

Distribution-Oblivious Databases

DBDBD, 2013-11-29

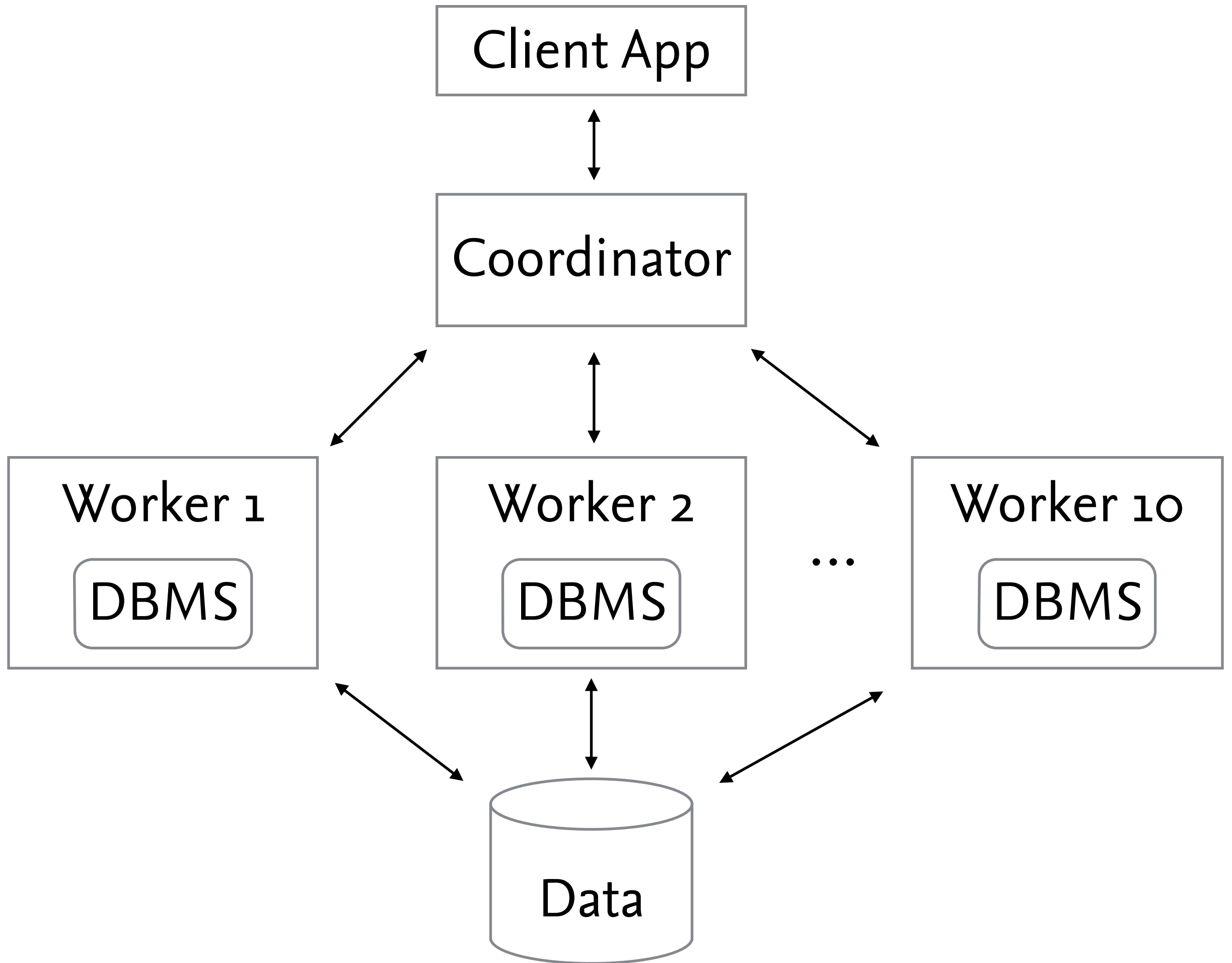
Environment

- Analytical queries and growing data volume
 - but scale-up (bigger box) limited (\$\$\$)
- SQL scale-out (more boxes)?
 - a.k.a. “making complex pieces of software even more complex”

Distribution “Fun”

- Task splitting
- Data placement
- Load Balancing
- Group Membership

Can we run SQL databases oblivious to distributed processing and still get scale-out efficiency benefits?



Client App

```
SELECT SUM(l_extendedprice *  
l_discount) FROM lineitem;
```

Coordinator

```
SELECT  
SUM(l_extendedprice *  
l_discount) FROM  
lineitem WHERE  
l_orderkey >= 0 AND  
l_orderkey <= 6000;
```

```
SELECT SUM(l_extendedprice  
* l_discount) FROM lineitem  
WHERE l_orderkey > 6000 AND  
l_orderkey <= 12000;
```

```
SELECT SUM(l_extendedprice  
* l_discount) FROM lineitem  
WHERE l_orderkey > 54000  
AND l_orderkey <= 60000;
```

Worker 1

DBMS

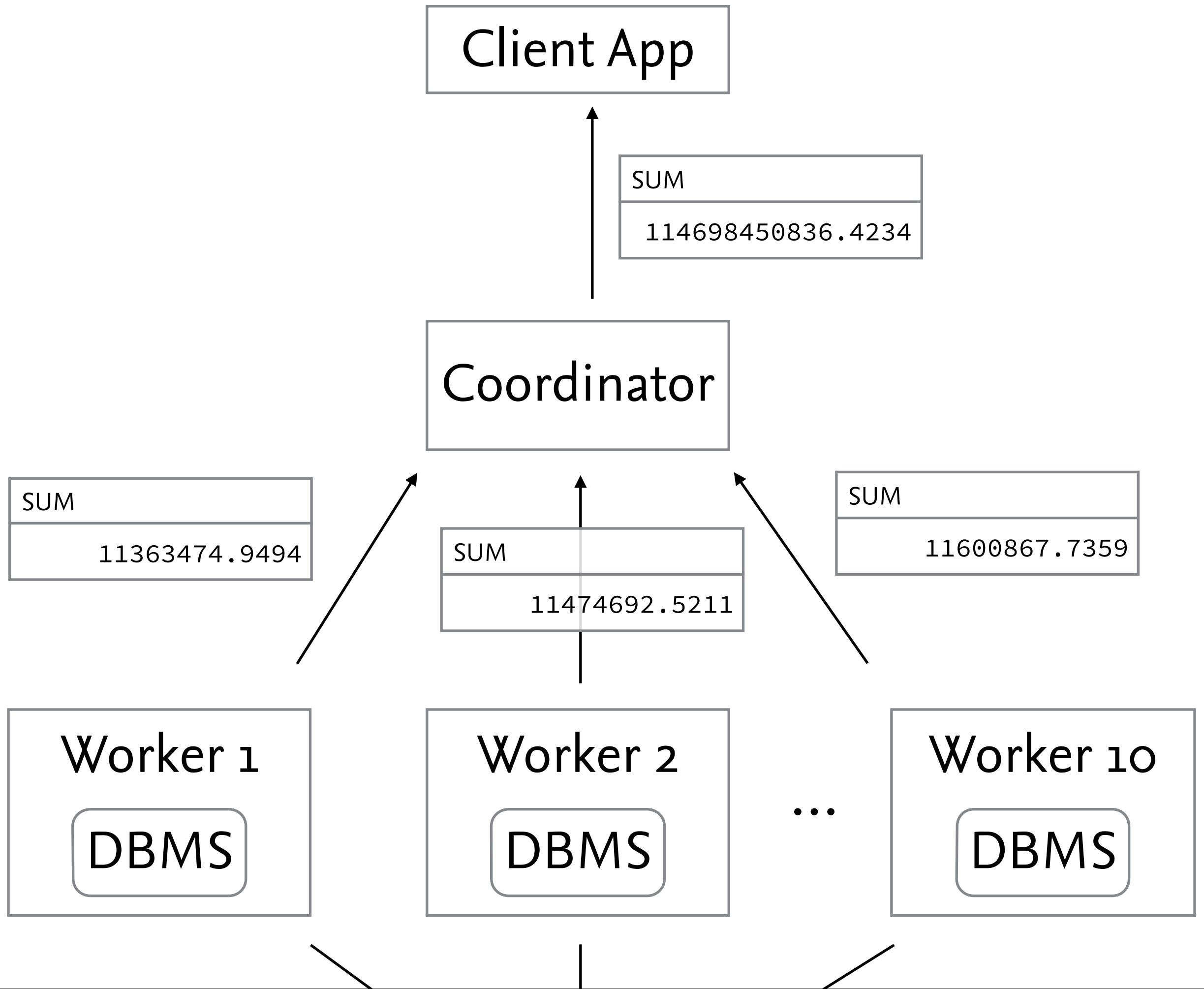
Worker 2

DBMS

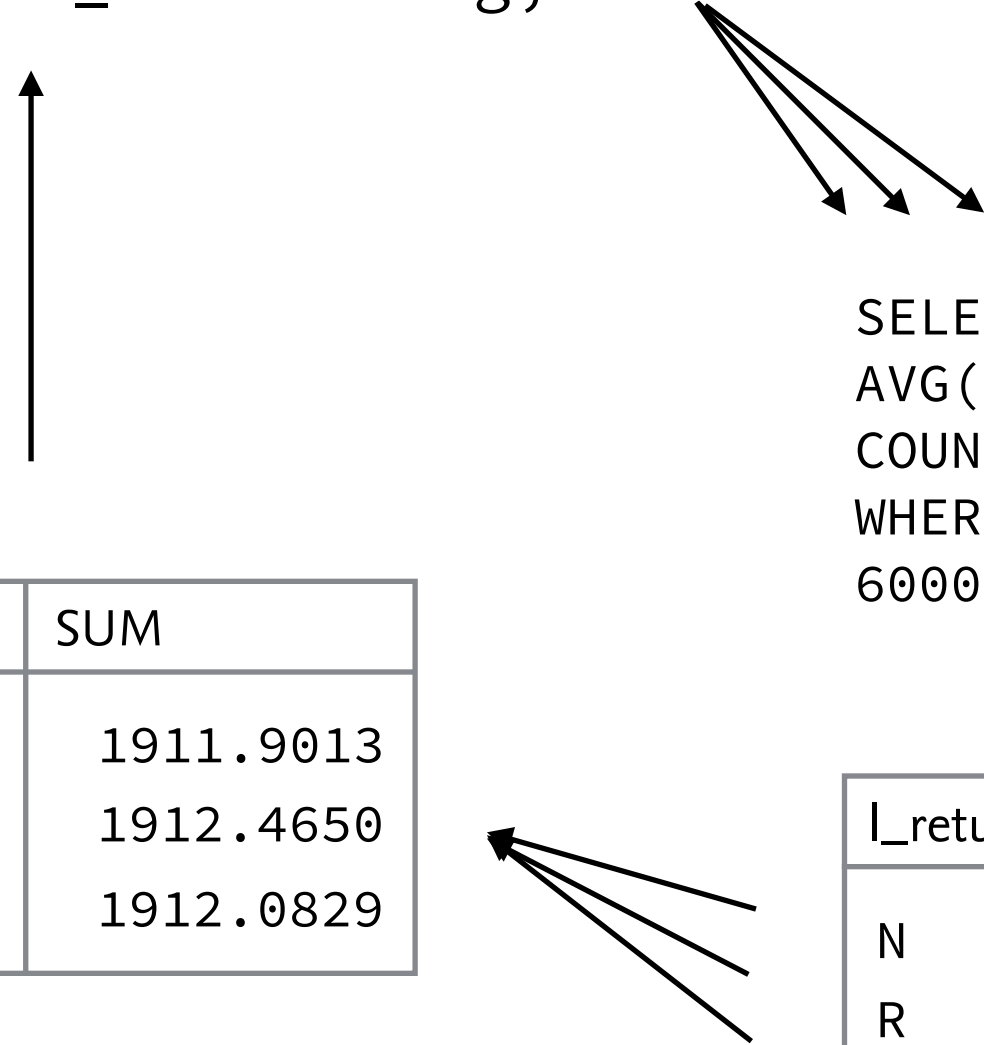
...

Worker 10

DBMS




```
SELECT l_returnflag,  
AVG(l_extendedprice *  
l_discount) FROM lineitem  
GROUP BY l_returnflag;
```



l_returnflag	SUM
N	1911.9013
R	1912.4650
A	1912.0829

```
SELECT l_returnflag,  
AVG(l_extendedprice * l_discount),  
COUNT(*) AS __groupcount FROM lineitem  
WHERE l_orderkey >= 0 AND l_orderkey <=  
6000 GROUP BY l_returnflag;
```

l_returnflag	SUM	__groupcount
N	1885.3028	3077
R	1867.0046	1459
A	1915.2755	1482

SQL level parallelism

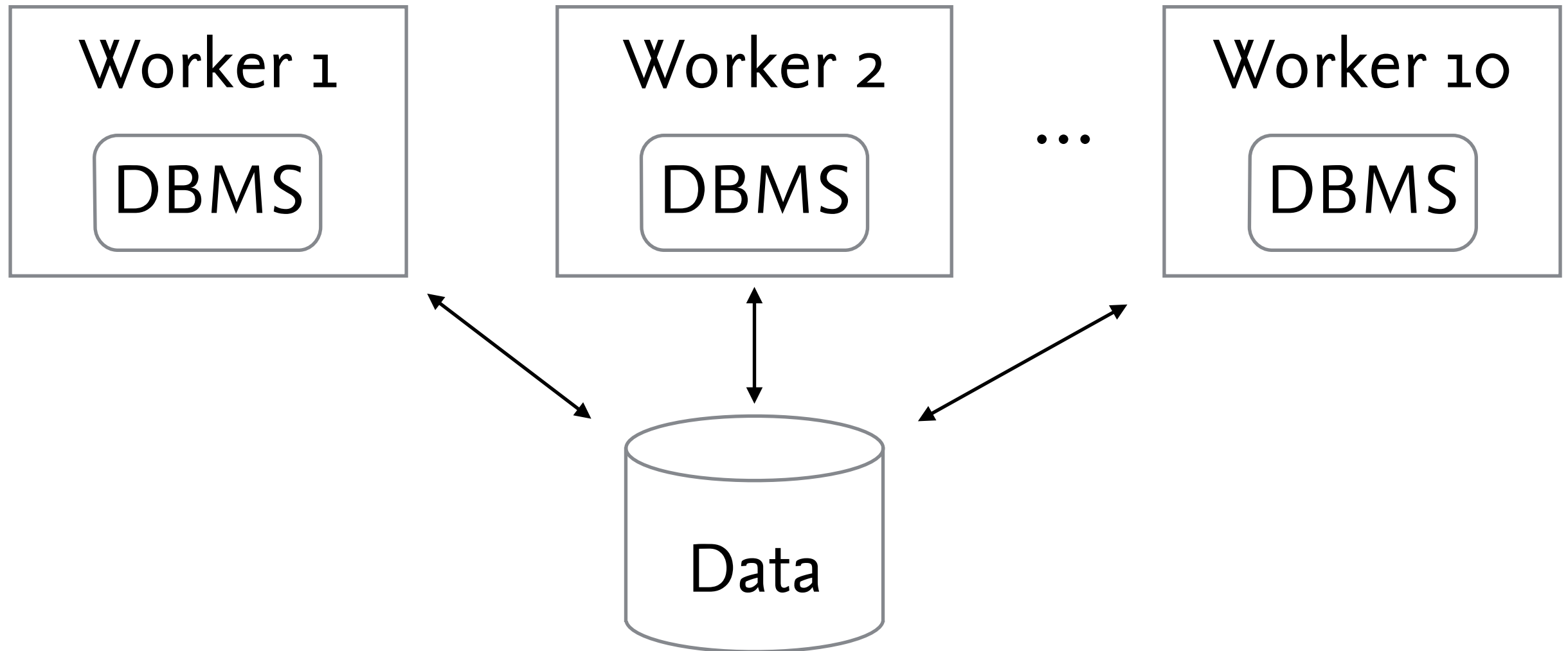
- Coordinator receives queries, splits them up into smaller subqueries
- Workers calculate result for subqueries, send result to coordinator
- Coordinator merges and delivers results

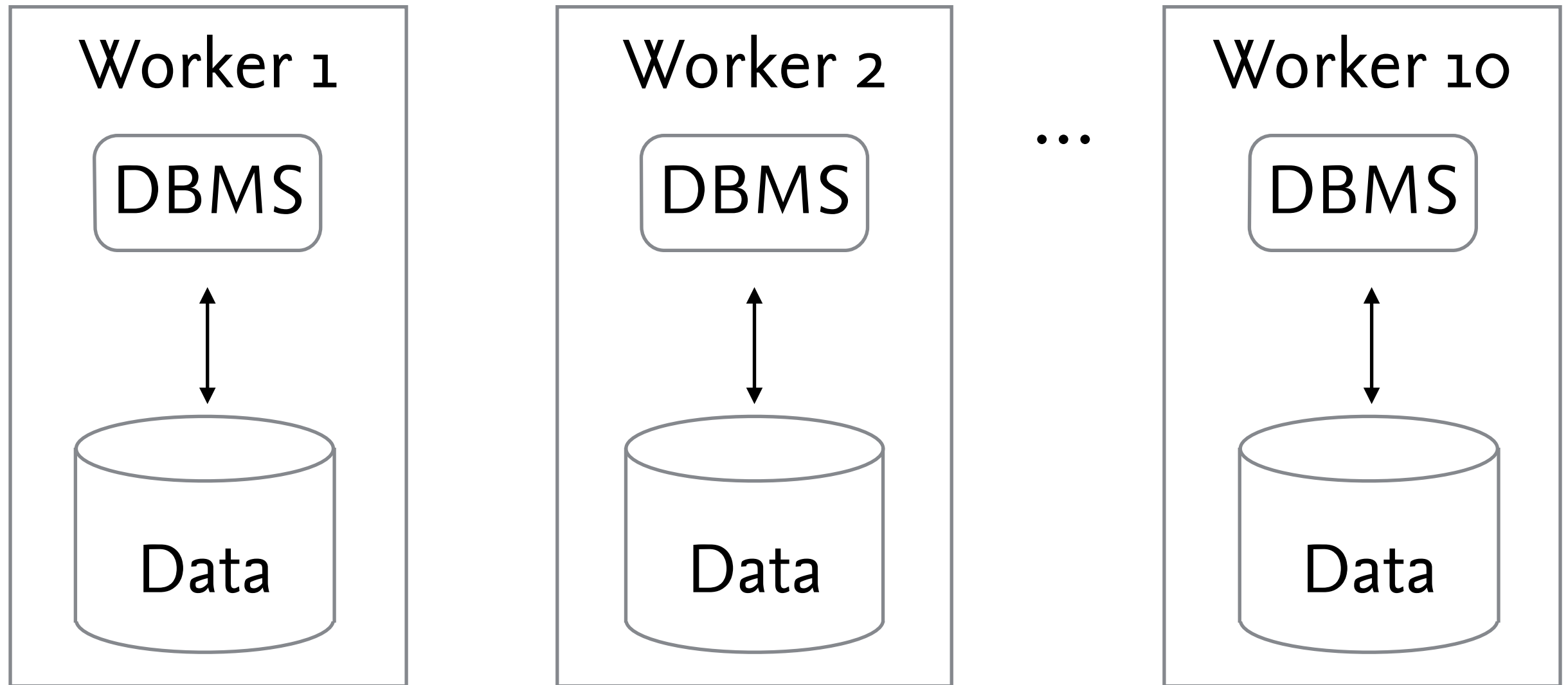
Yang, Yen, Tan & Madden:

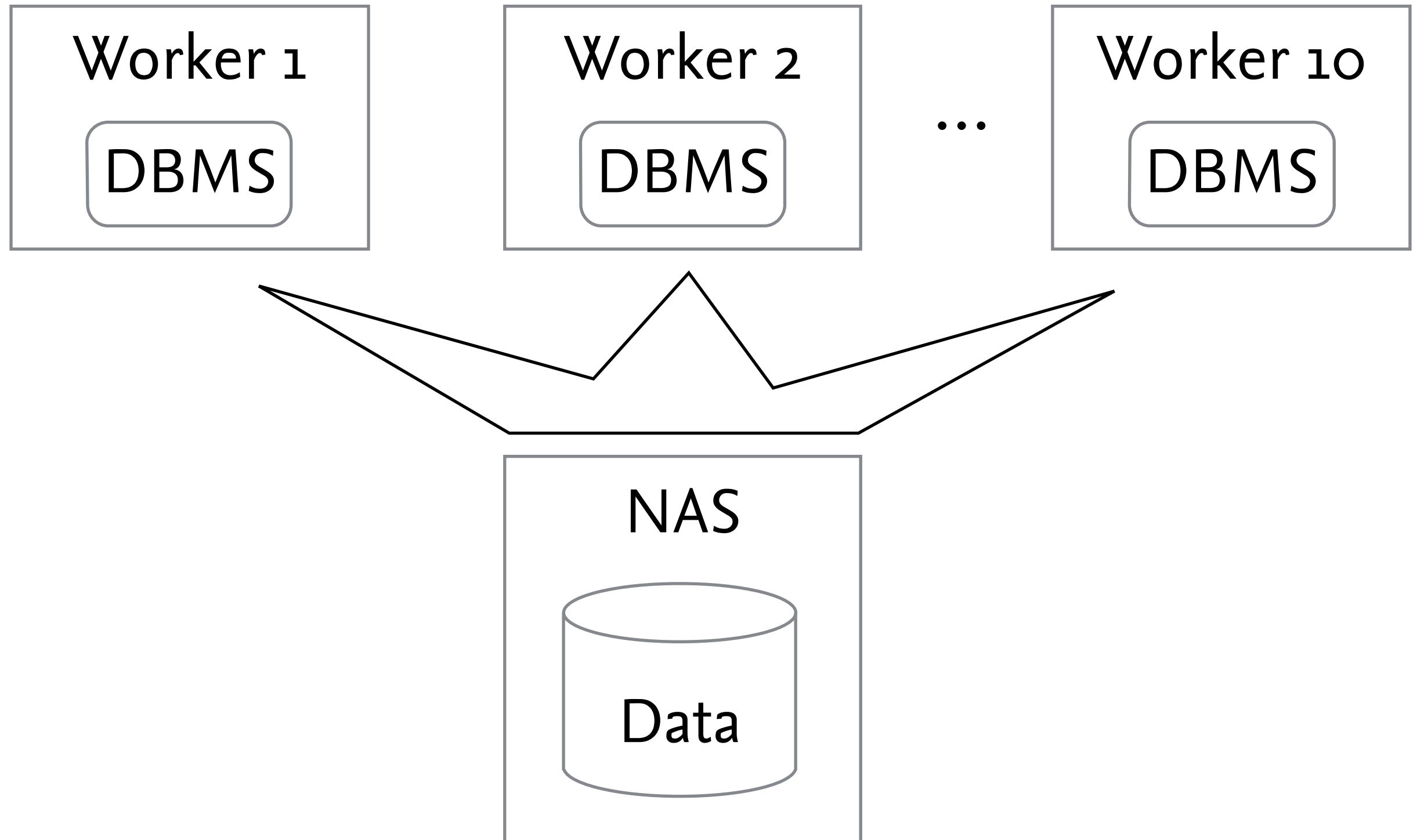
Osprey: Implementing MapReduce-style fault tolerance in a shared-nothing distributed database

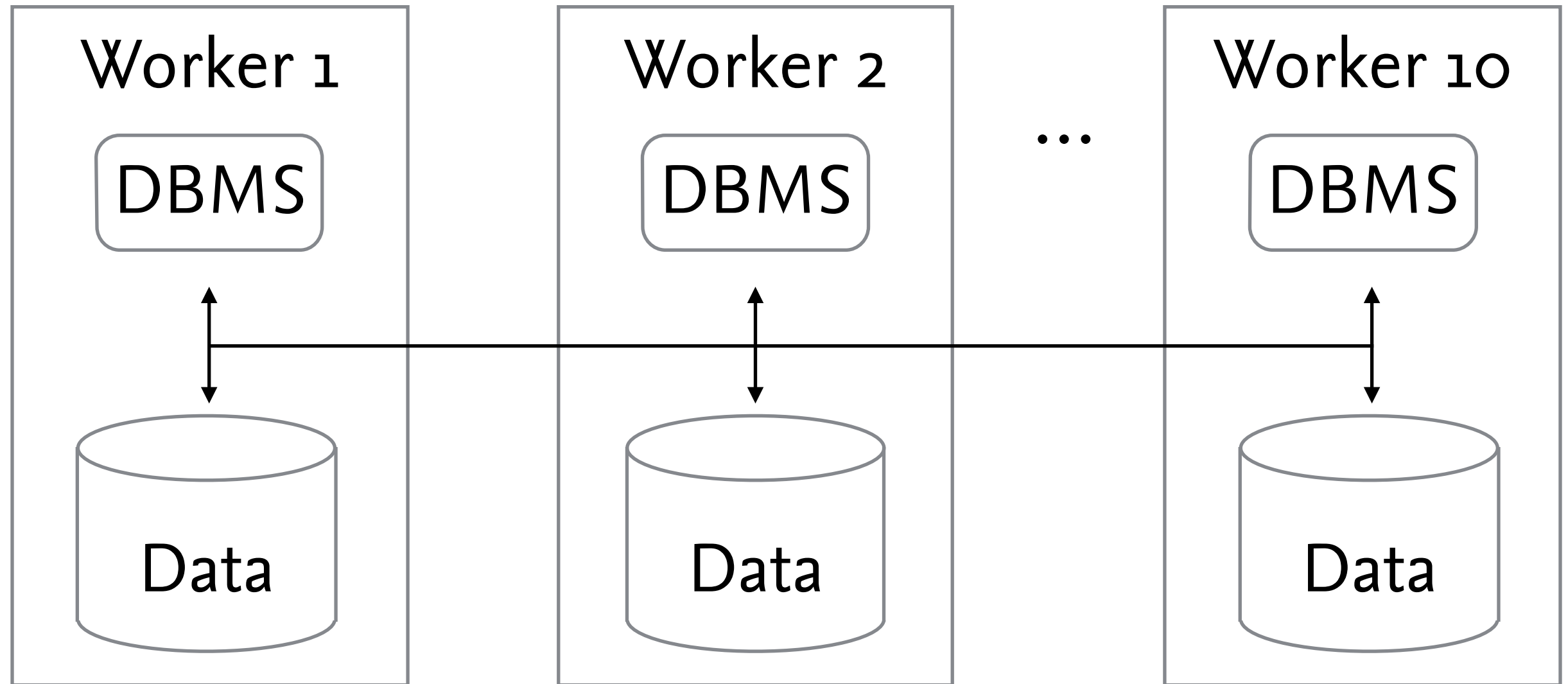
Lima, Mattoso & Valduriez:

Adaptive Virtual Partitioning for OLAP Query Processing in a Database Cluster









Caches?

- What about subquery locality?
- DB Buffer Mgr/OS implement LRU caches for DB data
- Can we predict FS IO from SQL queries?
 - We do not have to! All we need is mapping consistency.

Q₁

SELECT MIN(C2) FROM T1,T2
WHERE C5 > 42;

T₁

<u>C1</u>	C2	C3	C4

T₂

C2	C5

Q₂

SELECT MAX(C2) FROM T1,T3
WHERE C6 < 84;

T₁

<u>C1</u>	C2	C3	C4

T₃

C2	C6

$Q_{1/3}$

```
SELECT MIN(C2) FROM T1,T2
      WHERE C5 > 42
AND C1 >= 1000 AND C1 < 2000;
```

T₁

[illegible]

T₂

C ₂	C ₅

Q2/3

```
SELECT MAX(C2) FROM T1,T3
      WHERE C6 < 84
AND C1 >= 1000 AND C1 < 2000;
```

T1

<u>C₁</u>	C ₂	C ₃	C ₄

T₃

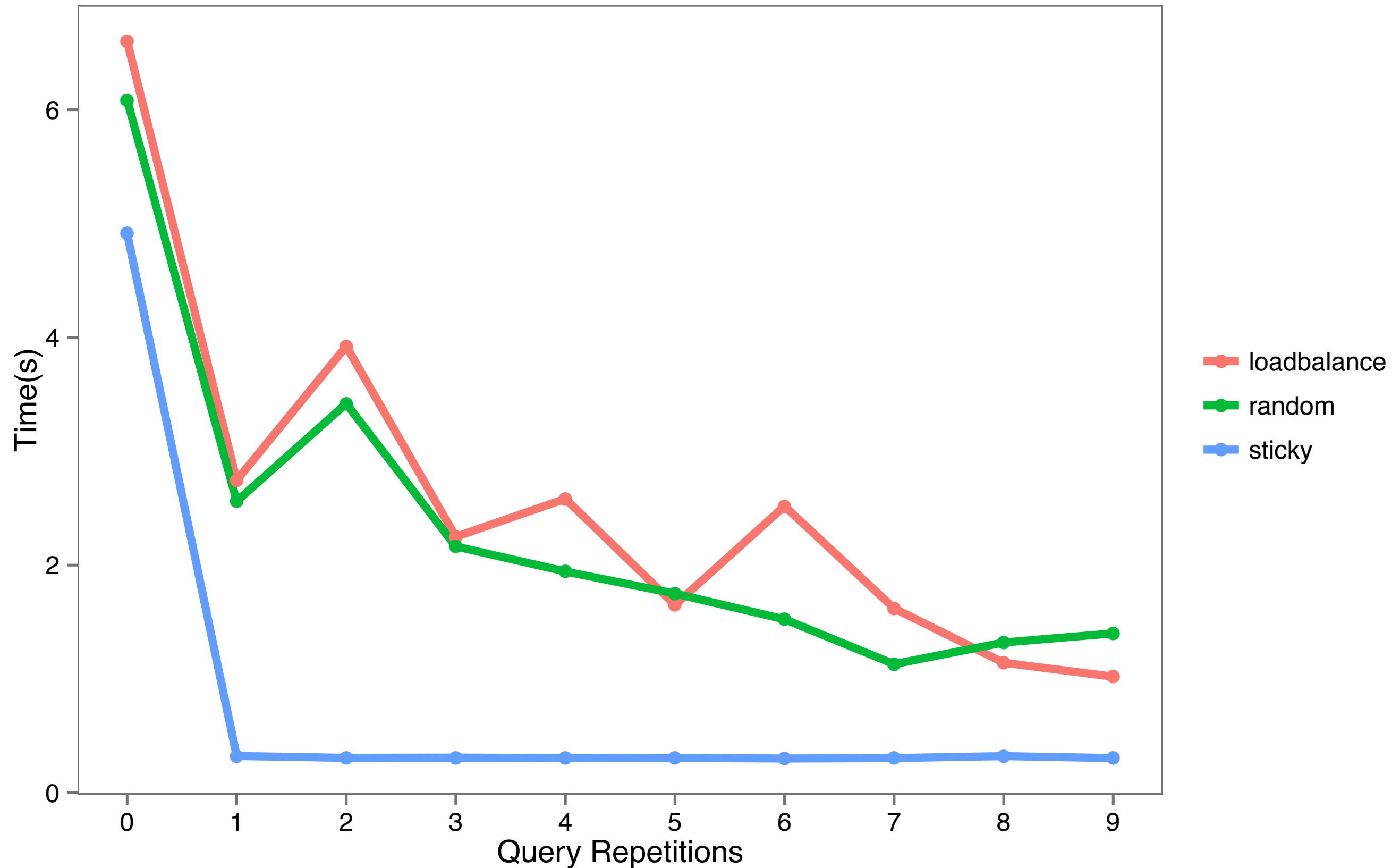
C2	C6

* Assuming some DBMS sanity

Subquery Assignment

- Hope: If we assign subqueries with equal fact table restrictions to same worker, we can make max. use of cached data
- Experiments on 10 nodes, SSBM Q₁₁, 10 times from cold start
 - Sticky
 - Random
 - Balance

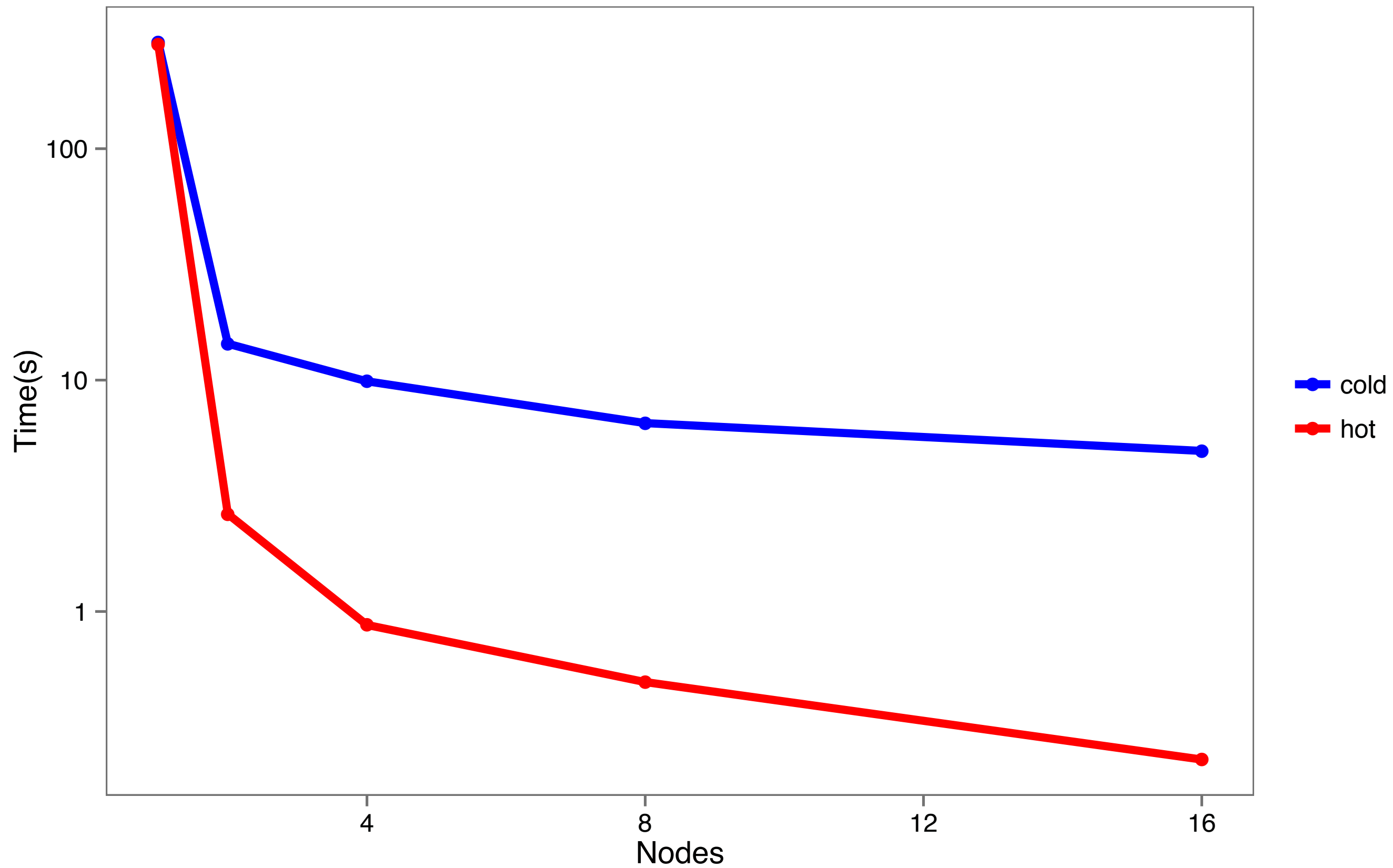
Subquery Assignment, SSBM Q₁₁, 10 Nodes



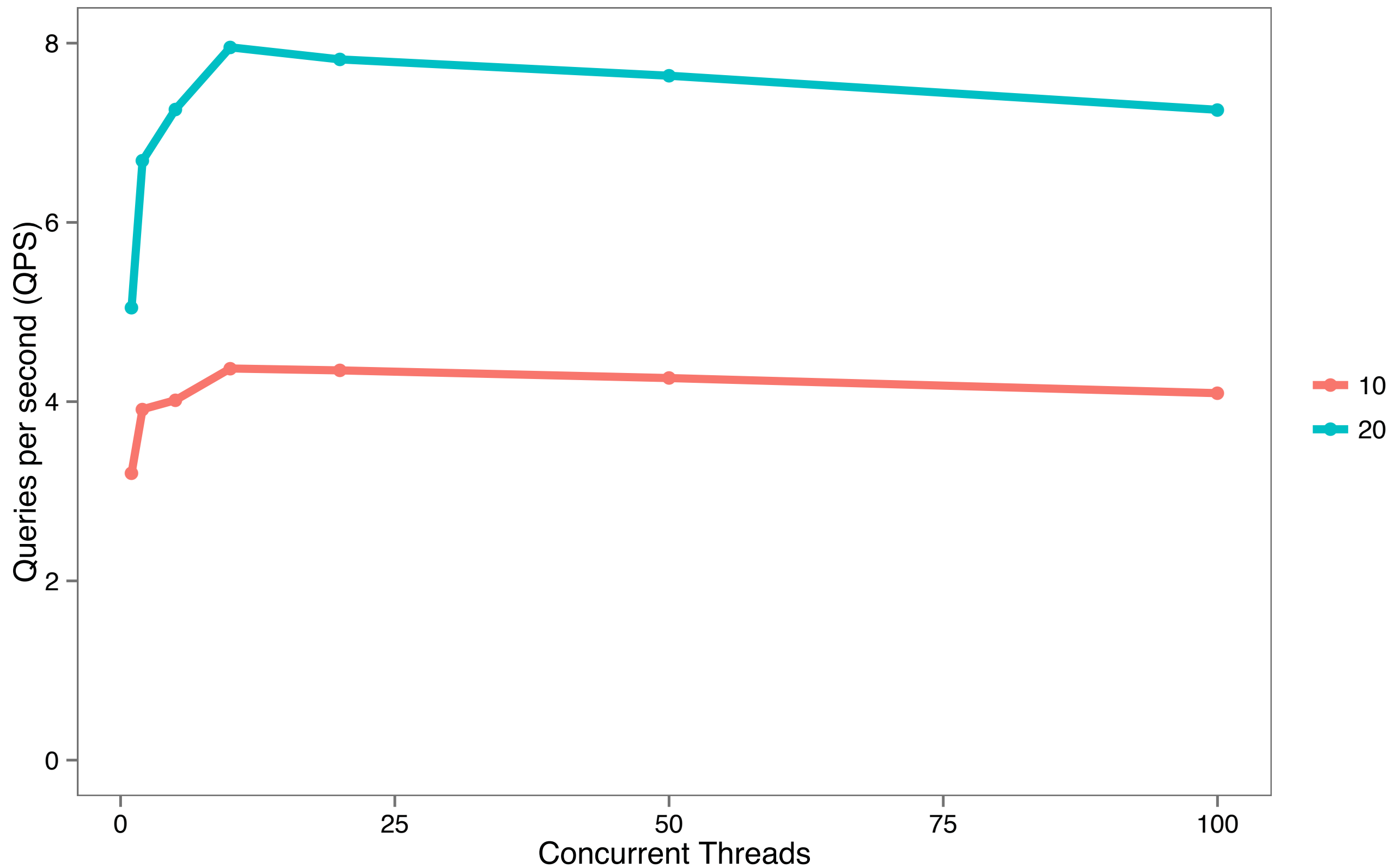
Scaleup

- Latency: More nodes, less query latency
- Throughput: More nodes, more parallel queries

Scaleup Latency, SSBM Q₁₁



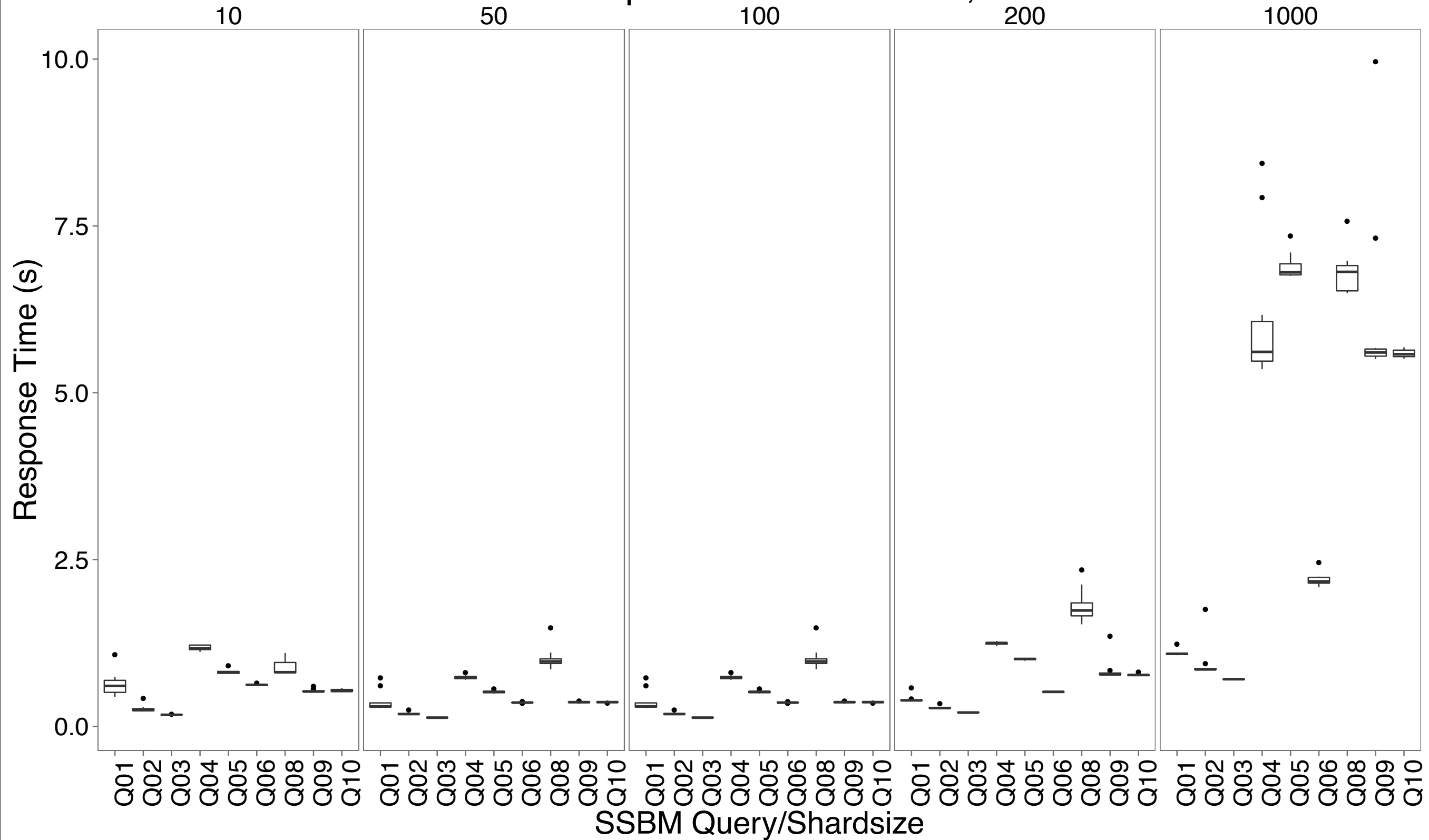
Scaleup Throughput, SSBM Q₁₁, Hot



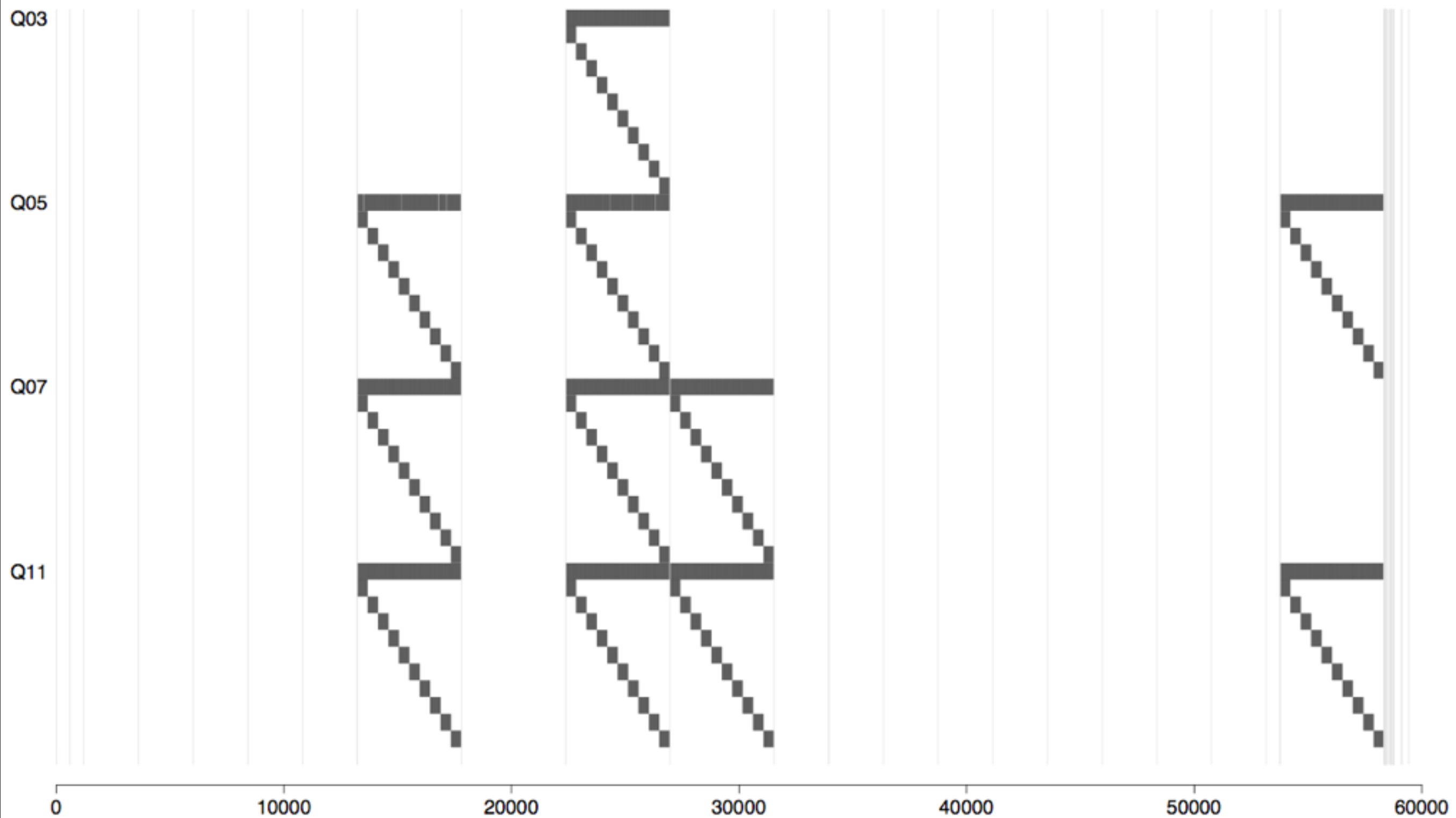
Thank You!

Questions?

Shard Size Impact – SSBM SF100, 10 rocks



File IO SSBM SF-100, MonetDB



File IO SSBM SF-100, PostgreSQL

