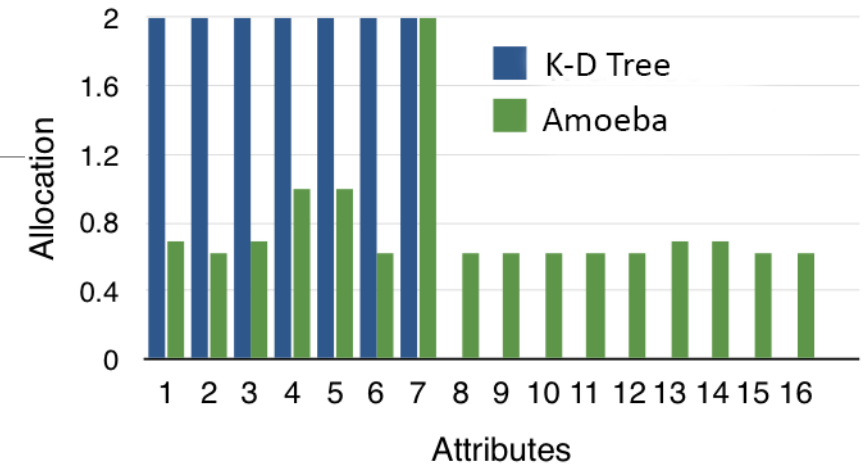
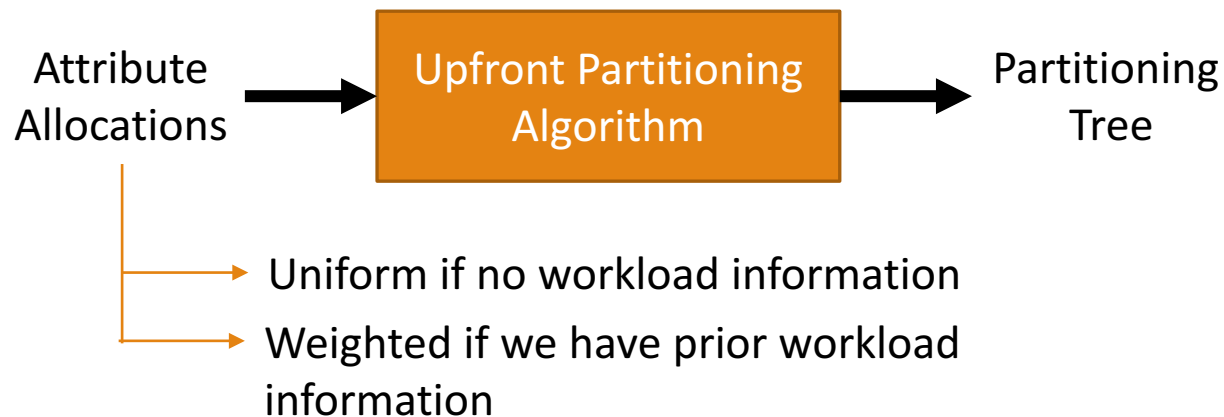


Allocation

Goal: Balance partitioning benefit across attributes

Allocation of attribute $i \sim$ average partitioning of an attribute j

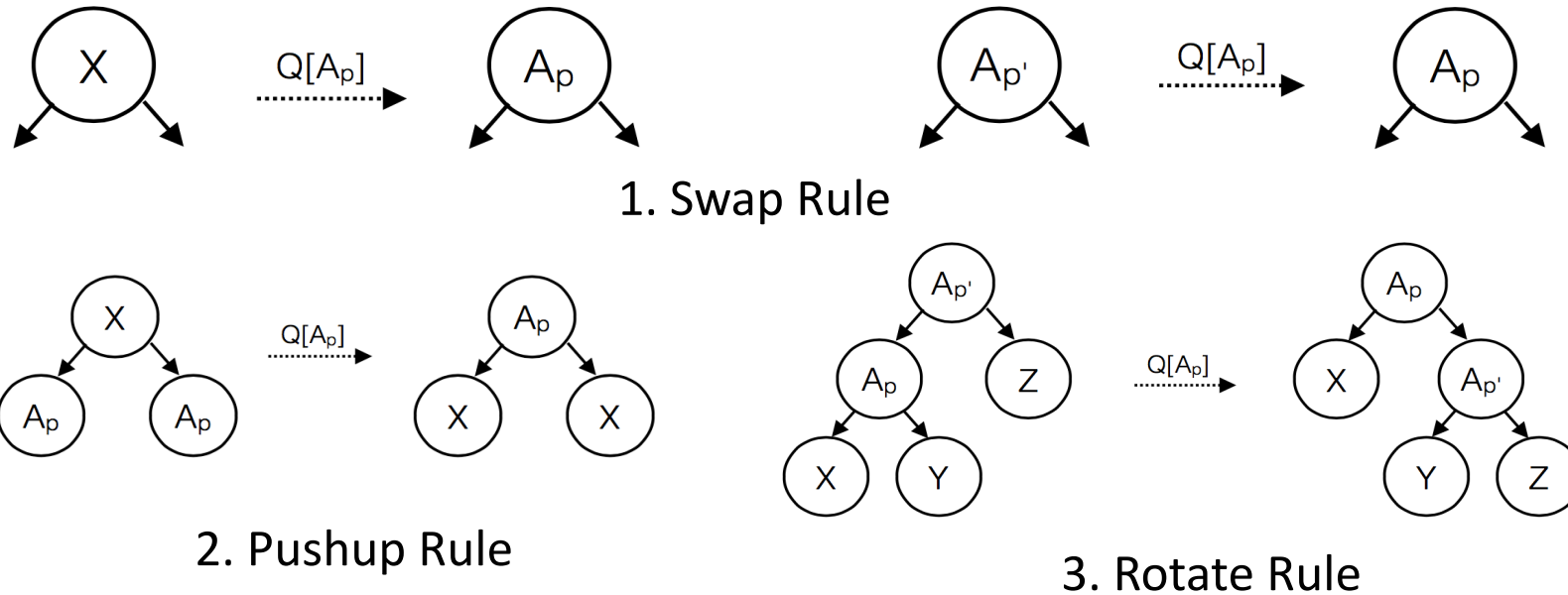
$$= \sum_{\text{all nodes } i} n_{ij} c_{ij}$$



2. Adaptive Query Executor

Goal: Return matching tuples + check if partitioning layout can be improved

Alternatives found via transformations on the partitioning tree



Getting a plan

Algorithm 2: getSubtreePlan

```
Input : Node node, Predicate pred
1 if isLeaf(node) then
2   return  $P_{none}(node)$ ;
3 else
4   if isLeftAccessed(node) then
5     leftPlan ← getSubtreePlan(node.lChild, pred);
6   if isRightAccessed(node) then
7     rightPlan ← getSubtreePlan(node.rChild, pred)
8   /* consider swap */
9   if leftPlan.fullyAccessed and rightPlan.fullyAccessed then
10    currentCost ←  $\sum_i \text{Cost}(node, q_i)$ ;
11    whatIfNode ← clone(node);
12    whatIfNode.predicate ← newPred;
13    swapNode(node, whatIfNode);
14    newCost ←  $\sum_i \text{Cost}(\text{whatIfNode}, q_i)$ ;
15    benefit = currentCost - newCost;
16    if benefit > 0 then
17      updatePlanIfBetter(node,  $P_{swap}(node, \text{whatIfNode})$ );
18  /* consider pushup */
19  if leftPlan.ptop and rightPlan.ptop then
20    updatePlanIfBetter(node,  $P_{pushup}(node, node.lChild, node.rChild)$ );
21  /* consider rotate */
22  if node.attribute == predicate.attribute then
23    if leftPlan.ptop then
24      updatePlanIfBetter(node,  $P_{rotate}(node, node.lChild)$ );
25    if rightPlan.ptop then
26      updatePlanIfBetter(node,  $P_{rotate}(node, node.rChild)$ );
27  /* consider doing nothing */
28  updatePlanIfBetter(node,  $P_{none}(node)$ );
29  return node.plan;
```

Algorithm 3: getBestPlan

Input : Tree *tree*, Predicate[] *predicates*

```
1 while predicates ≠  $\phi$  do
2   prevPlan ← tree.plan;
3   foreach p in predicates do
4     Plan newPlan ← getSubtreePlan(tree.root, p);
5     updatePlan(tree.root, newPlan, p);
6   if tree.plan ≠ prevPlan then
7     remove from predicates the newly inserted predicate;
8   else
9     break;
10 if tree.plan.benefit > tree.plan.cost then
11   return tree.plan;
12 else
13   return null;
```

Cost Model

The system maintain a window W of past queries

Compute Benefit and Repartitioning Cost for the best plan

Repartitioning **ONLY** happens when reduction in the total cost of the query workload is greater than re-partitioning cost.

Solves constant re-partitioning due to random query sequences and bounds the worse case impact.

$$\text{Cost}(T, q) = \sum_{b \in \text{lookup}(T, q)} n_b$$

$$\text{Benefit}(T') = \sum_{q \in W} \text{Cost}(T, q) - \sum_{q \in W} \text{Cost}(T', q)$$

$$\text{RepartitioningCost}(T, q) = \sum_{b \in B} c \cdot n_b$$

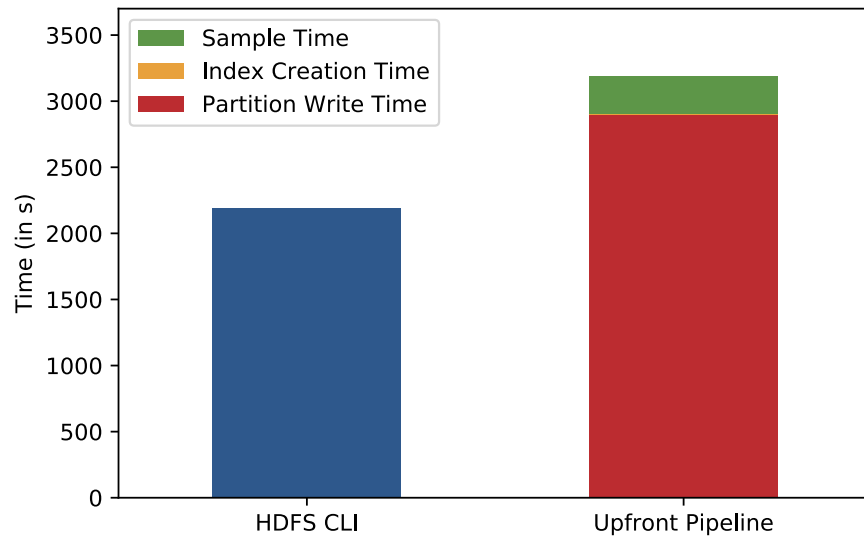
Performance

4 metrics

- 1) Load time
- 2) Time taken by first query
- 3) Aggregate runtime over a workload
- 4) Incremental improvement with workload hints

Load Time

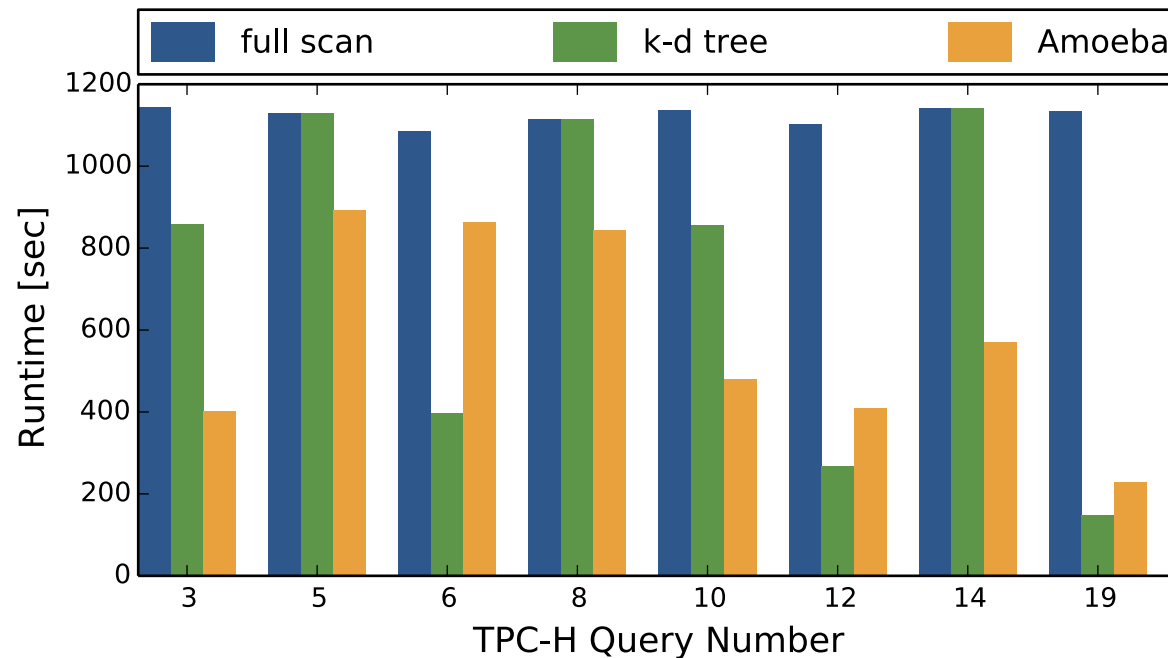
TPC-H: Scale Factor 200 + De-normalized. Data size: **1.4TB**



Loading performance: 1.38 times slower than HDFS

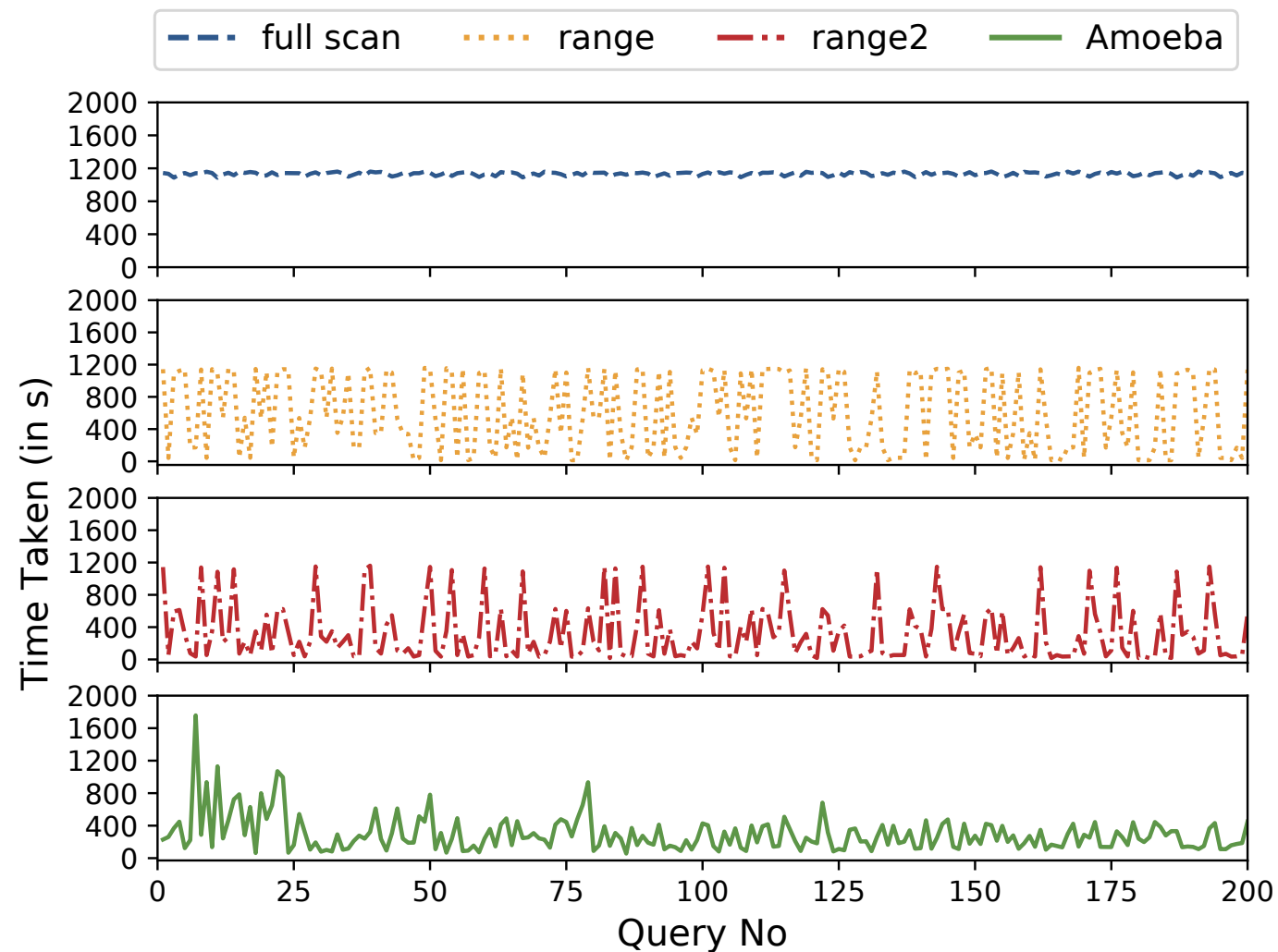
Load time scales almost linearly
with data size and independent of number of columns

Time taken by first query



On Average: 45% better than full scan
20% better than k-d tree

Aggregate Workload Runtime



Workload: 200 Queries generated from random initialization of 8 query templates of TPC-H benchmark

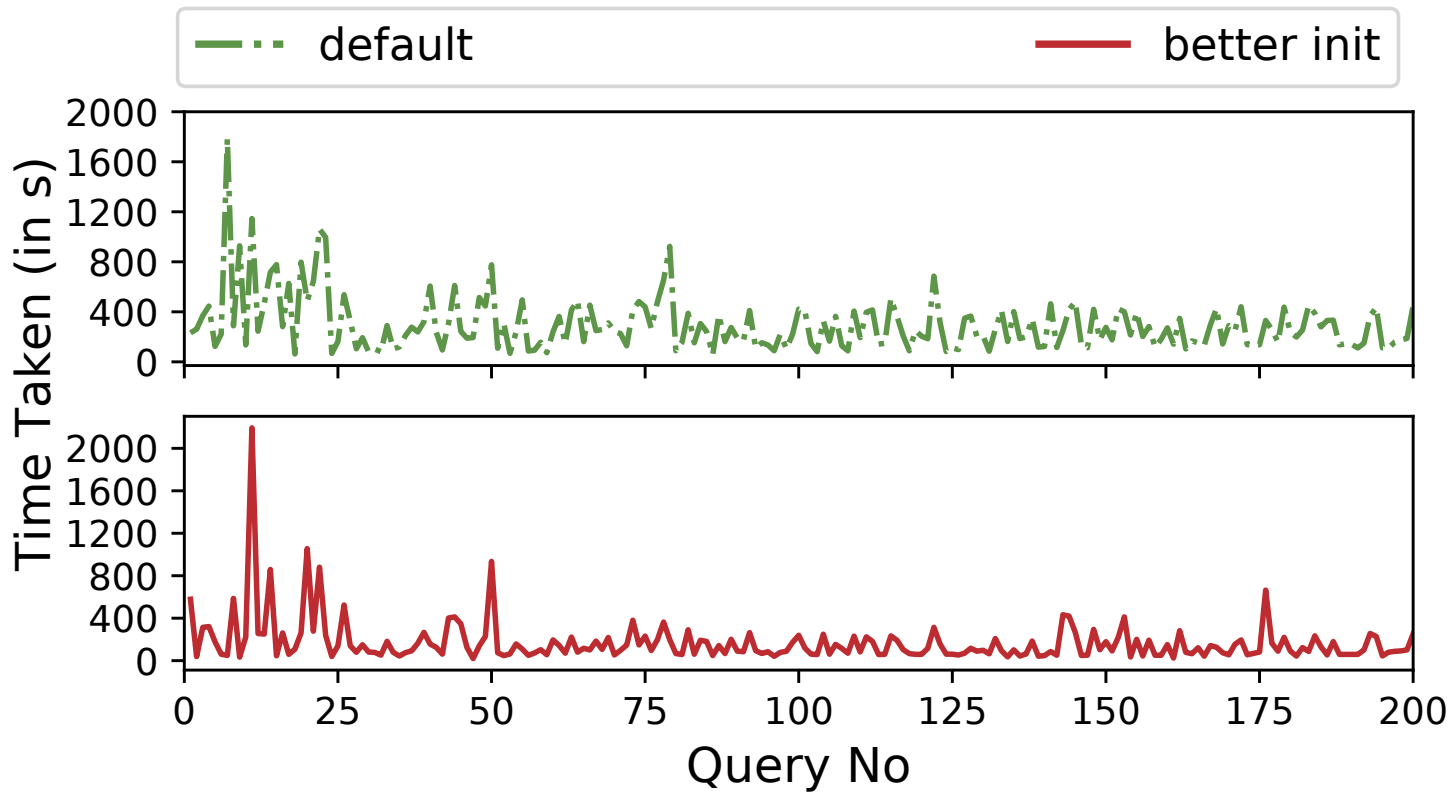
full scan – Baseline

range – partitions on orderdate (1 per date)
1.88x better

range2 – partitions on orderdate(64),
r_name(4),c_mktsegment(4),quantity(8)
3.48x better

Amoeba – 3.84x better than baseline

Workload Hints



Better Init:

Starts with custom allocation to
mimic *range2*

6.67x better than *fullscan*

Filtering ratio:

default : 0.81

better init : 0.9

Conclusion

- Amoeba is a distributed storage system based on an adaptive data partitioning scheme
 - Low loading overhead
 - Improved first query performance
 - Adapt to changes and significantly improvement to workload runtime
 - Can exploit workload hints
- Allows analysts to get started right away and reap benefits of partitioning without an upfront workload