

NOSQL DATABASES YOU HAVE TO KNOW

@TOM BUJOK



NOSQL DATABASES YOU HAVE TO KNOW

ABOUT ME

TOMEK BUJOK

HAZELCAST

@TOMBUJOK

WWW.REFICIO.ORG



Babun

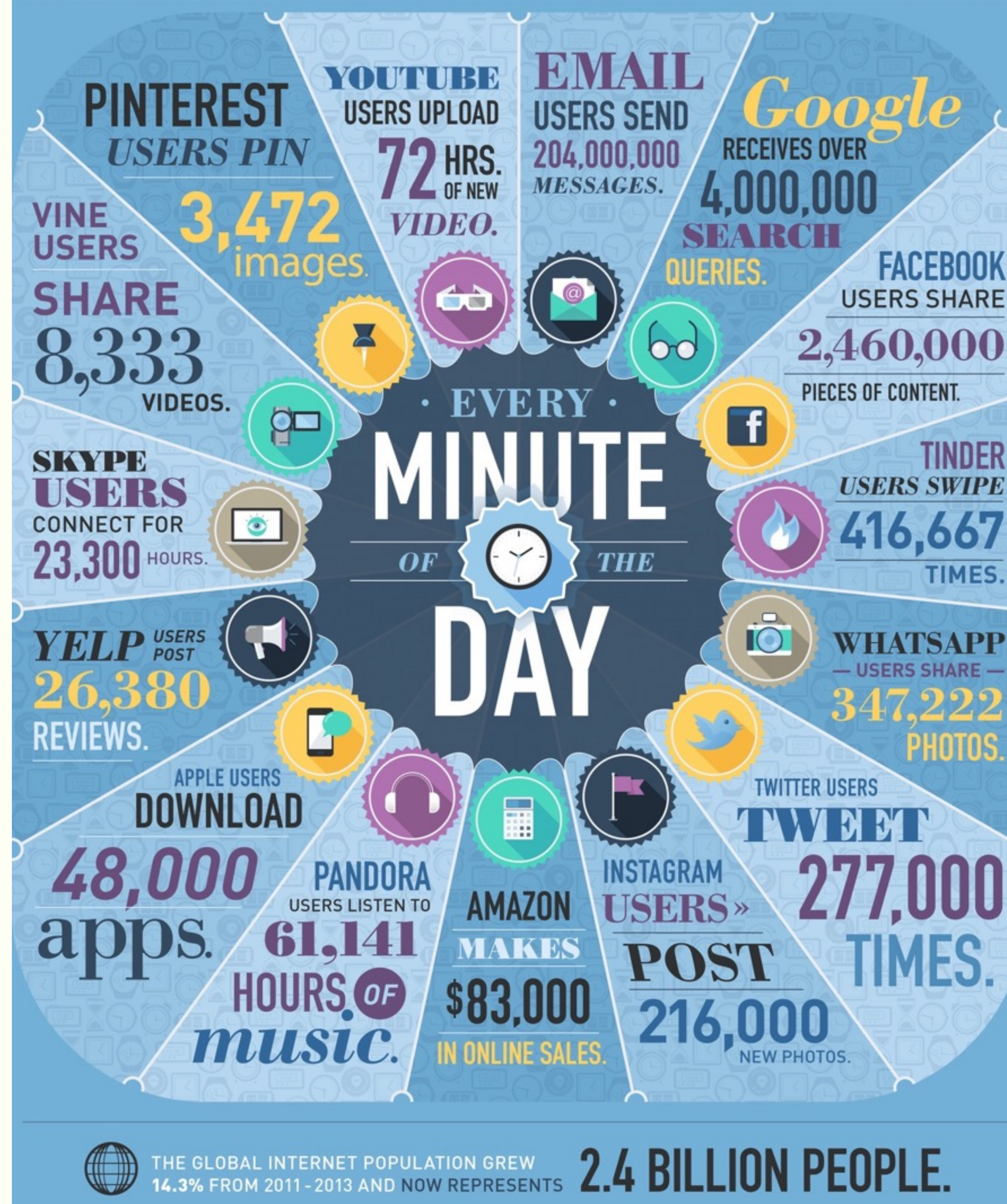
a Windows shell you will love!

 Download now

 Visit us on GitHub

(NOT) YET ANOTHER NOSQL TALK





FACEBOOK SCALE

- **1+ BILLION DAILY ACTIVE USERS** (04.2015)
- **600 TB DAILY ASYNCHRONOUS REPLICATION BETWEEN DATA CENTERS** (04.2014)
- **300+ PB IN TOTAL** (04.2014)
- **BILLIONS OF QUERIES PER SECOND** (09.2014)

CAN RDBMS HANDLE IT?

▸ RELATIONAL MODEL

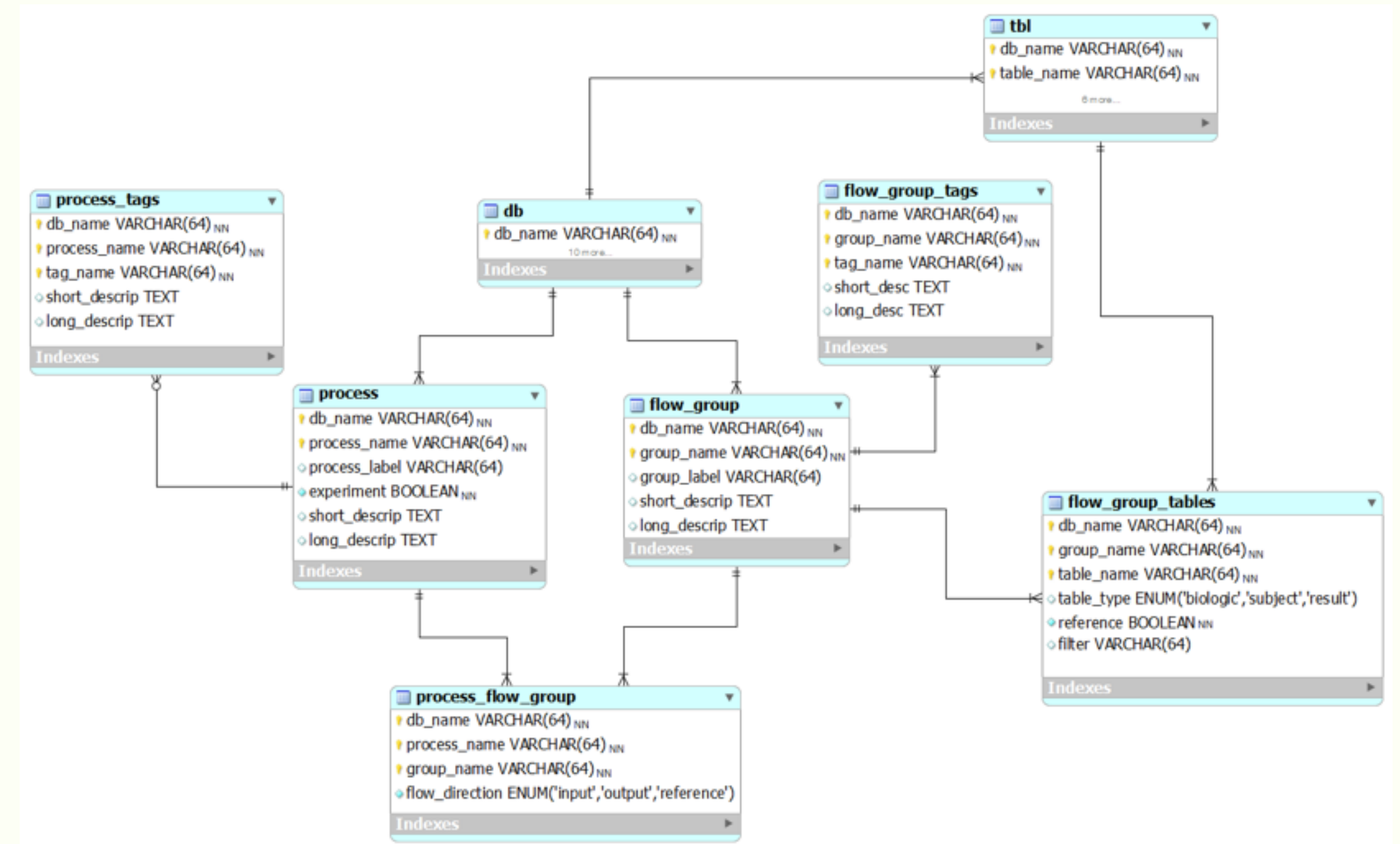
▸ ACID

▸ ATOMICITY - ALL OR NOTHING

▸ CONSISTENCY - ALWAYS VALID STATE

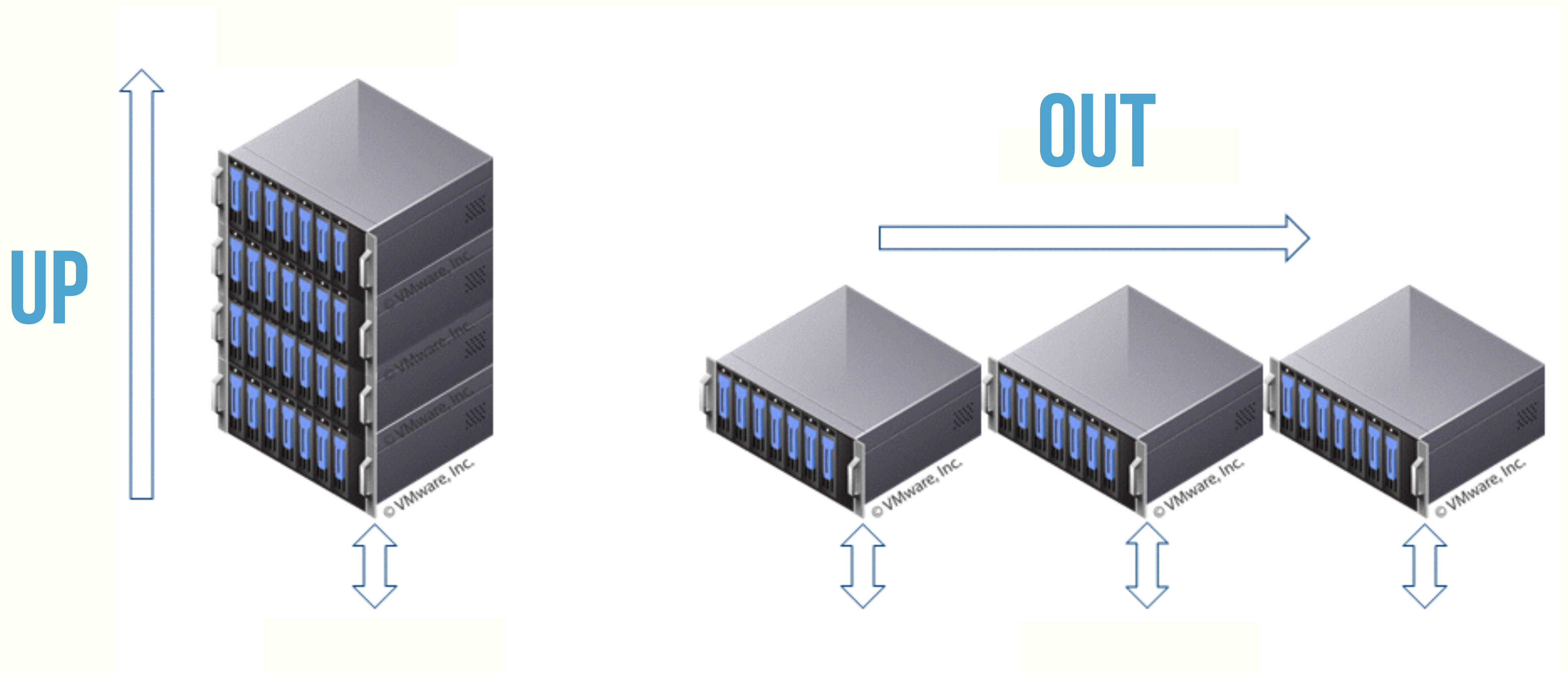
▸ ISOLATION - CONCURRENT TX -> STATE AS IF TX EXECUTED SERIALY

▸ DURABILITY - NO DATA LOSS AFTER TX COMMITED



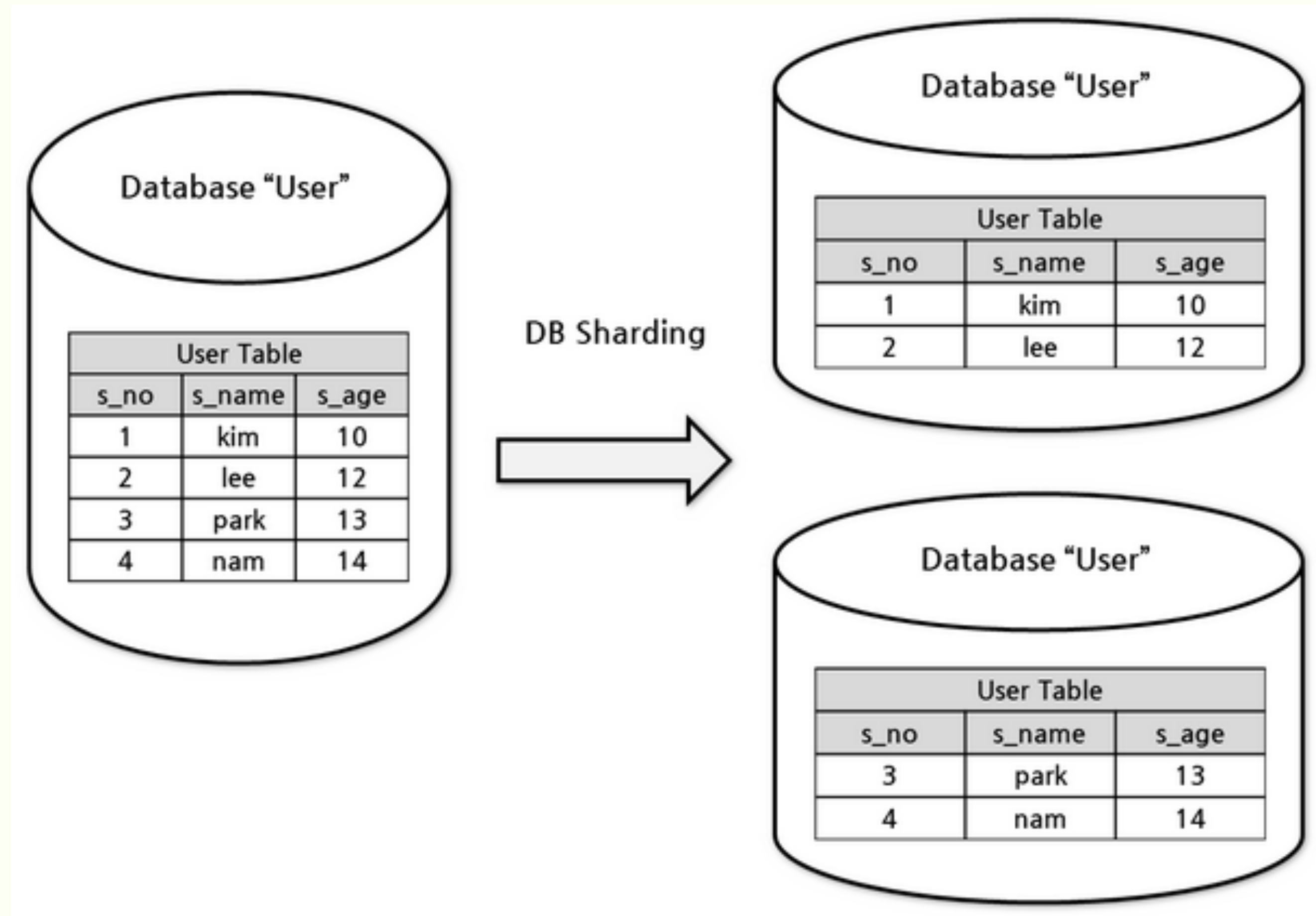
RDBMS “PROBLEMS”?

SCALING **UP** VS. SCALING **OUT**



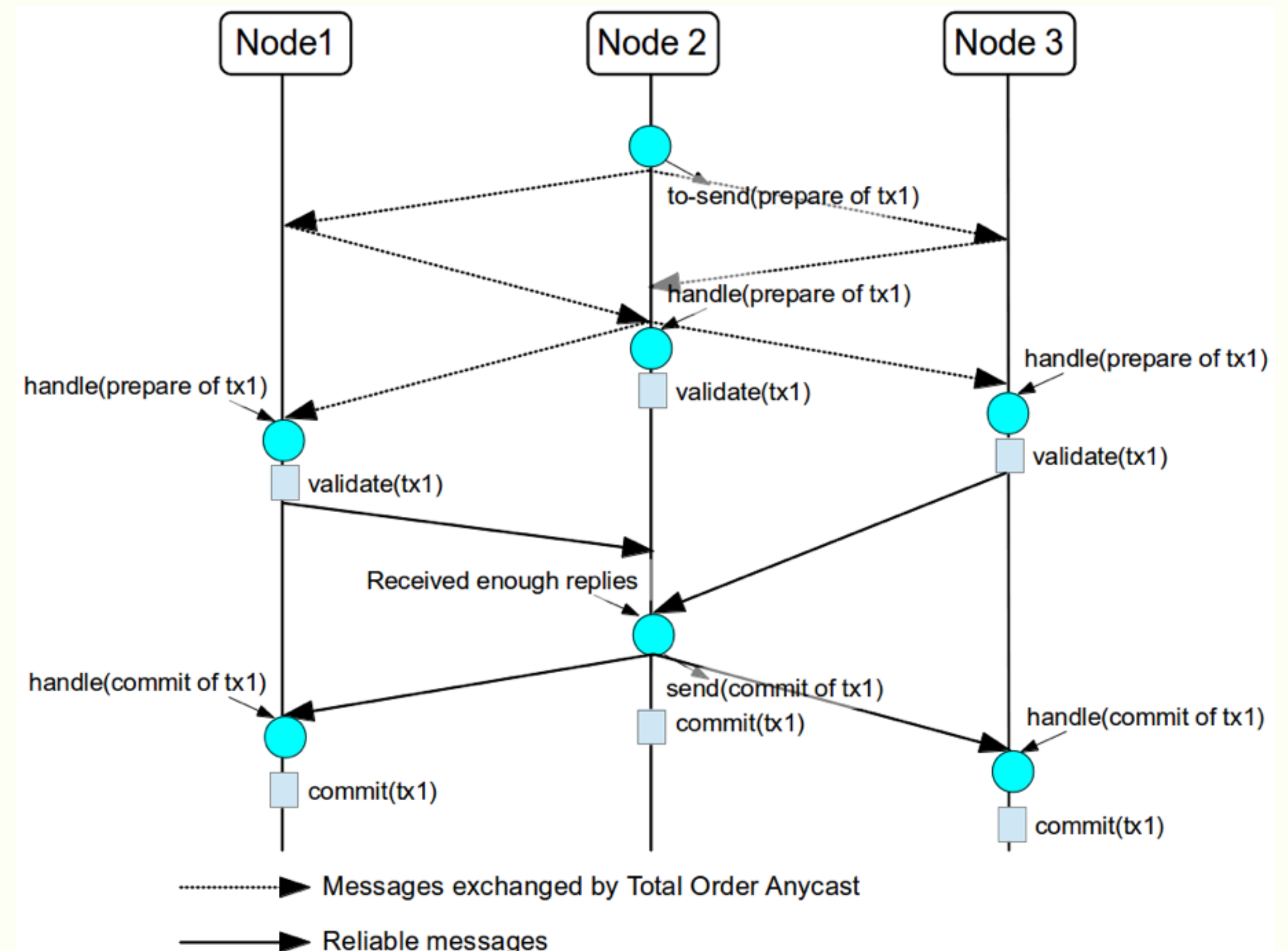
SHARDING / PARTITIONING

- ▶ DISTRIBUTED JOINS?
 - ▶ DISTRIBUTED REFERENTIAL?
- ## INTEGRITY?



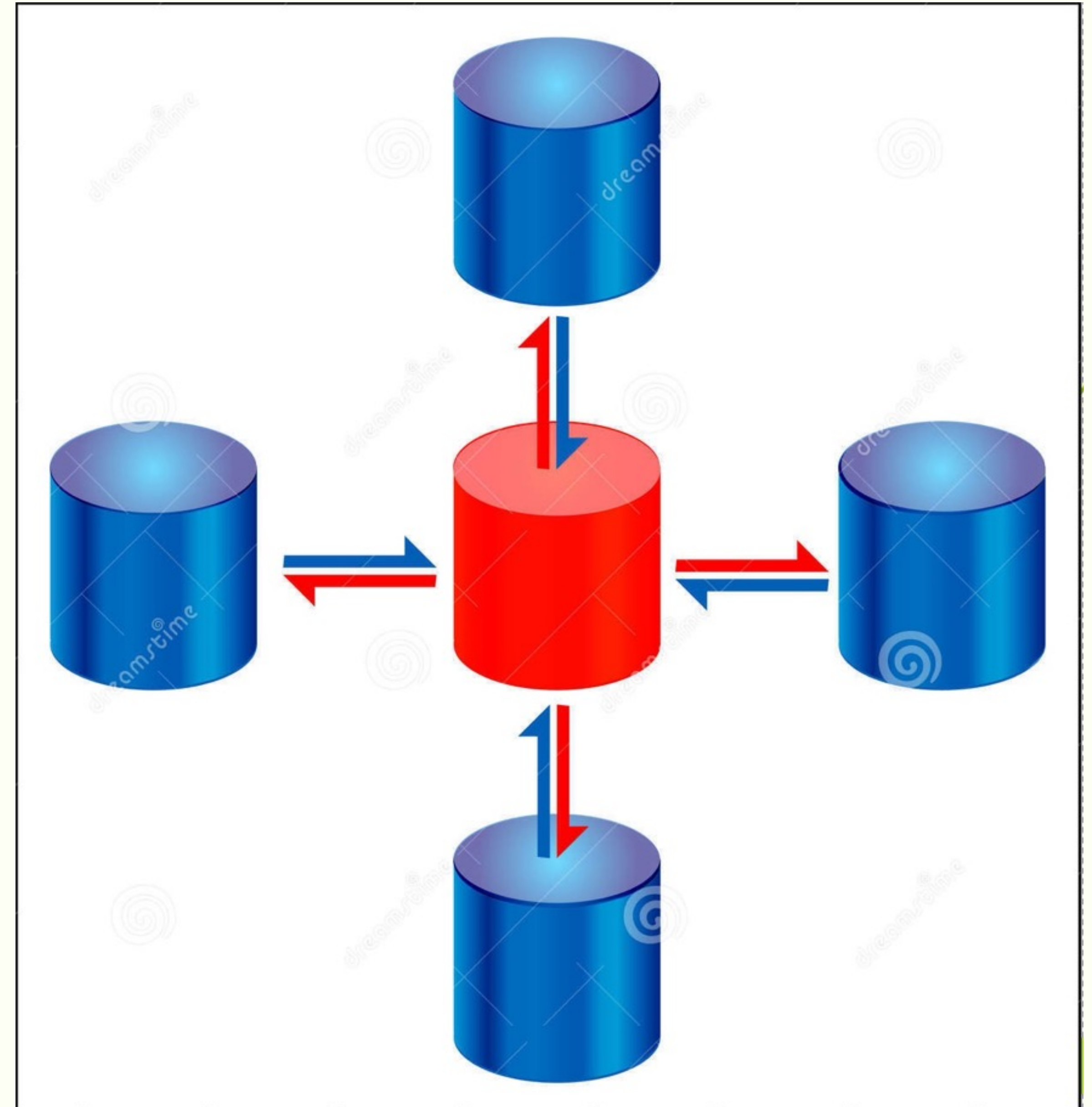
DISTRIBUTED TRANSACTIONS

- ATOMICITY
- LOCKS FOR THE WHOLE DURATION?



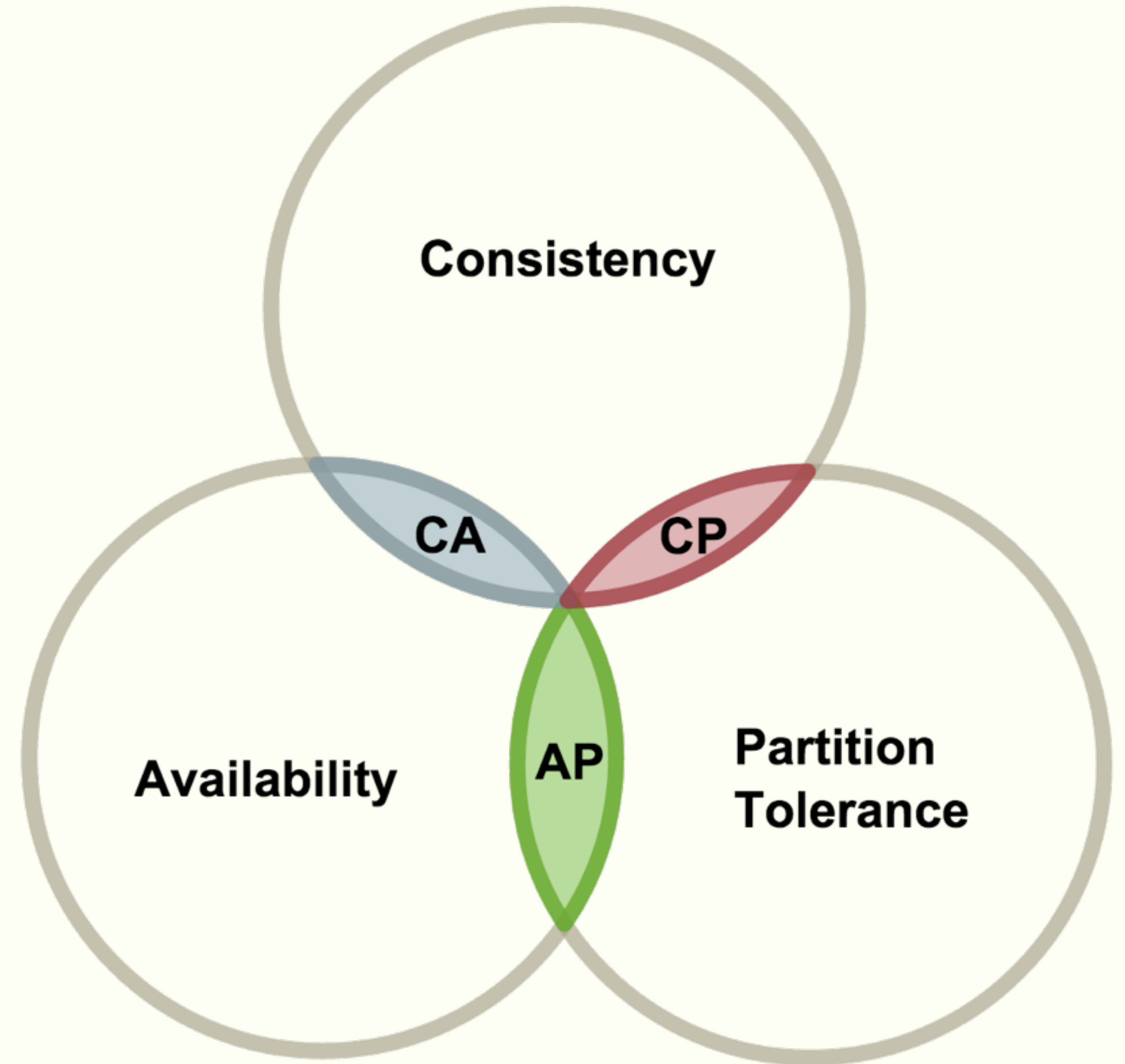
REPLICATION

- **CONSISTENCY -> HA AMONG REPLICAS**
- **ONLINE REPLICATION -> HIGH LATENCY**
- **POST-WRITE REPLICATION -> TRADEOFFS**



CAP THEOREM

- IN DISTRIBUTED SYSTEMS -> P IS OBLIGATORY
- IN CASE OF PARTITION
 - C vs A IS THE TRADE-OFF!
 - NOT A BINARY DECISION!



RDBMS TOTALLY SUCK!

FACEBOOK ARCHITECTURE

FACEBOOK ARCHITECTURE — MYSQL!

- NO JOINS OR LARGE SCANS
- FAVOR READ AND WRITE AVAILABILITY OVER ATOMICITY OR CONSISTENCY
- ASYNCHRONOUS REPLICATION BETWEEN DATA CENTERS (EVENTUAL CONSISTENCY)
- MULTI-SHARD TRANSACTIONS PROHIBITED
- TAO = DATA STORE FOR THE SOCIAL GRAPH (MYSQL + MEMCACHED)

RDBMS TOTALLY SUCK?
PRANKED!



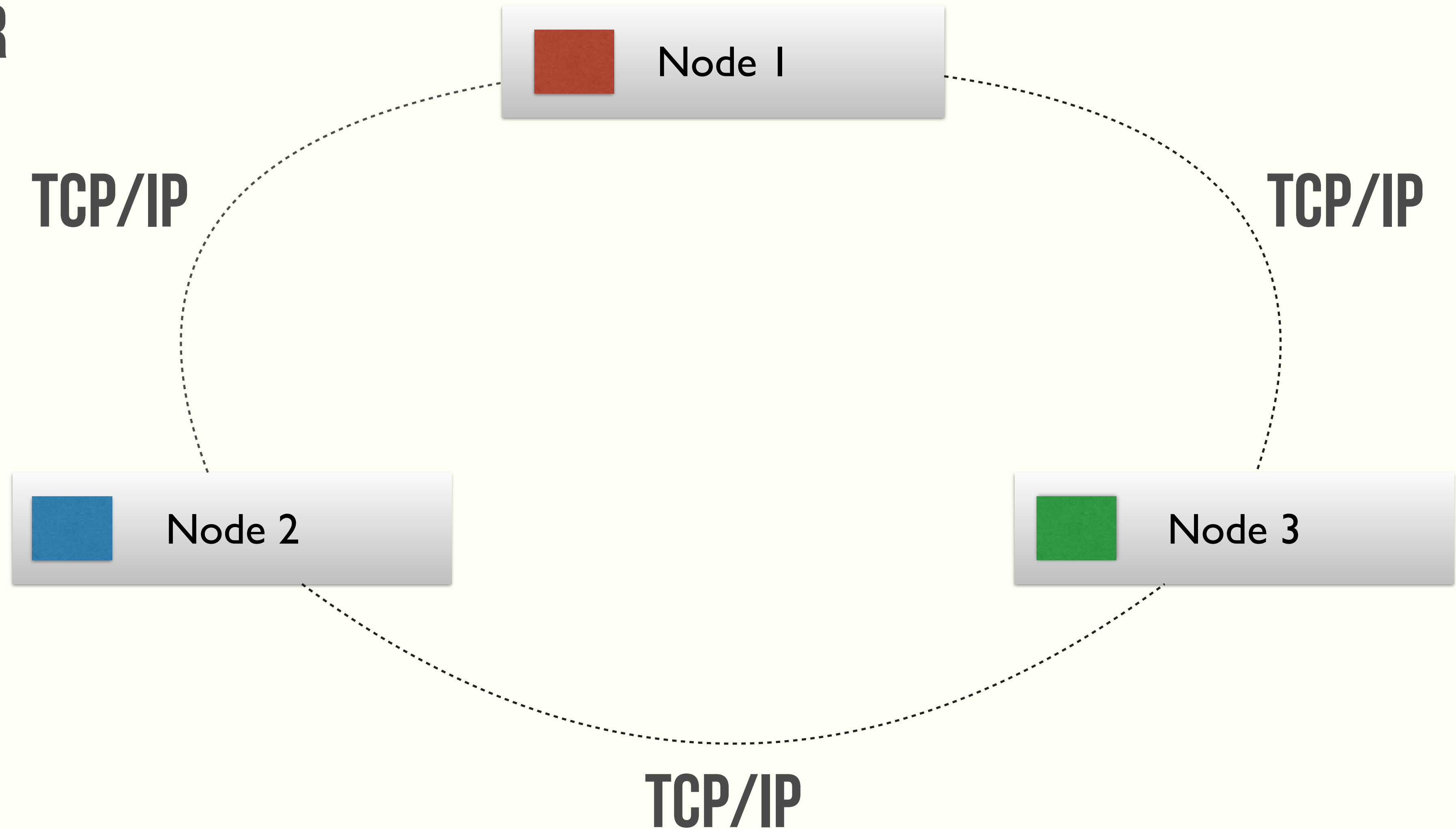
NOSQL

HAZELCAST

- **LEADING OPEN-SOURCE IN-MEMORY DATA / COMPUTE GRID**
- **JCACHE PROVIDER**
- **KEY-VALUE (IMAP, JCACHE)**
- **JUST A JAR!**

TOPOLOGY

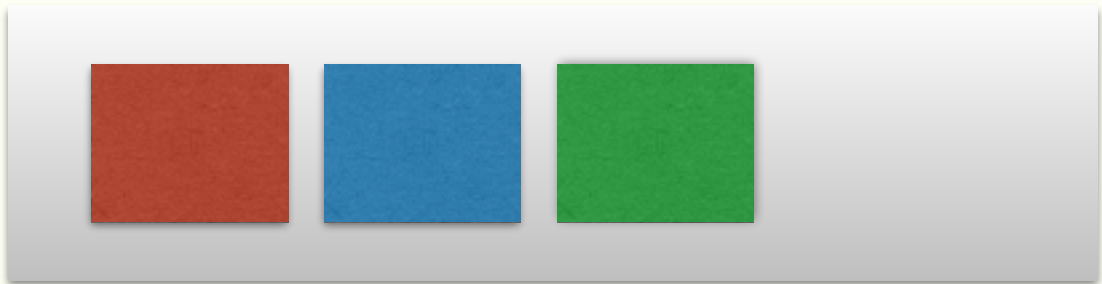
▸ PEER-TO-PEER



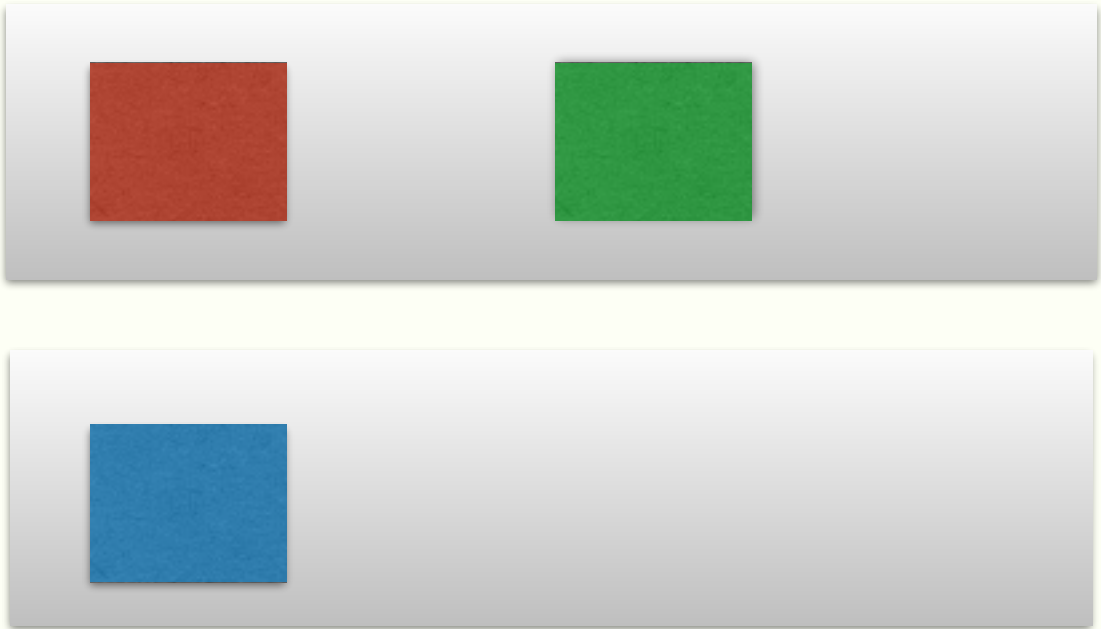
PARTITIONING / SHARDING



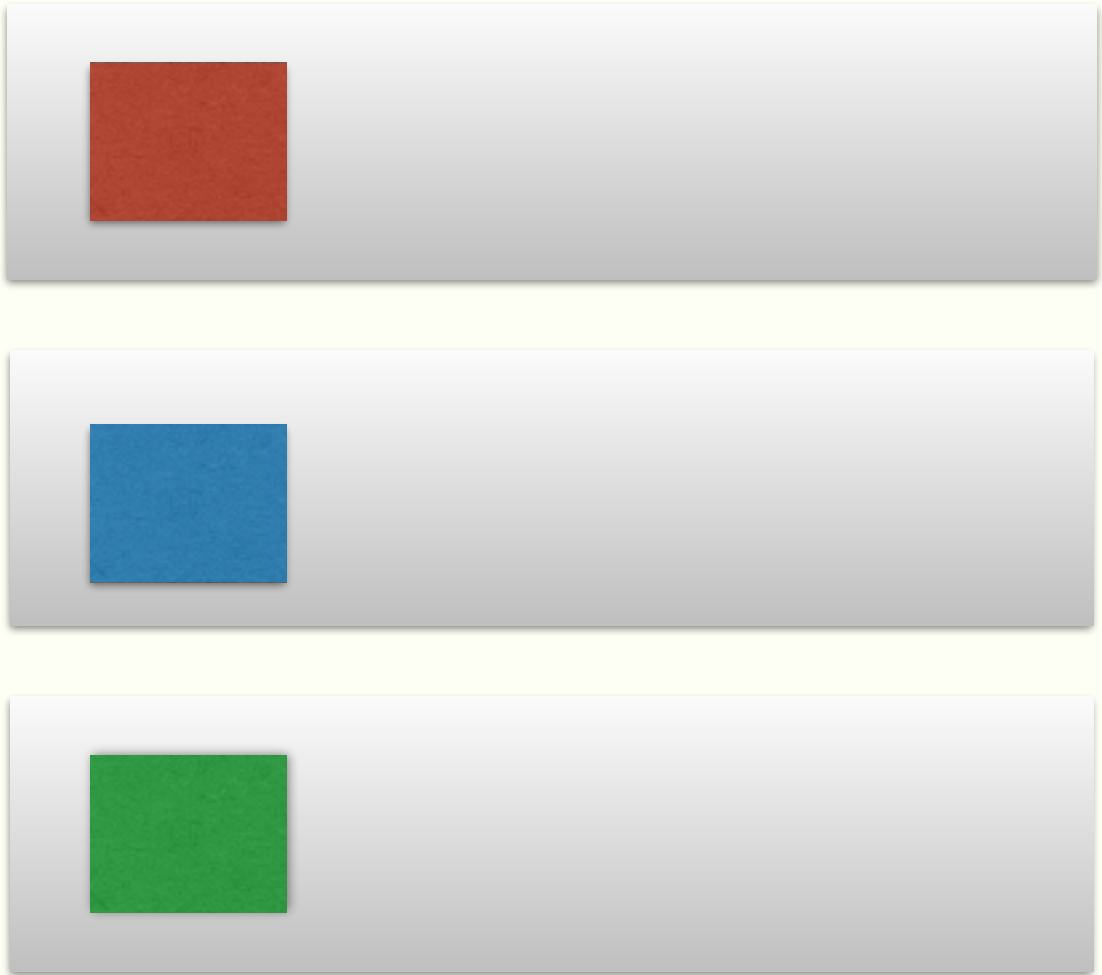
Step 1



Step 2



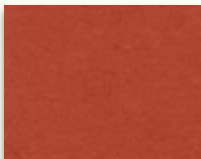
Step 3



Node 1

Node 2

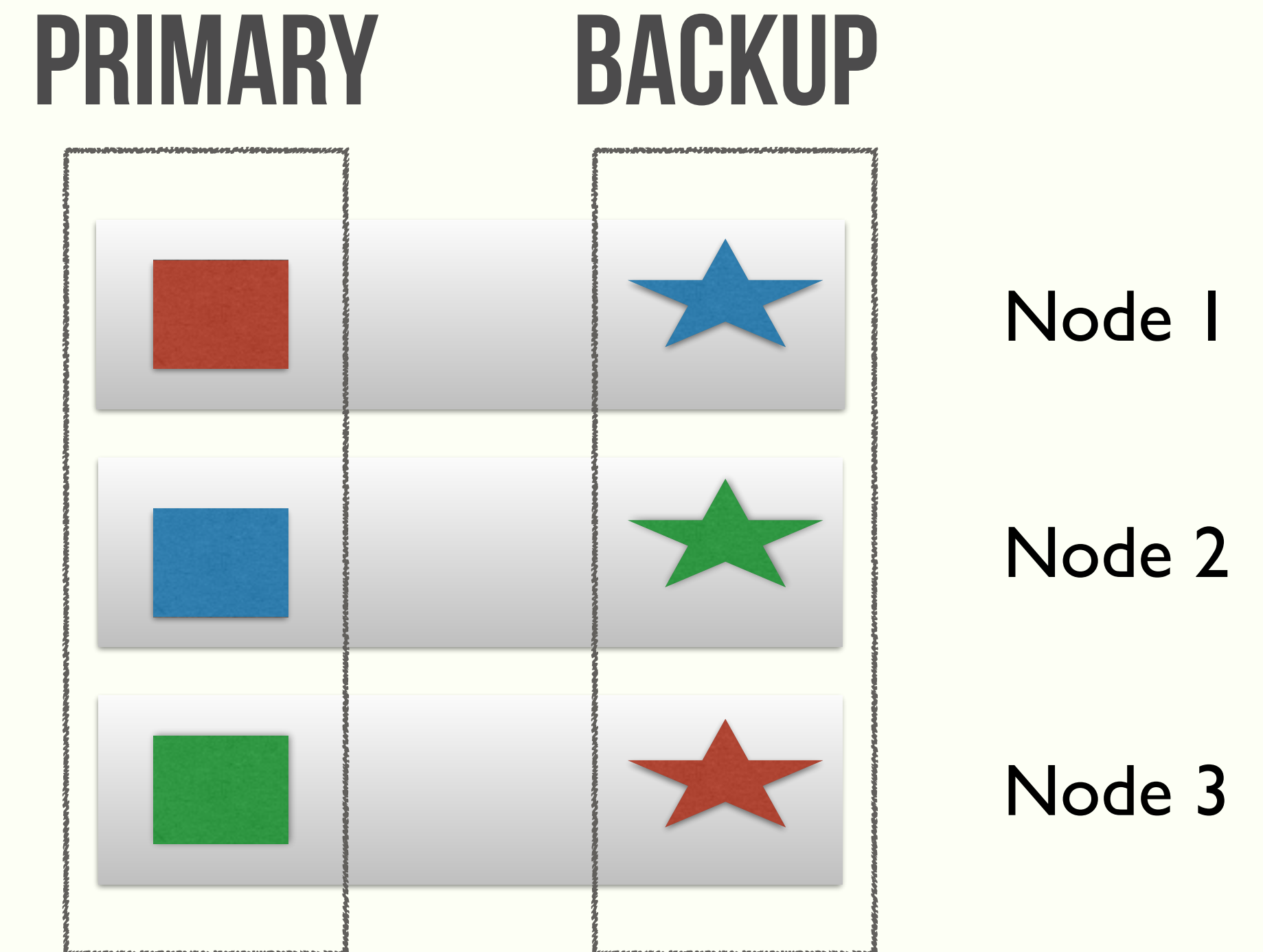
Node 3

   ... 271 PARTITIONS

$$\text{PARTITION_ID} = \text{HASH (KEY)} \% \text{PARTITION_COUNT}$$

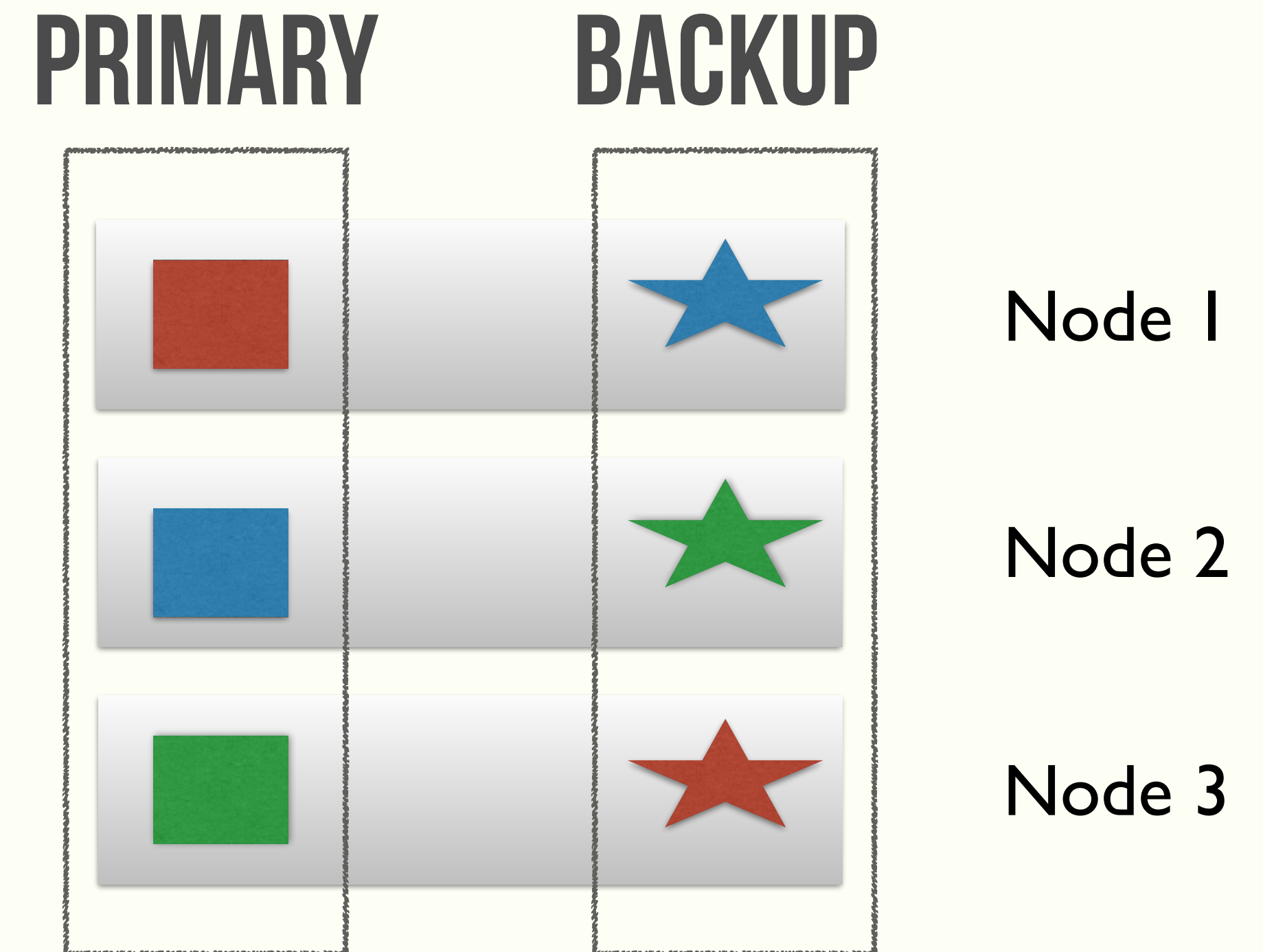
REPLICATION

- PRIMARY VS. BACKUP
 - SYNC
 - ASYNC
- BACKUP COUNT
 - PER DATA STRUCTURE



WRITING

- PUT -> SERIALISABLE ON PARTITION THREAD
- OPTIMISTIC "LOCKING"
 - REPLACE OPERATION
- PESSIMISTIC LOCKING
 - LOCK AND UNLOCK



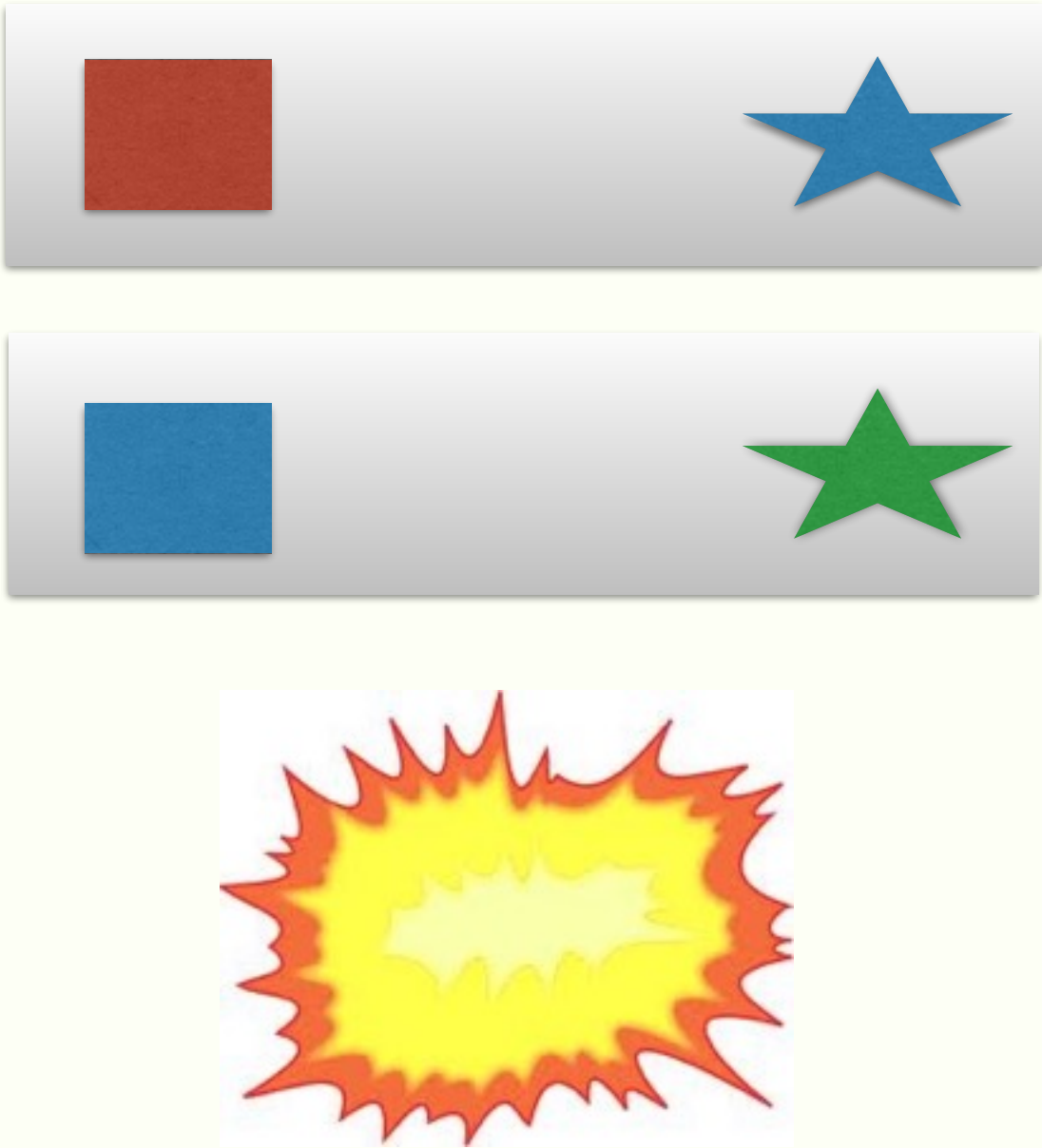
FAILOVER



Step 1



Step 2



Step 3

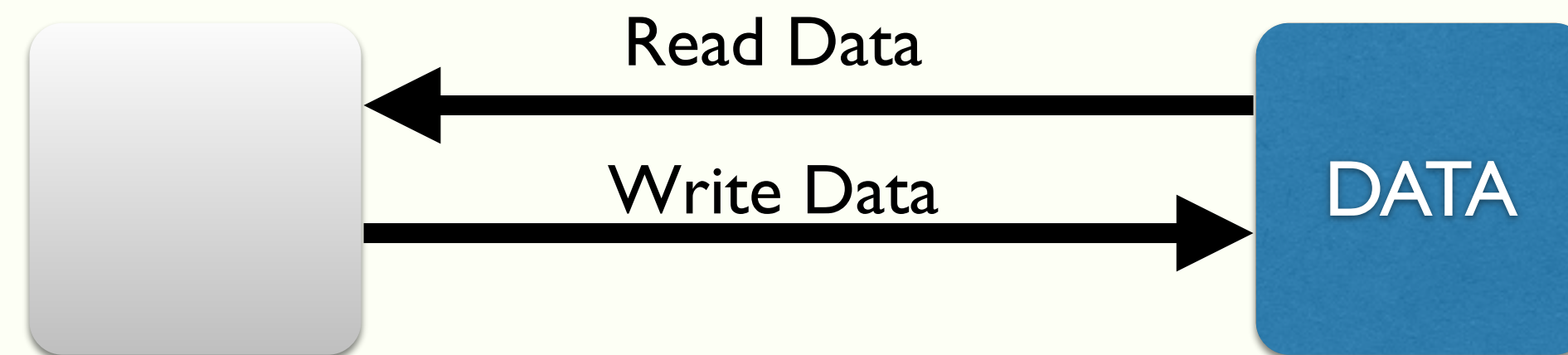


Node 1

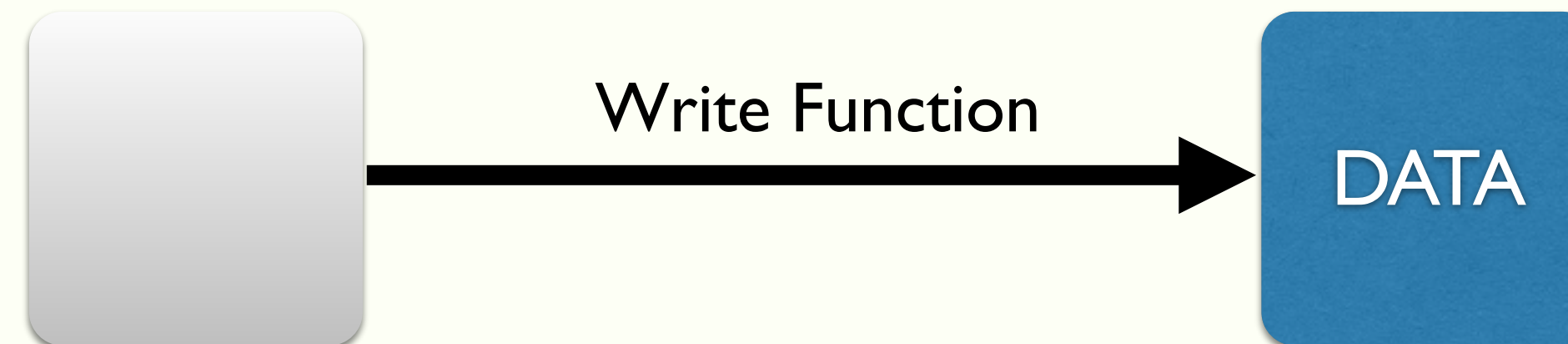
Node 2

Node 3

DATA PROCESSING

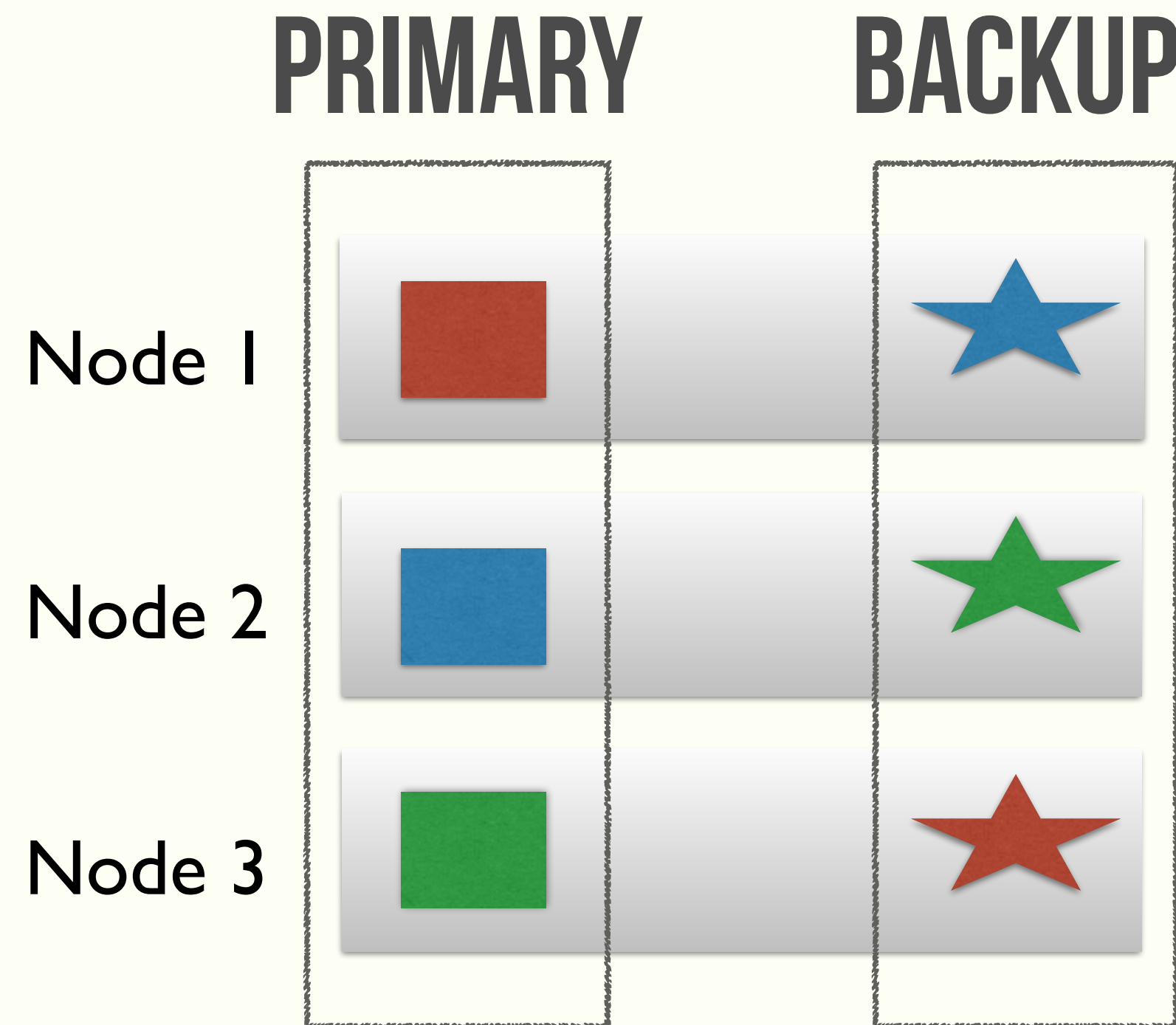


BAD: SEND DATA TO FUNCTION



GOOD: SEND FUNCTION TO DATA

ENTRY PROCESSOR



```
class SalaryProcessor
    extends AbstractEntryProcessor<> {

    int amount;

    SalaryProcessor(int amount) {
        this.amount = amount;
    }

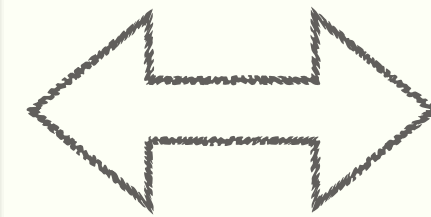
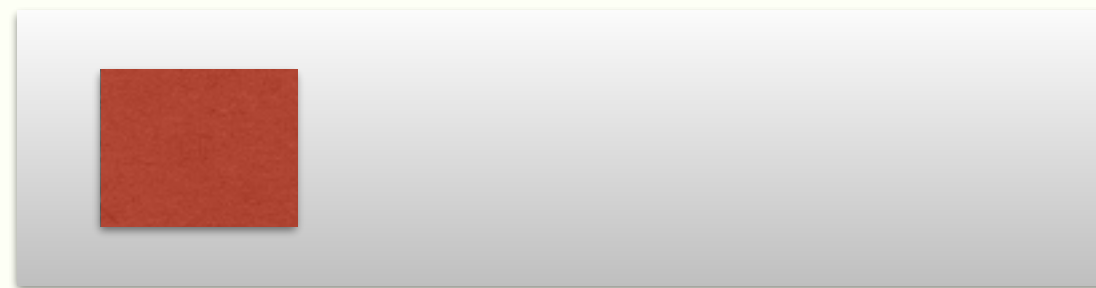
    @Override
    public Object process(Map.Entry<> entry) {
        entry.getValue().salary += amount;
        return null;
    }
}
```

DATA AFFINITY



MAP

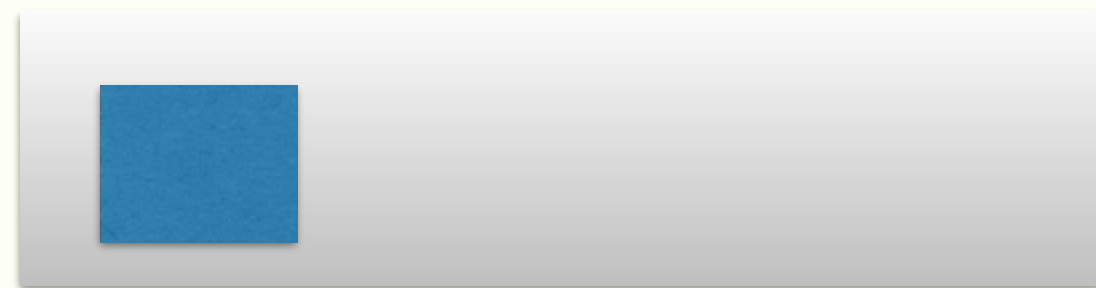
Node 1



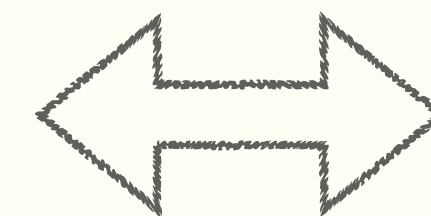
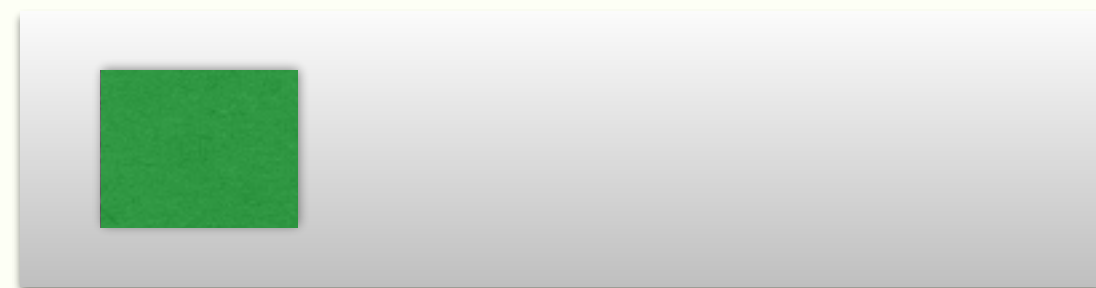
James Bond

Personal
Salary
Notes
Address

Node 2



Node 3



Eathen Hunt

Personal
Salary
Notes
Address

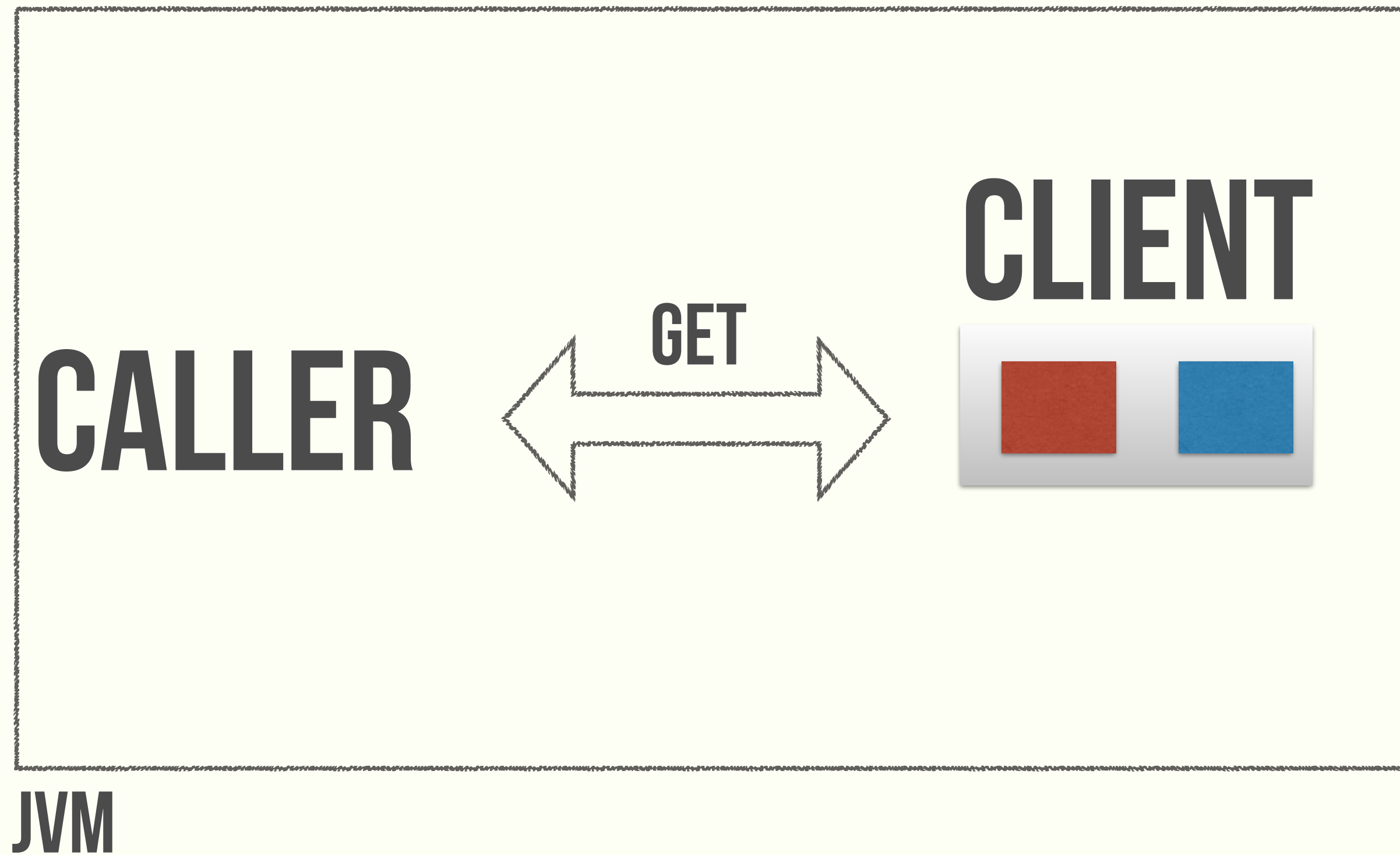
```
public class SalaryKey implements  
    Serializable, PartitionAware {
```

```
    private final long employeeId;  
    private final long salaryId;
```

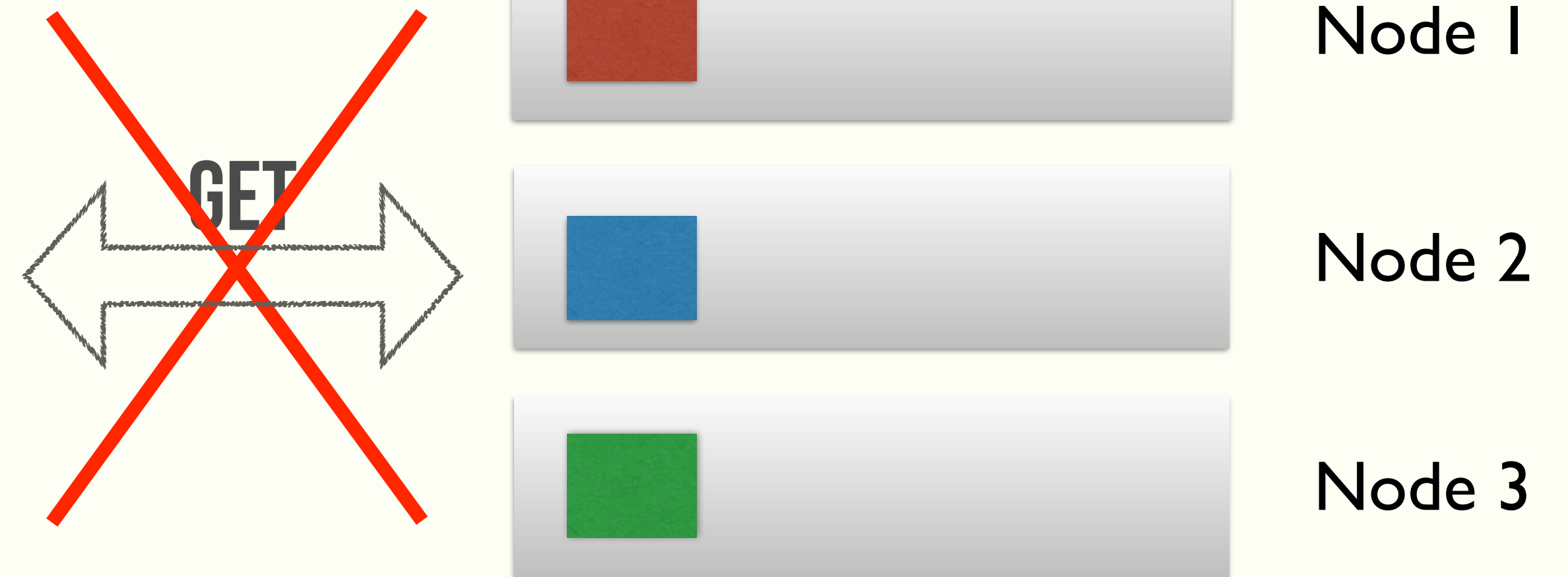
```
@Override
```

```
public Object getPartitionKey()  
{  
    return employeeId;  
}
```

WITH NEAR CACHE



hazelcast MAP



QUERYING - PREDICATES



```
public Set<Employee> getWithNameOrAge( int age, double salary ) {  
  
    Predicate age = Predicates.between( "age", 25, 35 );  
    Predicate salary = Predicates.lessThan( "salary", 50000 );  
    Predicate query = Predicates.or( age, salary );  
  
    // SqlPredicate query = new SqlPredicate(  
    //     "(age BETWEEN 25 AND 35) OR (salary < 50000)"  
    // )  
  
    return employees.values( query );  
  
}
```

HAZELCAST USED BY OPEN-SOURCE

- **VERT.X**
- **APACHE CAMEL, SHIRO, KARAF**
- **MULE ESB**
- **WSO2**
- **ALFRESCO**



mongoDB

OVERVIEW

- DOCUMENT ORIENTED
- JSON-STYLE DOCUMENTS
 - MAX 16MB
- DYNAMIC (NO) SCHEMAS
- AD-HOC QUERY LANGUAGE

Collection
↓
`db.users.insert(`

Document
↓

```
{  
  name: "sue",  
  age: 26,  
  status: "A",  
  groups: [ "news", "sports" ]  
}
```

)

Document

```
{  
  name: "sue",  
  age: 26,  
  status: "A",  
  groups: [ "news", "sports" ]  
}
```

insert

Collection

{ name: "al", age: 18, ... }
{ name: "lee", age: 28, ... }
{ name: "jan", age: 21, ... }
{ name: "kai", age: 38, ... }
{ name: "sam", age: 18, ... }
{ name: "mel", age: 38, ... }
{ name: "ryan", age: 31, ... }
{ name: "sue", age: 26, ... }

users

DOCUMENTS

- BSON DOCUMENT
- DOCUMENT

```
{  
  "a" : "MongoDB",  
  "b" : [ 1, 2 ]  
}
```



mongoDB

```
new BsonDocument()  
  .append("a", new BsonString("MongoDB"))  
  .append("b", new BsonArray(  
    Arrays.asList(new BsonInt32(1), new BsonInt32(2))  
  ));
```

```
new Document().append("a", "MongoDB")  
  .append("b", Arrays.asList(1, 2));
```

WRITING

- ATOMICITY ON SINGLE DOC
- NO MULTI-DOC TX
- \$ISOLATED



mongoDB

Collection
↓
`db.users.insert(`

Document
↓

```
{  
  name: "sue",  
  age: 26,  
  status: "A",  
  groups: [ "news", "sports" ]  
}
```

)

Document

```
{  
  name: "sue",  
  age: 26,  
  status: "A",  
  groups: [ "news", "sports" ]  
}
```

insert

Collection

{ name: "al", age: 18, ... }
{ name: "lee", age: 28, ... }
{ name: "jan", age: 21, ... }
{ name: "kai", age: 38, ... }
{ name: "sam", age: 18, ... }
{ name: "mel", age: 38, ... }
{ name: "ryan", age: 31, ... }
{ name: "sue", age: 26, ... }

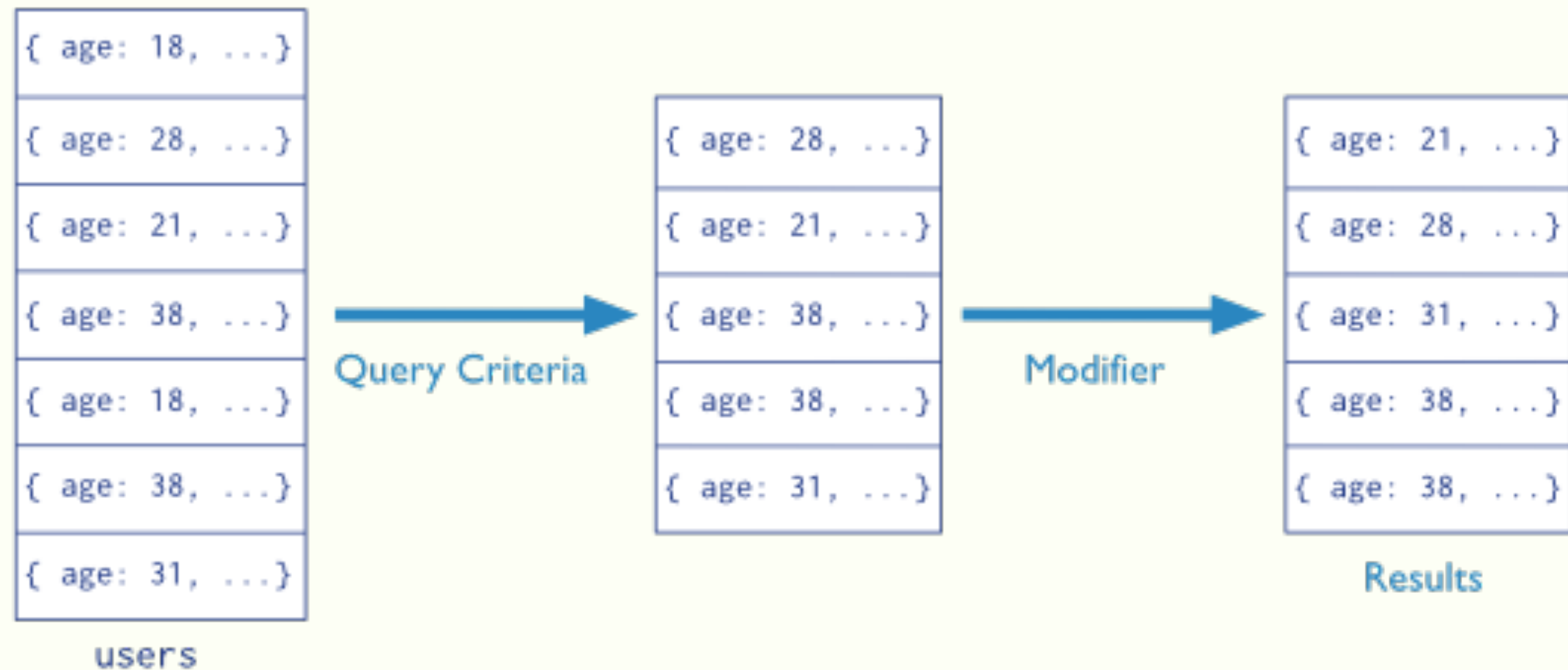
users



QUERYING

- NO SQL
- NO JOINS
 - \$LOOKUP
- SINGLE COLLECTION
- PROJECTIONS
- AGGREGATIONS

Collection Query Criteria Modifier
`db.users.find( { age: { $gt: 18 } } ).sort( {age: 1 } )`



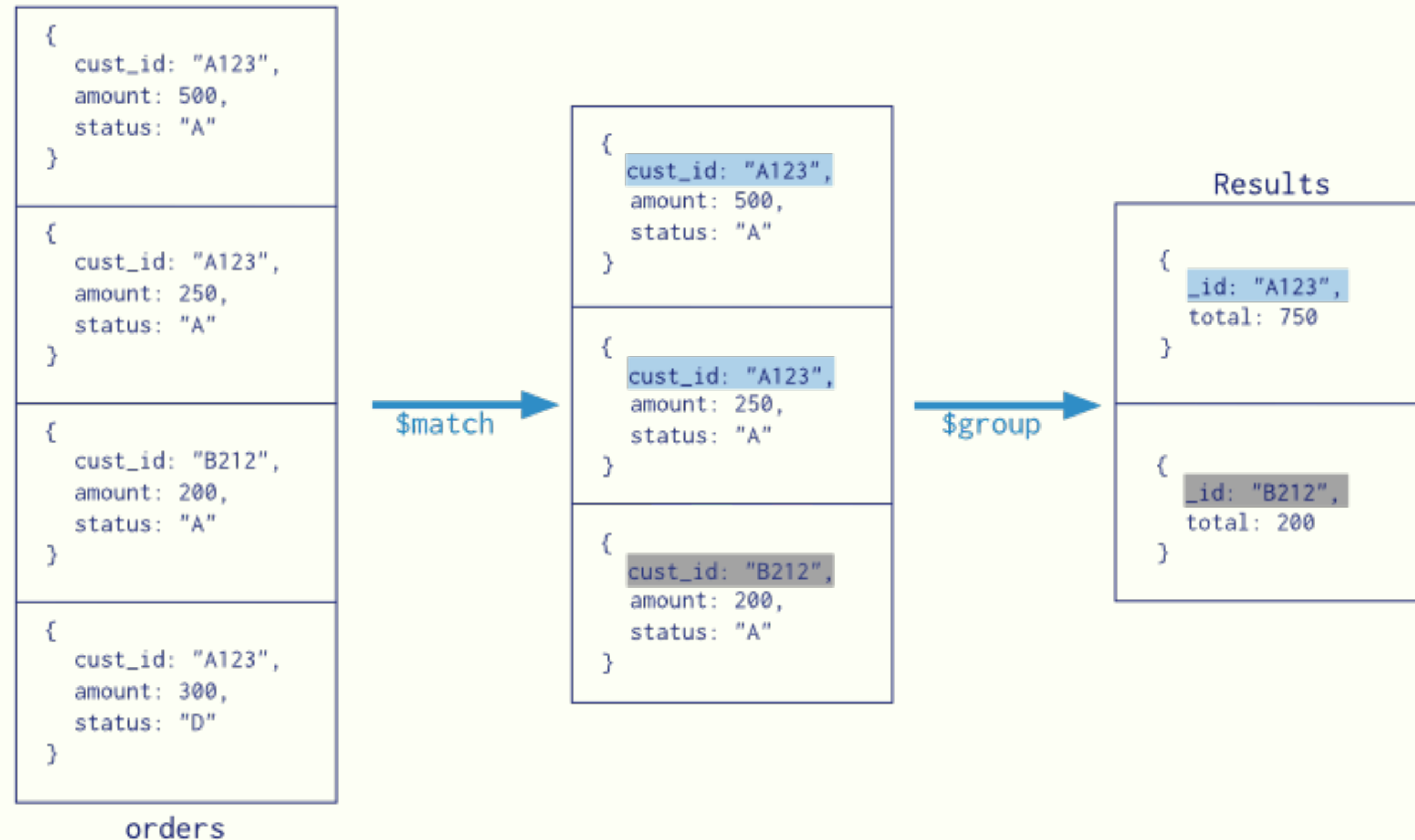
AGGREGATIONS

- ▶ MULTIPLE STAGES
- ▶ \$LOOKUP



mongoDB

Collection
↓
`db.orders.aggregate( [`
 `$match stage → { $match: { status: "A" } },`
 `$group stage → { $group: { _id: "$cust_id", total: { $sum: "$amount" } } }`
 `]`)

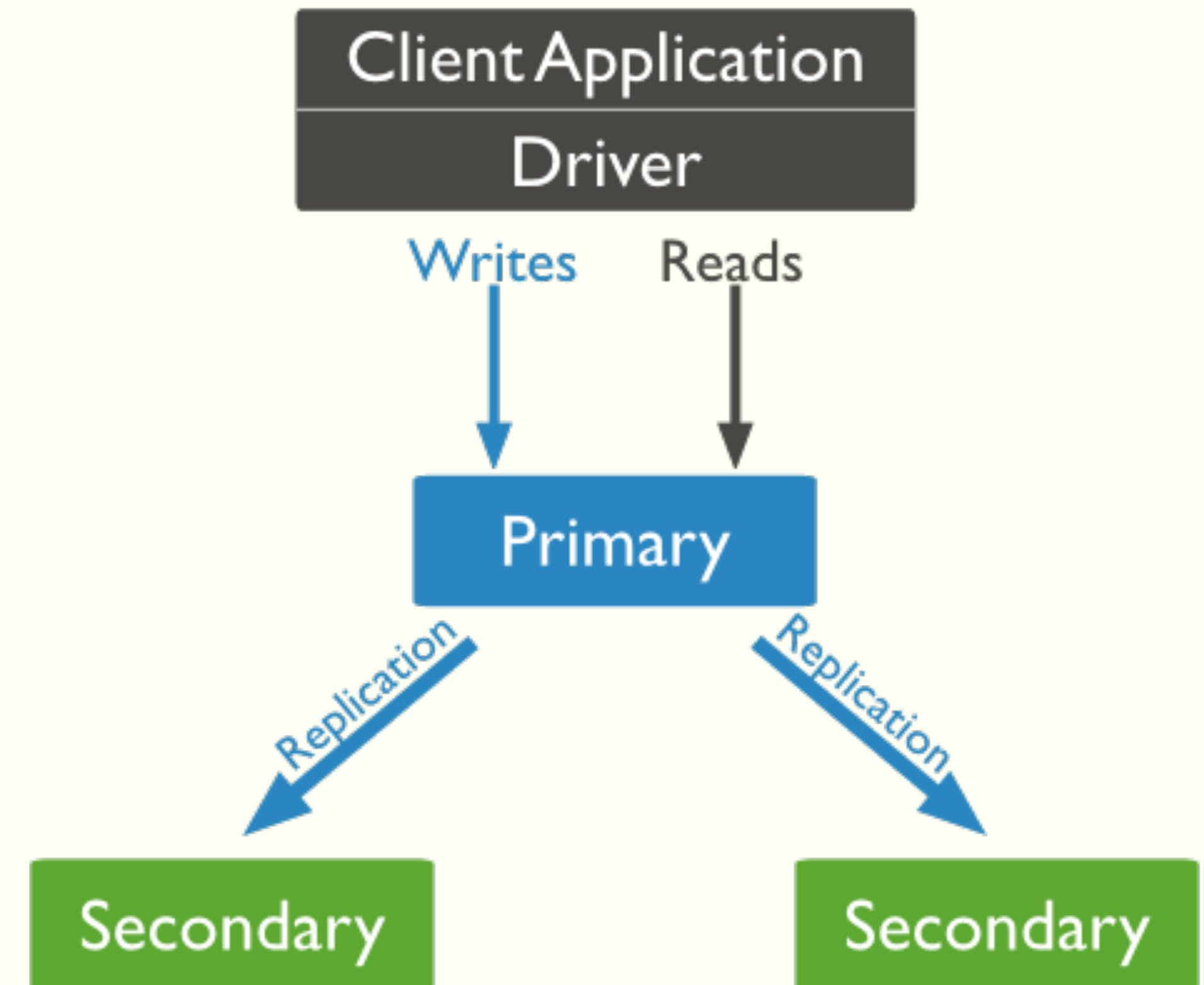


REPLICATION

- ONLY ONE PRIMARY
- ASYNC REPLICATION
- WRITE TO PRIMARY ONLY
- READ FROM PRIMARY -> STRICT CONSIST.
- READ FROM SECONDARY -> EVENTUAL CONSIST.
- AUTOMATIC FAILOVER -> REELECTION



mongoDB

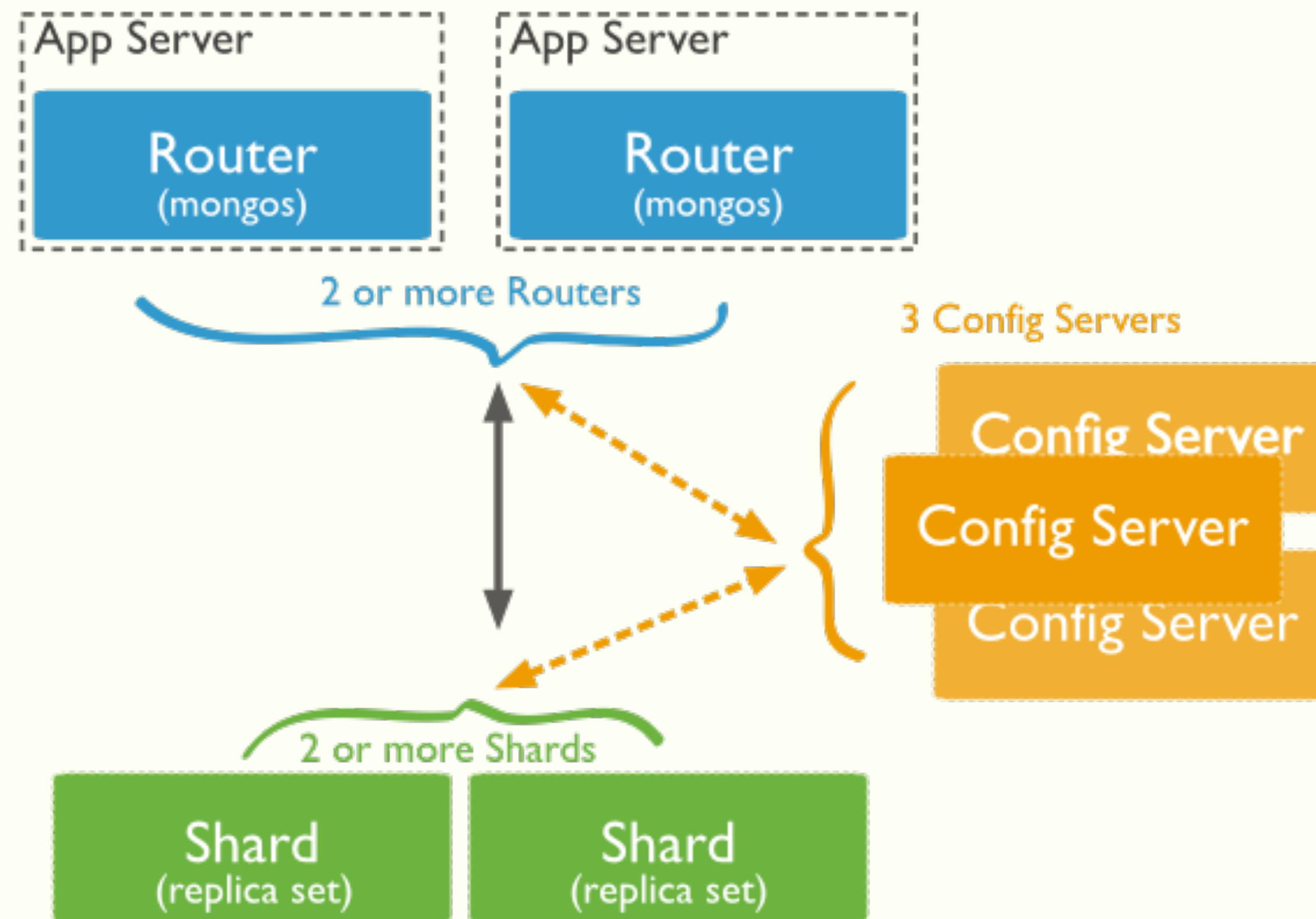


SHARDING

- SHARDS
- CONFIG SERVERS
- ROUTERS (MONGOS)
- SHARDS



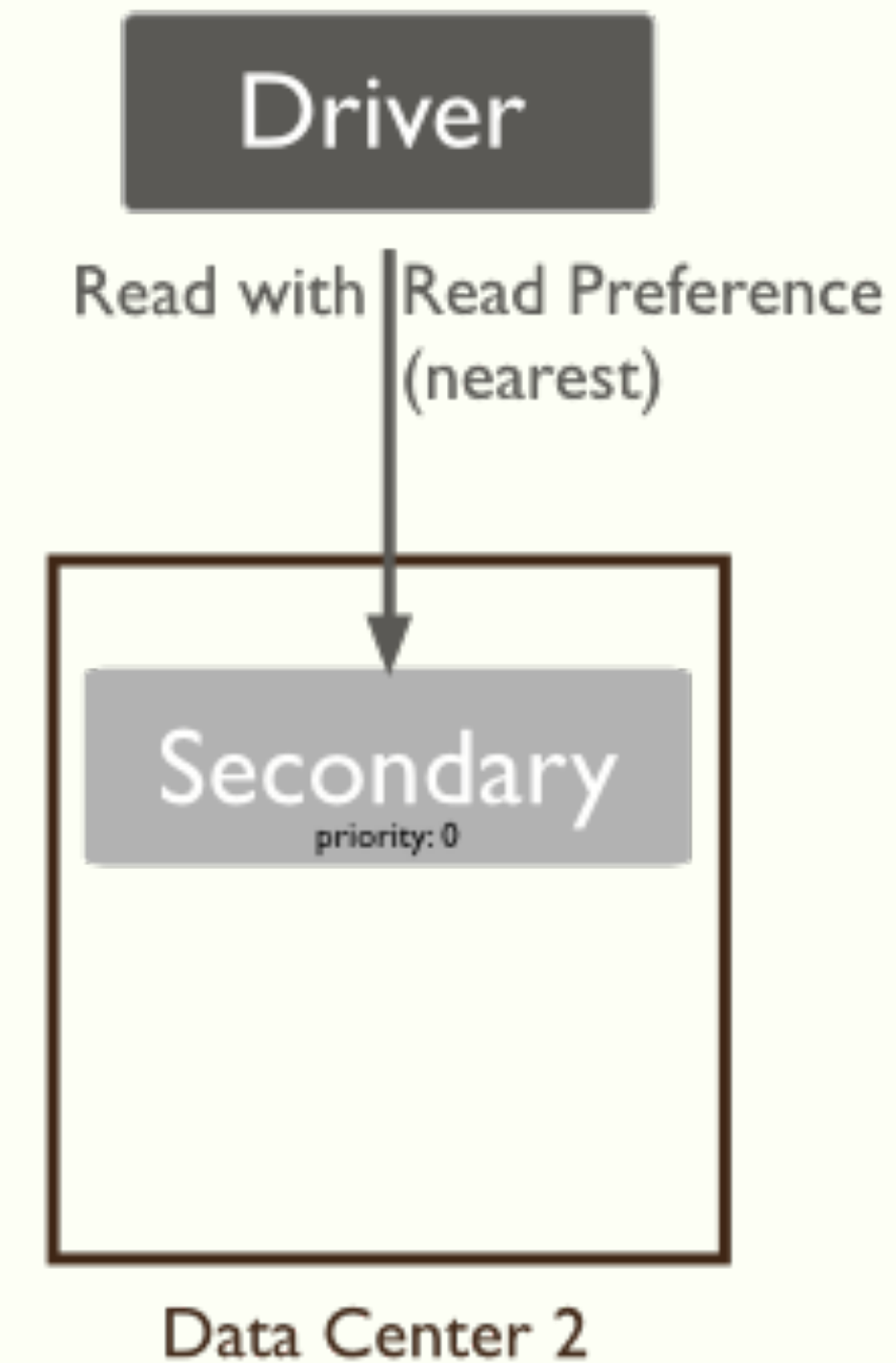
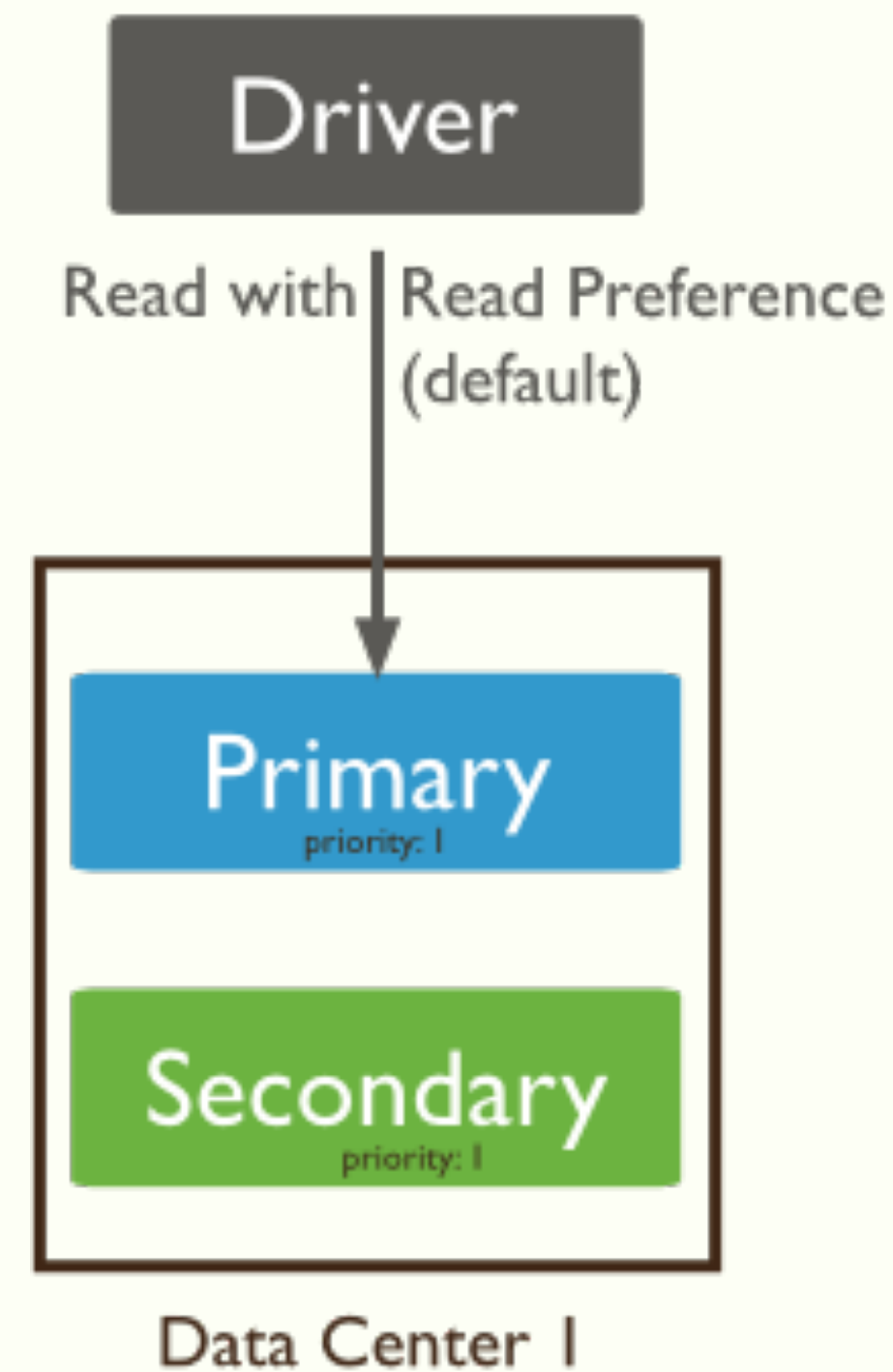
mongoDB





READ PREFERENCE

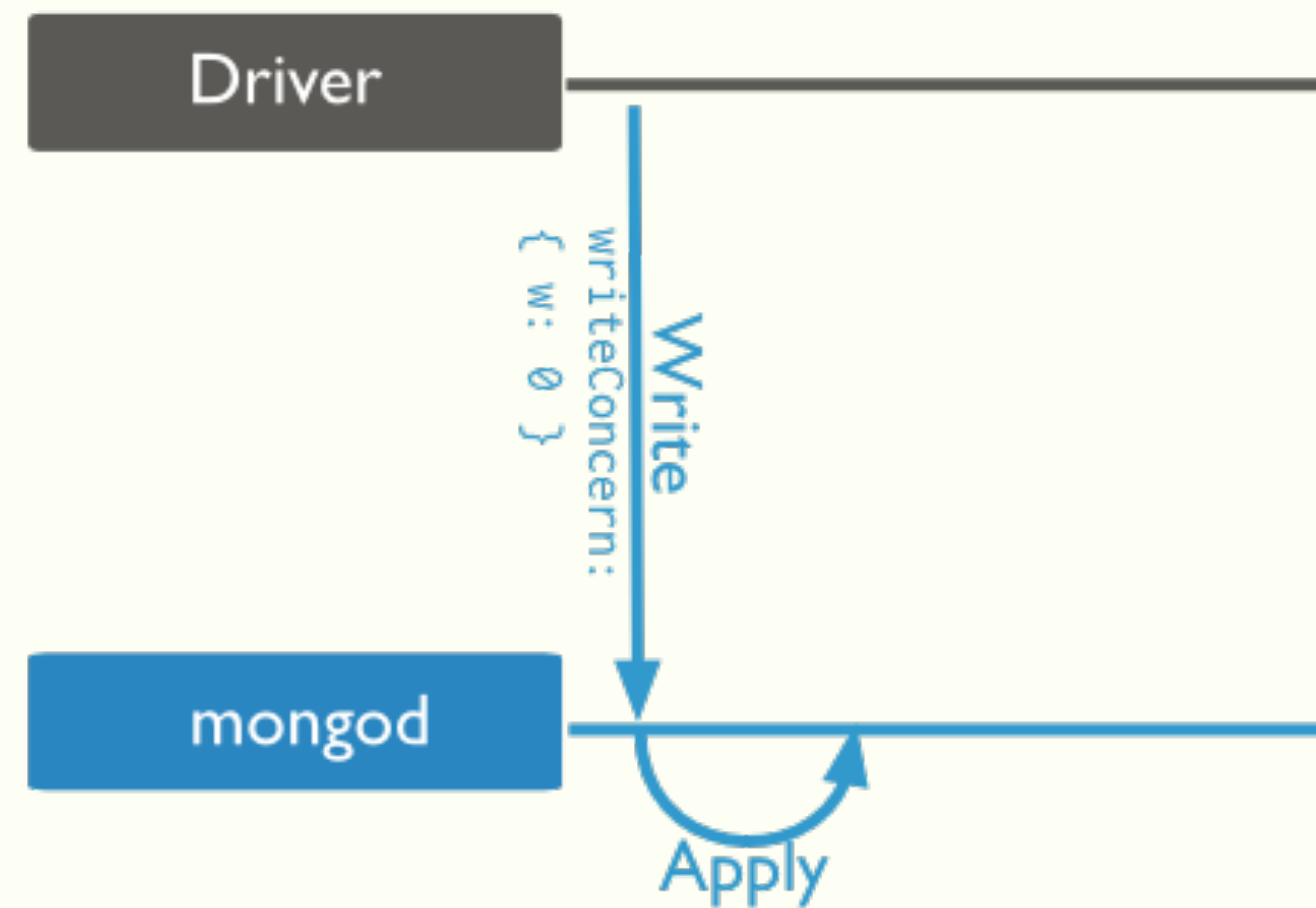
- ▶ PRIMARY (DEFAULT)
- ▶ PRIMARY_PREFERRED
- ▶ SECONDARY
- ▶ SECONDARY_PREF.
- ▶ NEAREST



WRITE CONCERN

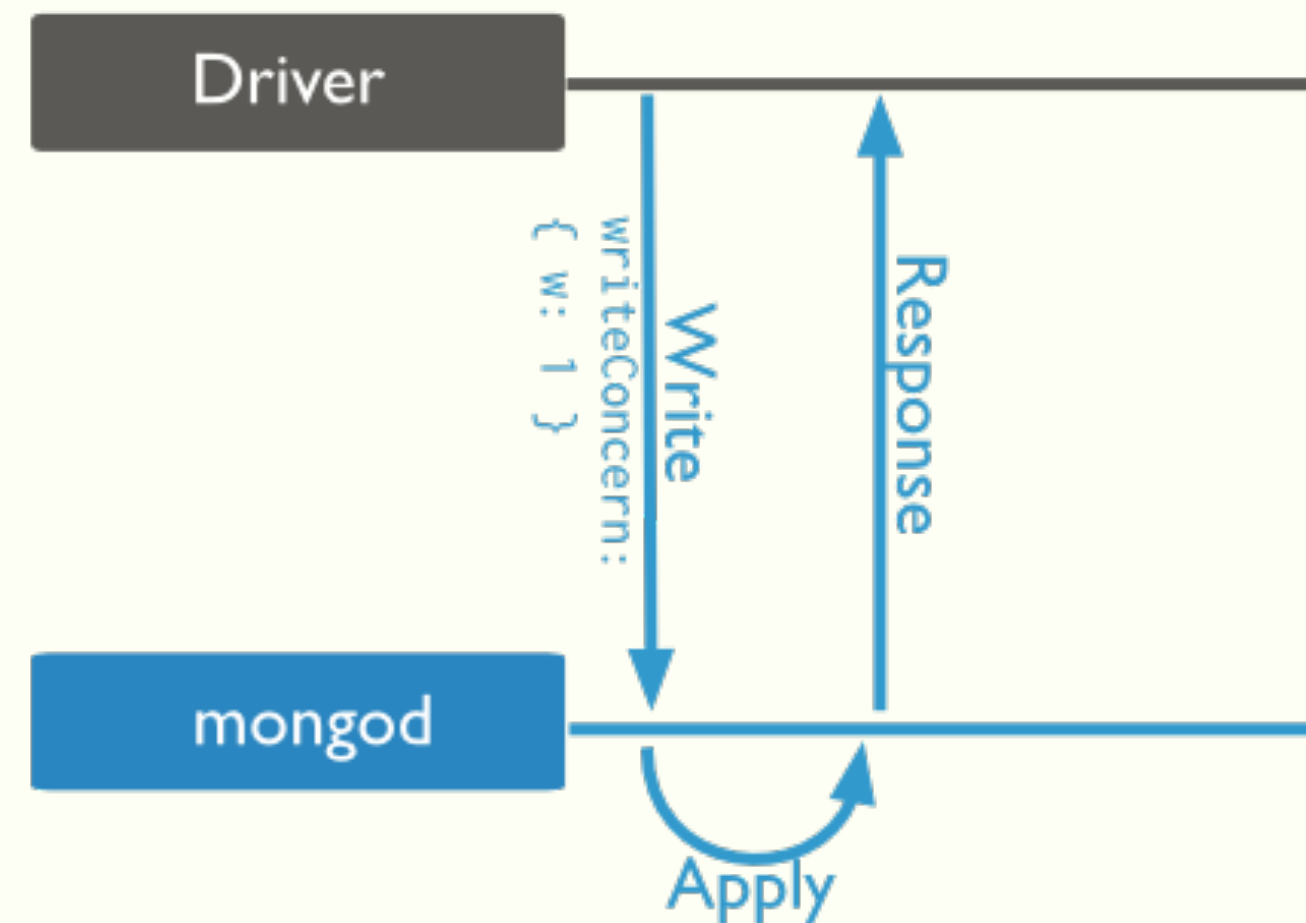


mongoDB



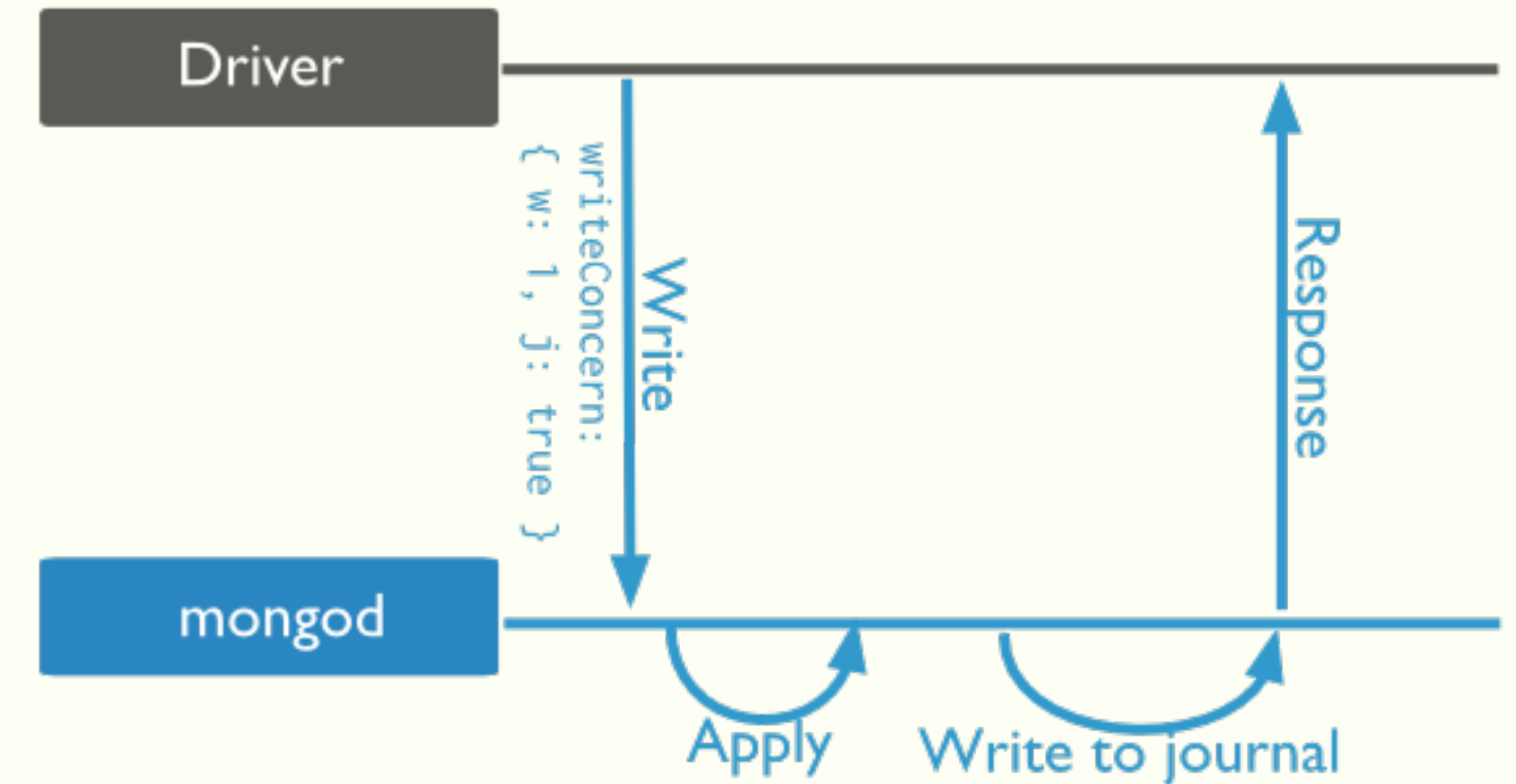
▸ UNACKNOWLEDGED

$W == 0$



▸ ACKNOWLEDGED

$W == 1$



▸ JOURNALED

$W == 1 ; J = \text{TRUE}$



USE CASES

- **WEB-BASED APPS**
- **EBAY**
 - **METADATA STORAGE**
 - **CLOUD MANAGEMENT**
 - **MERCHANDISING CATEGORIZATION**
- **THE WEATHER CHANNEL**
 - **NEAR REAL TIME WEATHER ALERTS**

APACHE HBASE



- ▶ BIG TABLE STORE
- ▶ COLUMN-ORIENTED
- ▶ STRONGLY CONSISTENT RW
- ▶ AUTOMATIC SHARDING

Row Key	Customer		Sales	
	Customer Id			
	Name	City	Product	Amount
101	John White	Los Angeles, CA	Chairs	\$400.00
102	Jane Brown	Atlanta, GA	Lamps	\$200.00
103	Bill Green	Pittsburgh, PA	Desk	\$500.00
104	Jack Black	St. Louis, MO	Bed	\$1600.00

Column Families

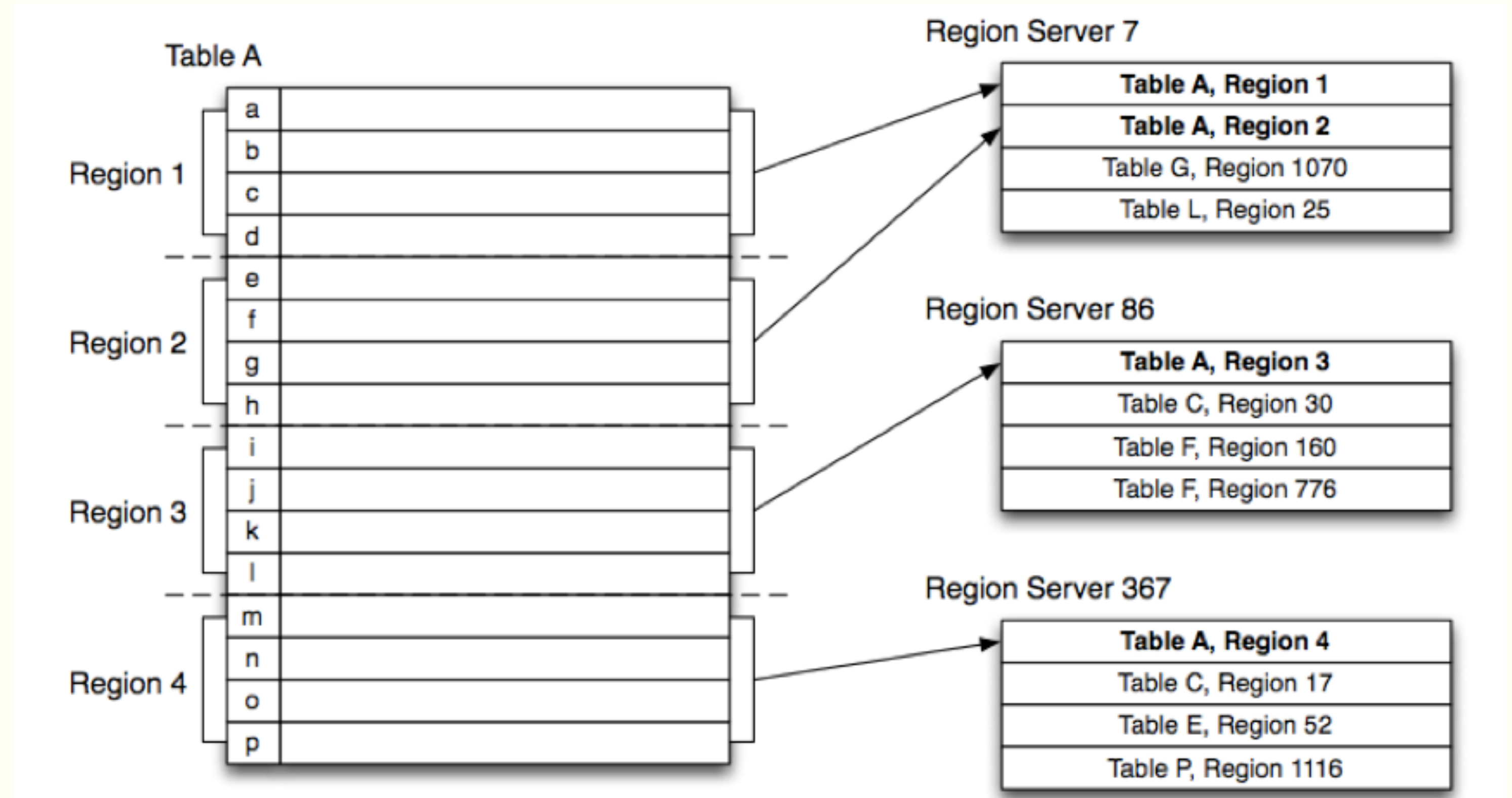
STORAGE

- ▶ SCHEMA-LESS
- ▶ COLUMN-ORIENTED
- ▶ WIDE AND SPARSE TABLES



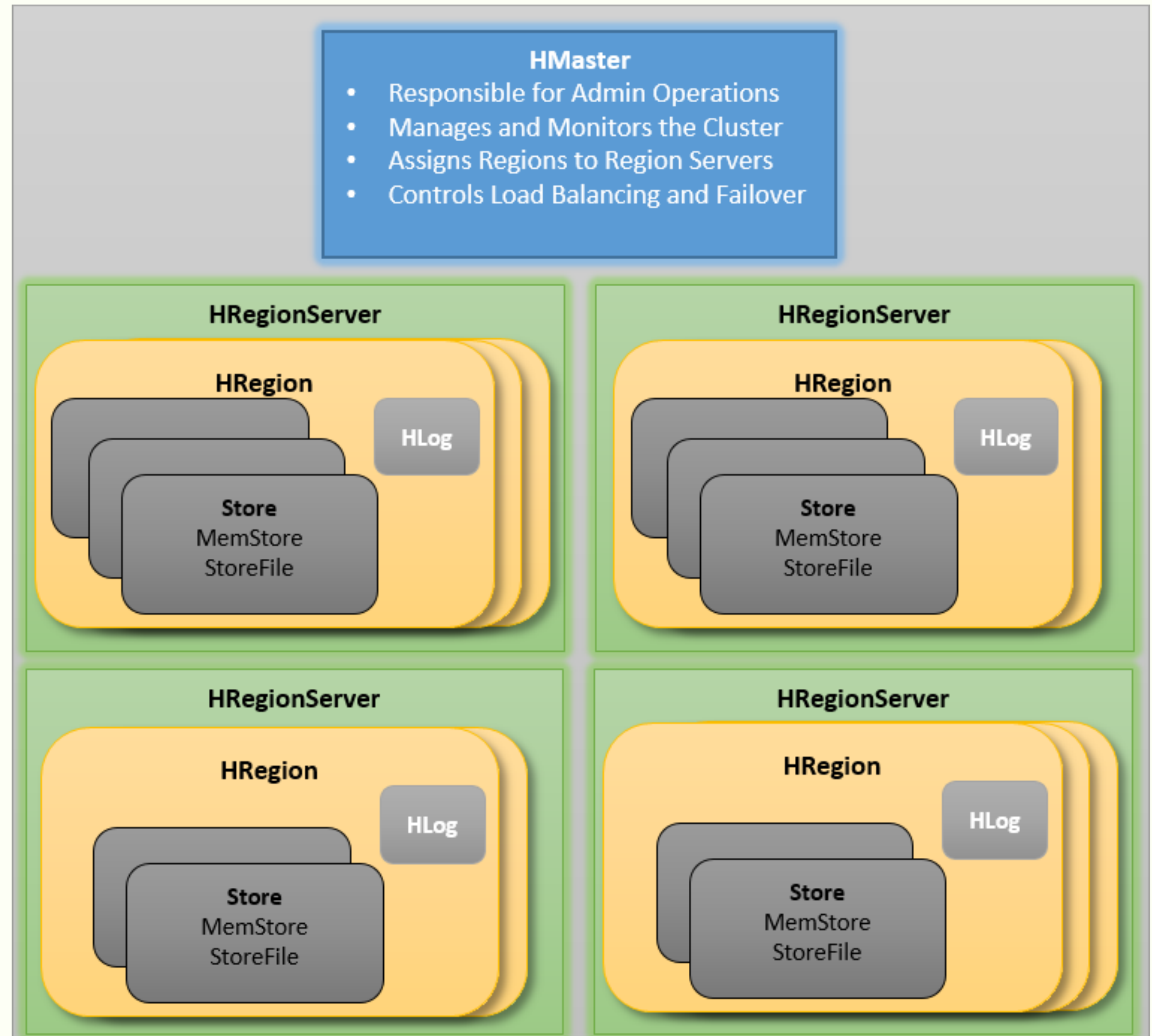
Row ID	Customer	Product	Amount
101	John White	Chairs	\$400.00
102	Jane Brown	Lamps	\$500.00
103	Bill Green	Lamps	\$150.00
104	Jack Black	Desk	\$700.00
105	Jane Brown	Desk	\$650.00
106	Bill Green	Desk	\$900.00

HBASE SHARDING

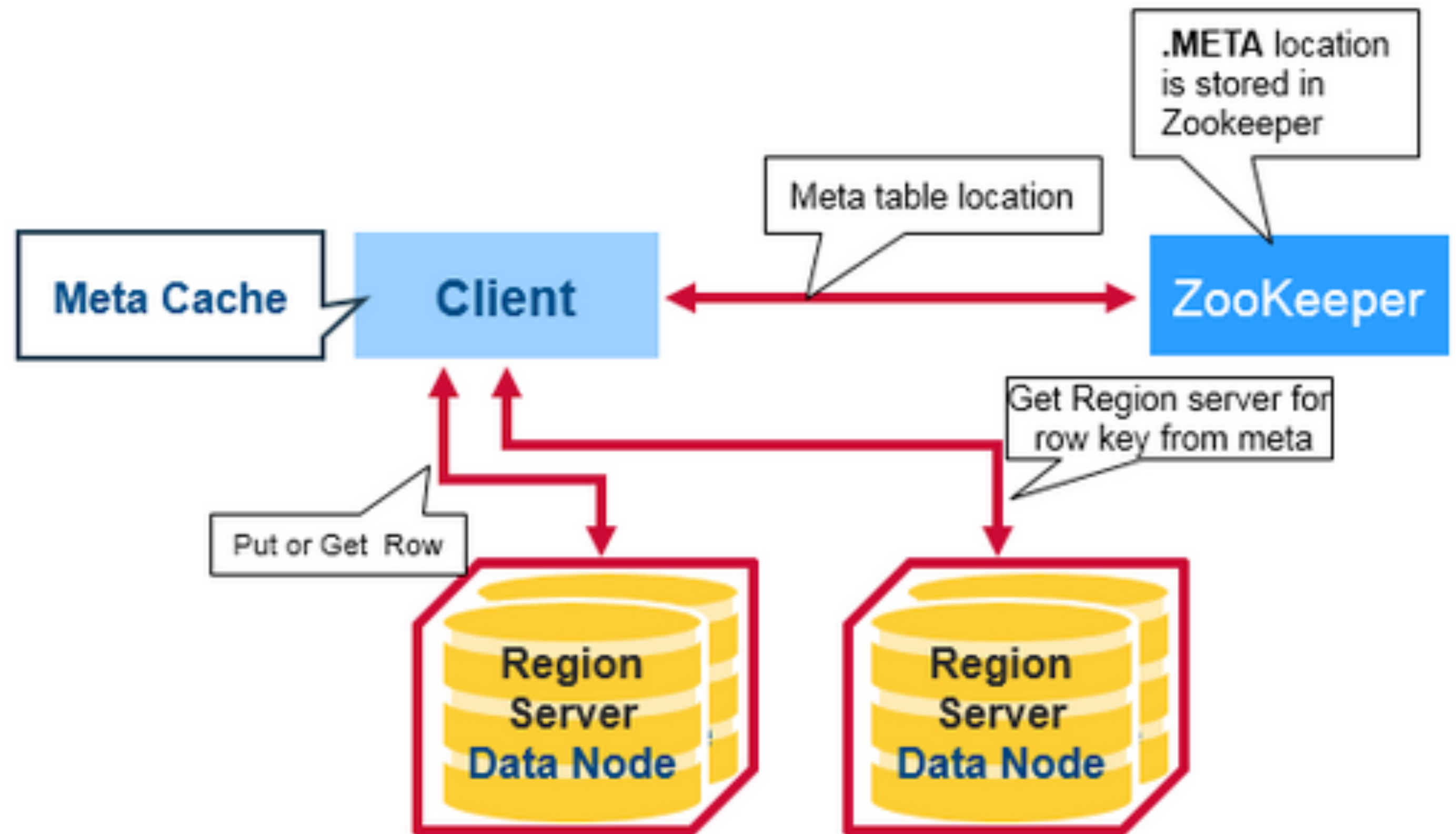


HBASE ARCHITECTURE

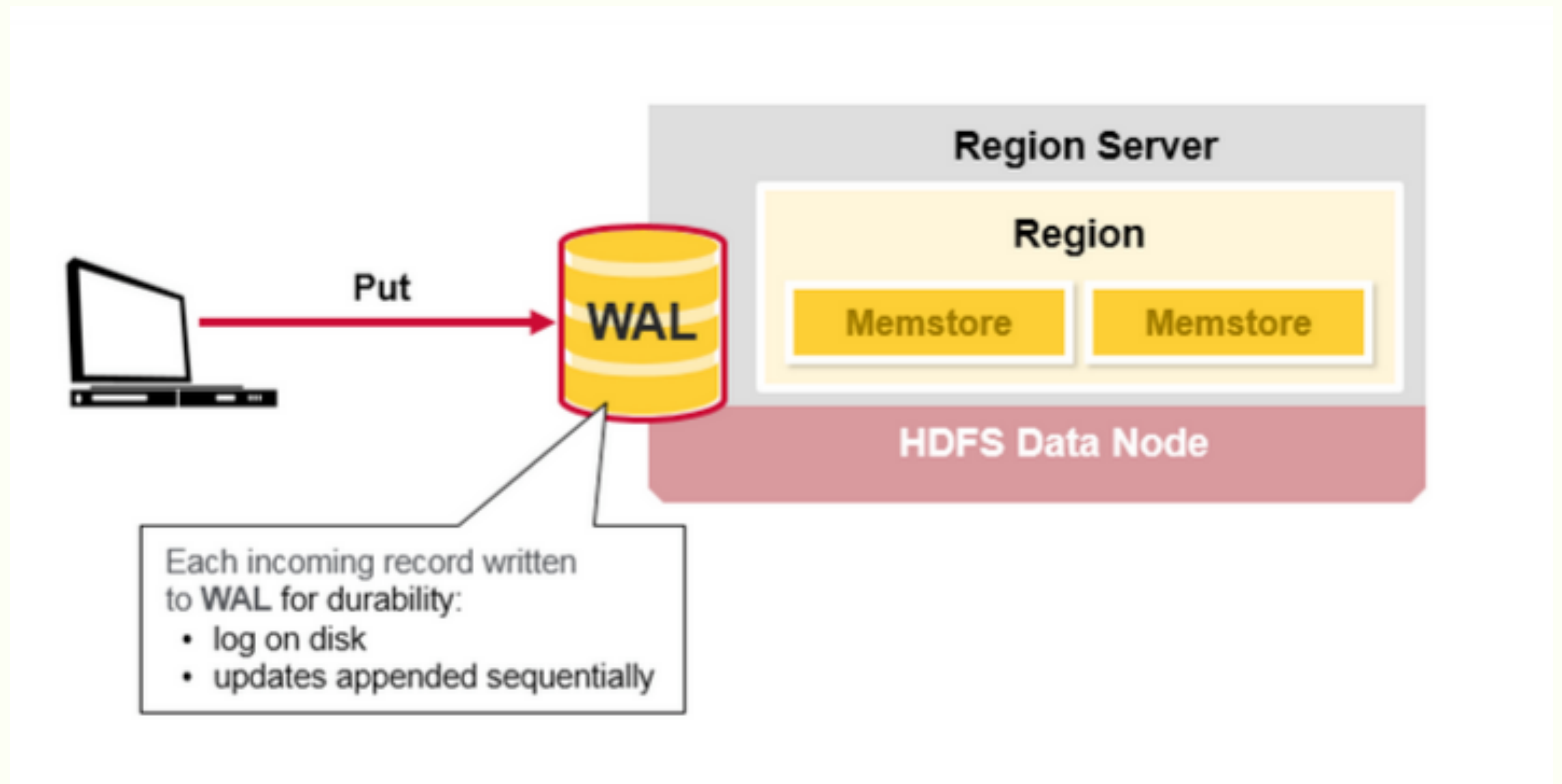
- MASTER
 - ASSIGNS THE REGIONS
 - LOAD-BALANCING
 - FAIL-OVER
- REGION SERVER
 - MANAGE/SPLIT REGIONS
 - READ/WRITE OPERATIONS



HBASE READ / WRITE

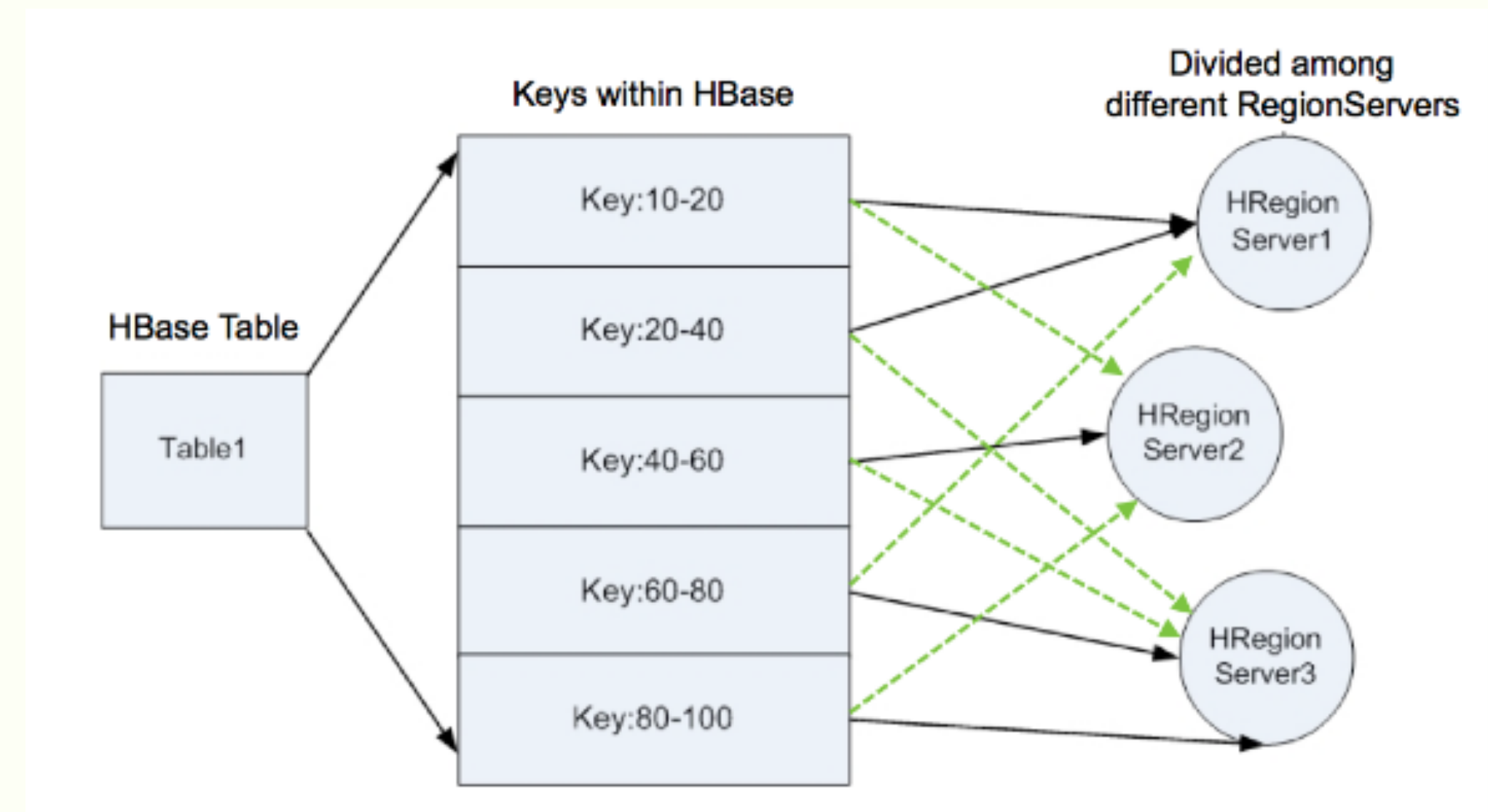


HBASE READ / WRITE



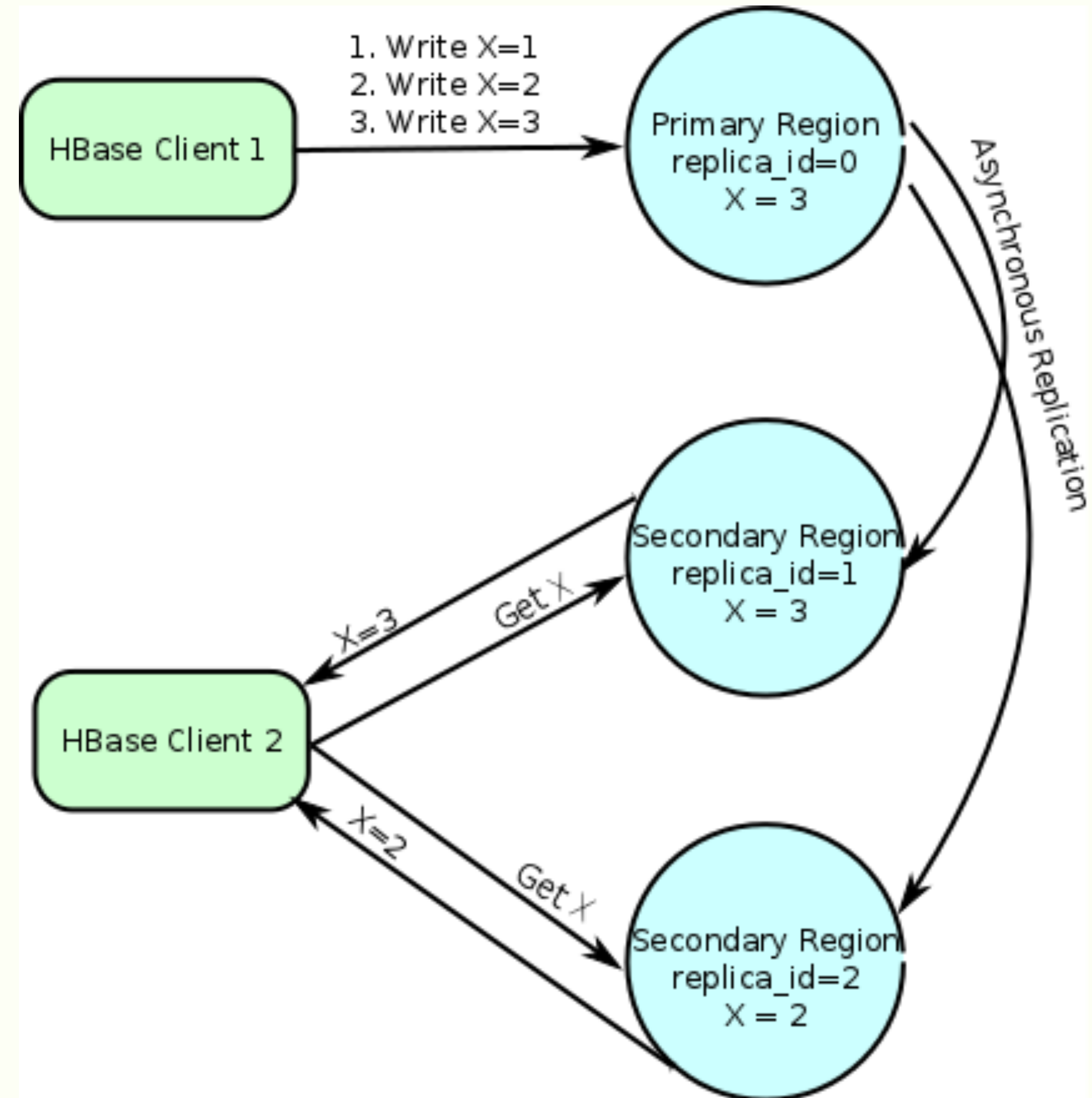
HBASE REPLICATION & SHARDING

- ASYNC REPLICATION
 - HDFS REPLICATION FACTOR
- PRIMARY REGION -> READ-WRITE
- SECONDARY REGION -> READ ONLY
- SECONDARY -> MAY BE STALE



HBASE CONSISTENCY

- ▶ **STRONG**
- ▶ **TIMELINE** (FROM 01.2015)

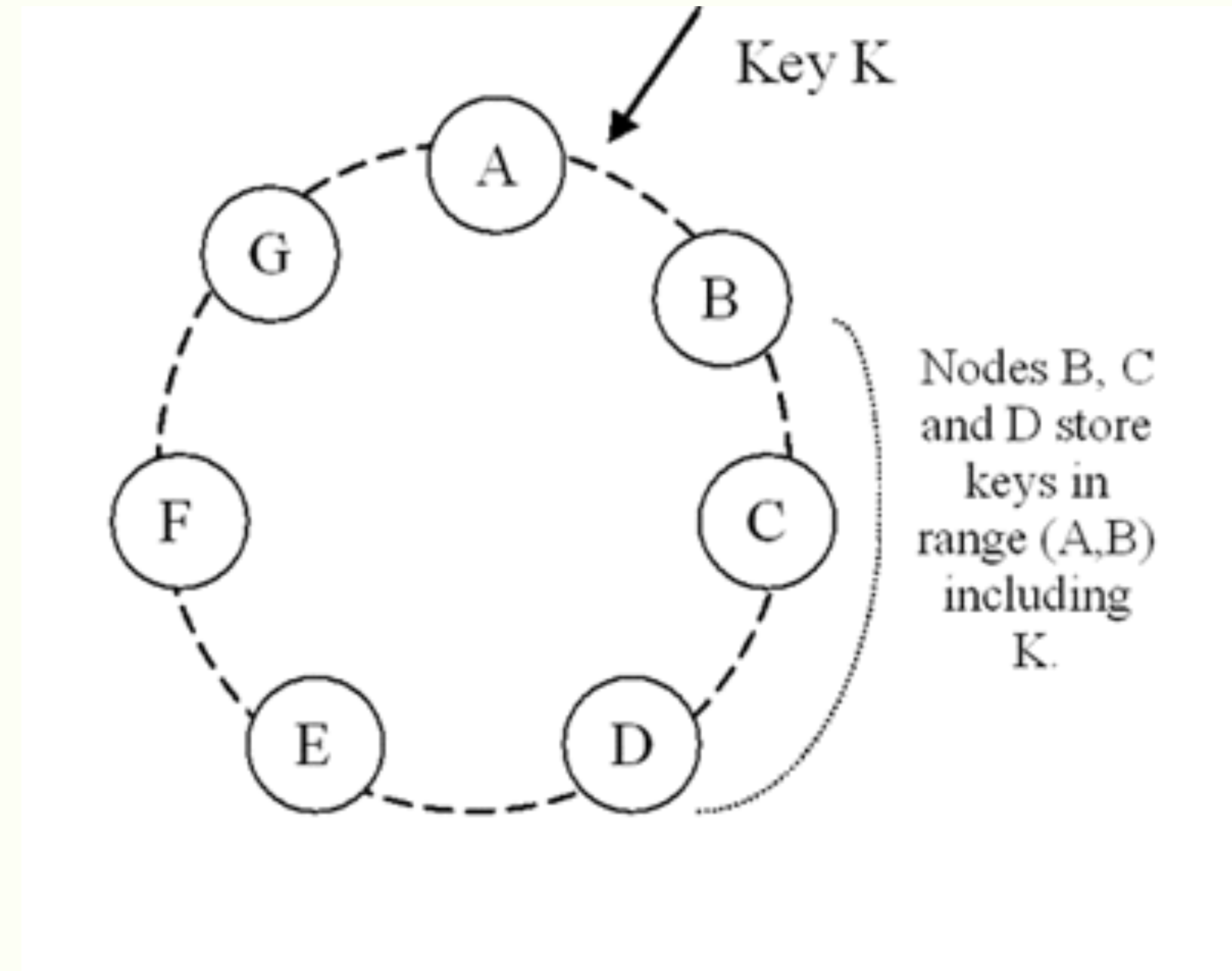


HBASE USE CASES

- **FACEBOOK MESSAGING**
- **HIGH WRITE THROUGHPUT**
- **LOW LATENCY RANDOM READS**
- **STRONG CONSISTENCY**
- **HA VIA AUTOMATIC FAILOVER**

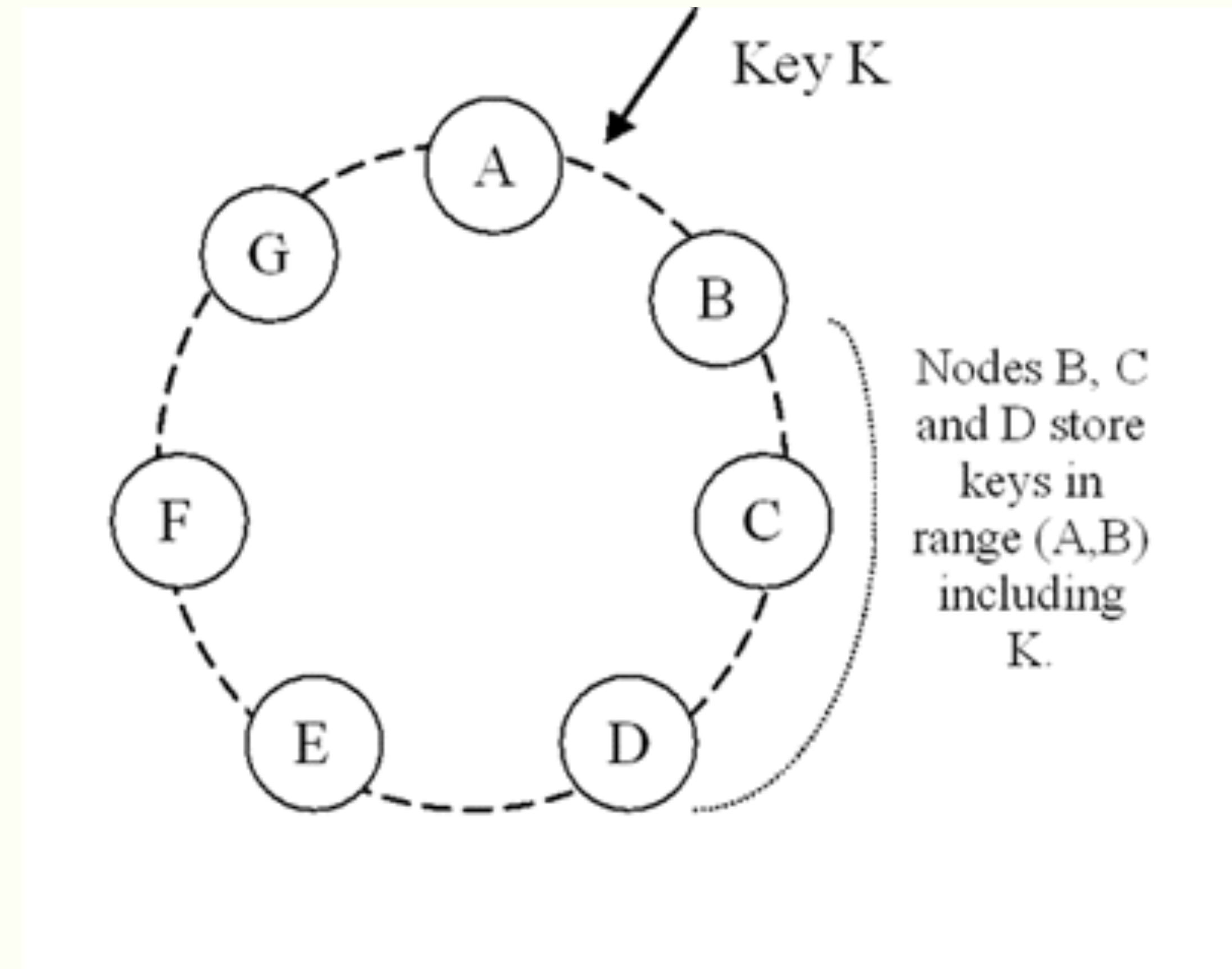
DYNAMO STORAGE

- HA KEY-VALUE STRUCTURED STORAGE
- DATA PARTITIONED USING CONSISTENT HASHING
- LINEAR SCALABILITY
- MASTER-MASTER
- CONSISTENCY FACILITATED BY OBJECT VERSIONING
- HINTED-HANDOFF
- ANTI-ENTROPY USING MERKLE-TREE



APACHE CASSANDRA

- BIG-TABLE + DYNAMO
- HIGH AVAILABILITY
- INCREMENTAL SCALABILITY
- EVENTUALLY CONSISTENT
- TUNNABLE CONSISTENCY AND LATENCY
- NO SINGLE POINT OF FAILURE

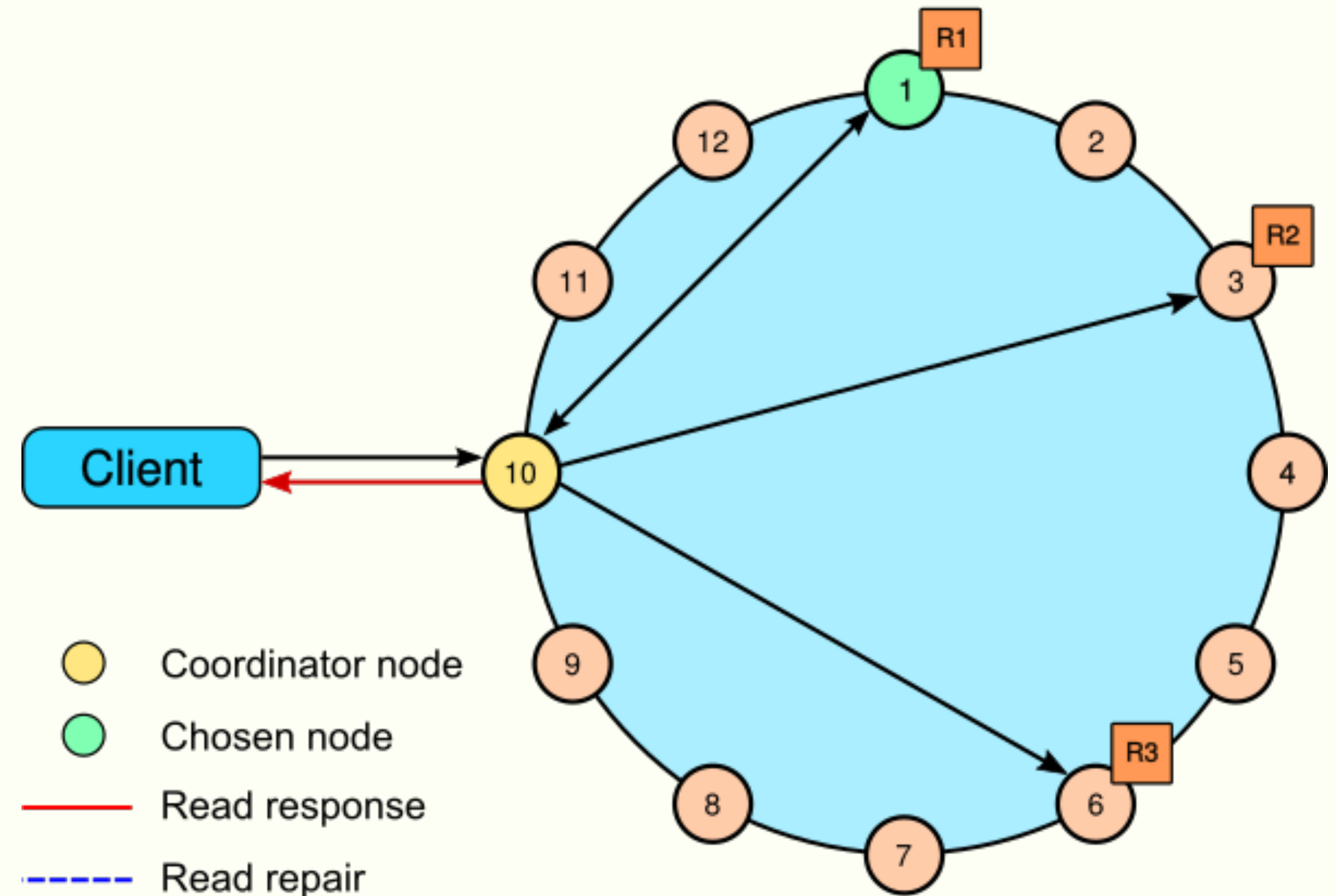


KEY CONCEPTS

- ▶ MASTER-MASTER
- ▶ COORDINATOR
- ▶ QUERY “PROXYING”



Single data center cluster with 3 replica nodes and consistency set to ONE

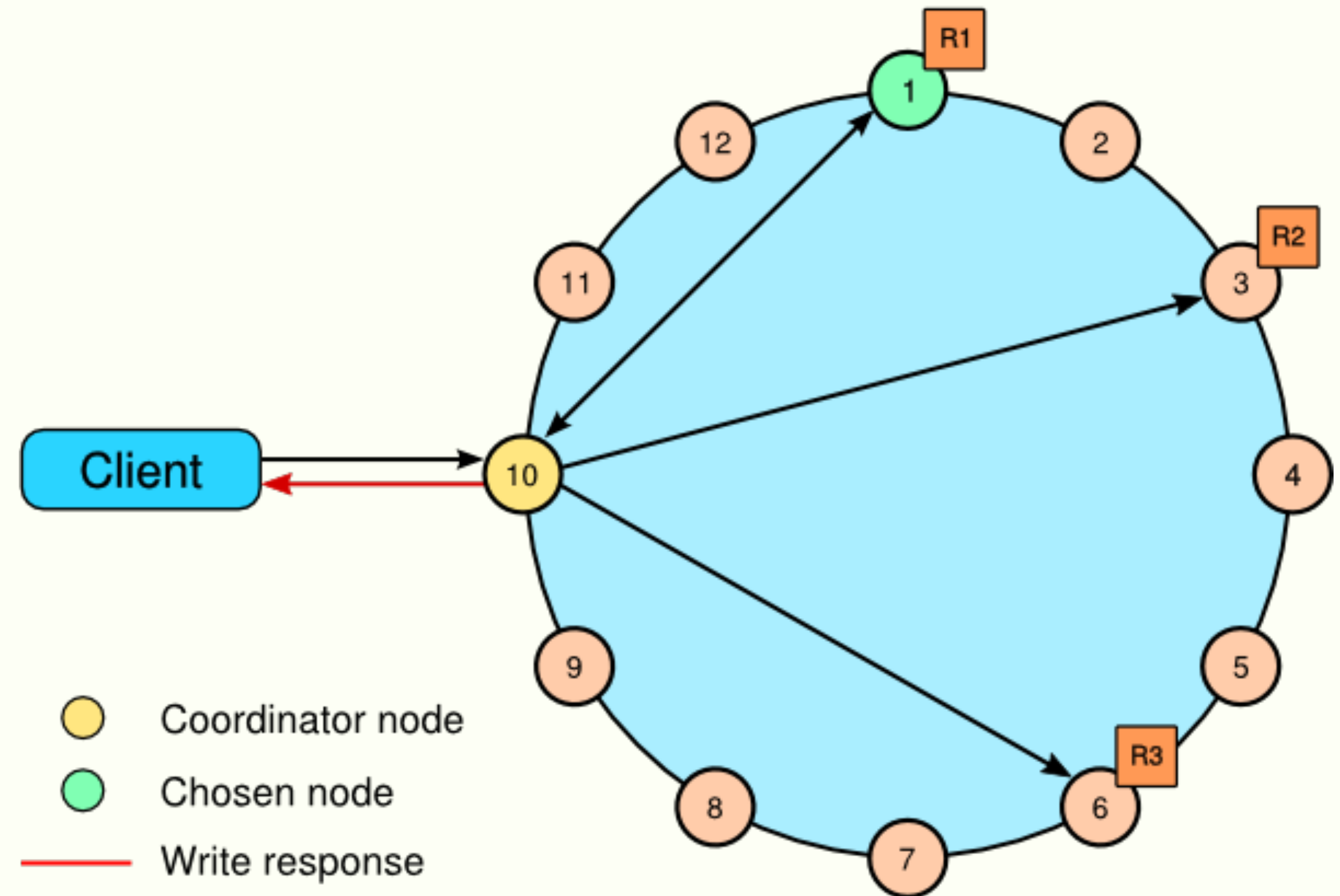


REPLICATION FACTOR

- RF = TOTAL NUMBER OF REPLICAS
- REPLICATION STRATEGY



Single data center cluster with 3 replica nodes and consistency set to ONE



CONSISTENCY LEVEL

▸ TUNABLE CONSISTENCY

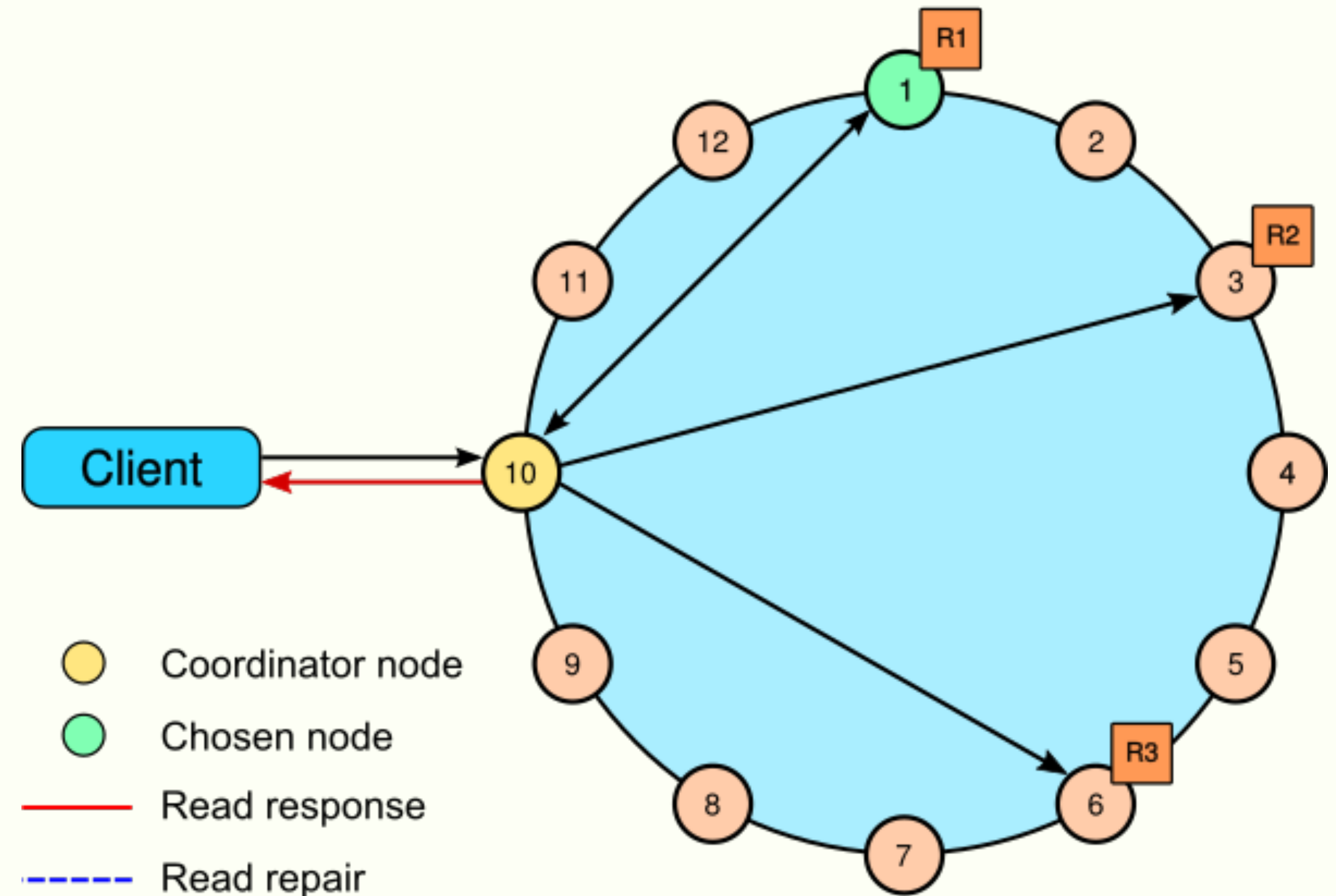
▸ READ

▸ WRITE



Cassandra

Single data center cluster with 3 replica nodes and consistency set to ONE



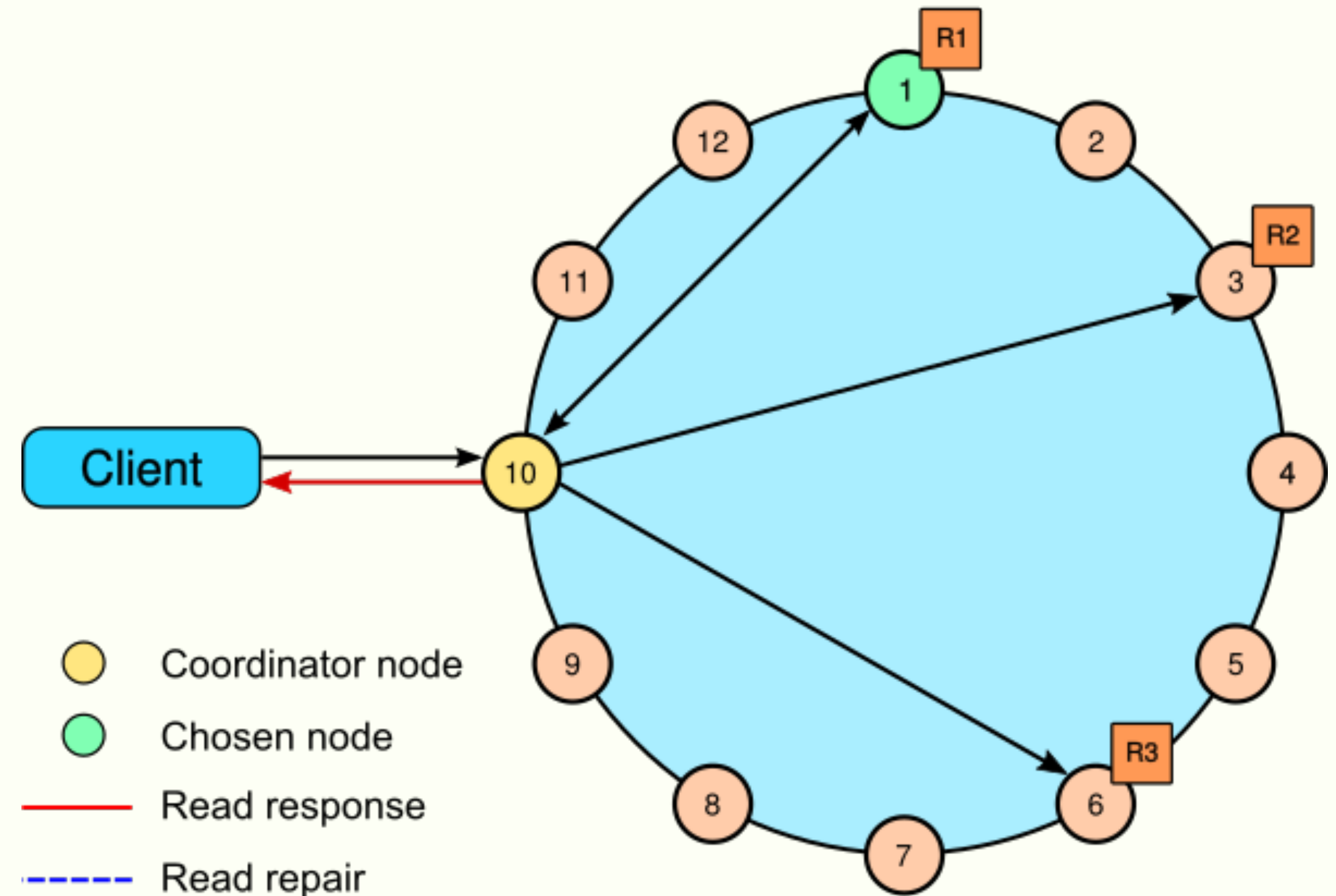
CONSISTENCY LEVEL

- READ REQUEST
- DIGEST QUERY
- REPAIRED READ
- CL=1



Cassandra

Single data center cluster with 3 replica nodes and consistency set to ONE



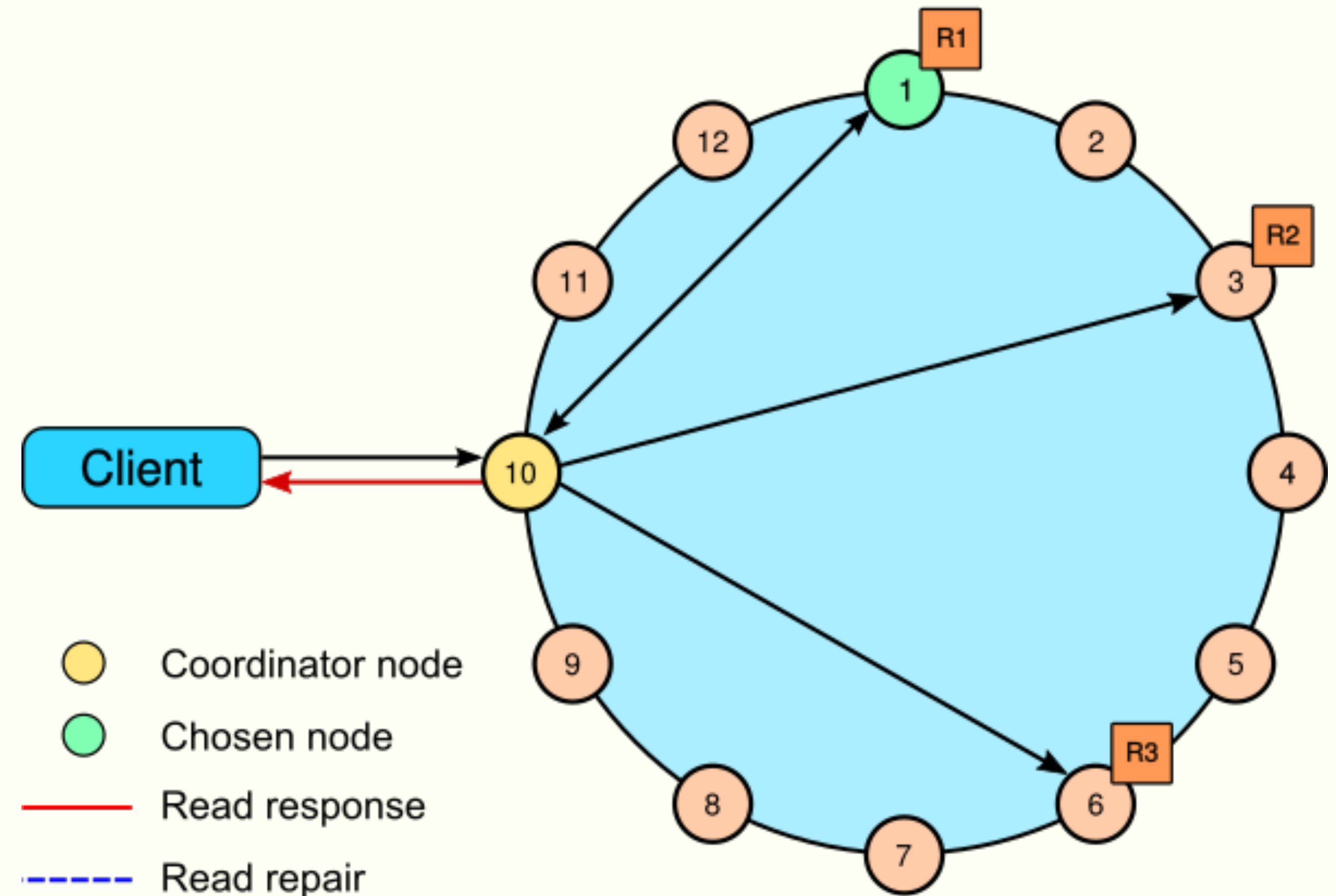
READ

- RF = 3
- CL=1



Cassandra

Single data center cluster with 3 replica nodes and consistency set to ONE



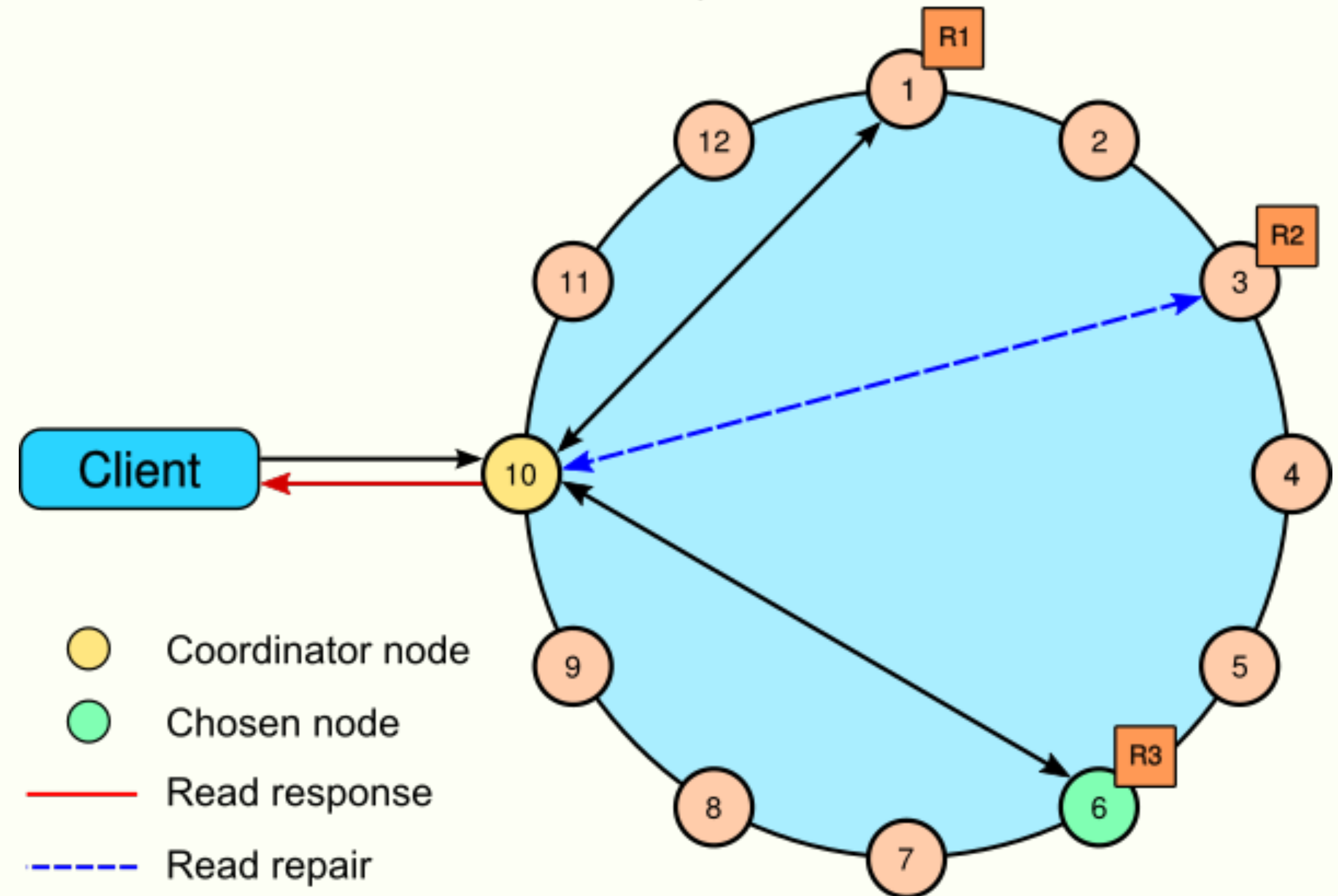
READ

- RF = 3
- CL= QUORUM



Cassandra

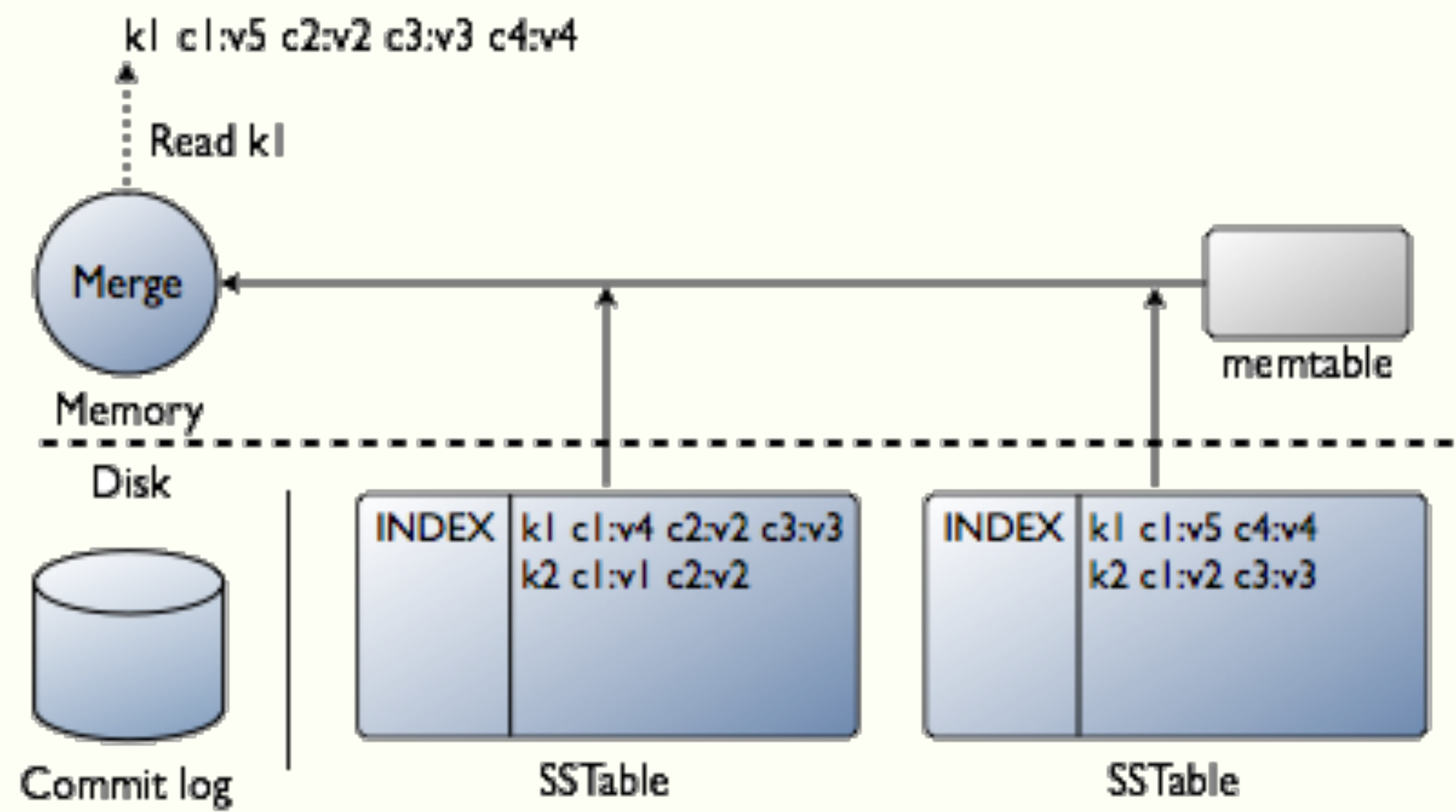
Single data center cluster with 3 replica nodes and consistency set to QUORUM



STORAGE CONCEPTS



READS



ANTI ENTROPY

- HINTED HANDOFF
- READ REPAIR
- FULL VS. INCREMENTAL REPAIR



USE CASES - **NETFLIX**

- 90 CLUSTERS, 2700 NODES, 4 DC
- NEARLY ALL DATA
- FILM METADATA
- USER RATINGS
- RECOMMENDATIONS



GRAPH DATABASES

NOT A SWISS-ARMY KNIFE



ANY QUESTIONS?

THANK YOU

@TOM BUJOK

