



Using Hadoop Together with MySQL for Data Analysis

Alexander Rubin
Principle Architect, Percona
April 18, 2015

About Me

Alexander Rubin, Principal Consultant, Percona

- Working with MySQL for over 10 years
 - Started at MySQL AB, Sun Microsystems, Oracle (MySQL Consulting)
 - Worked at Hortonworks (Hadoop company)
 - Joined Percona in 2013

Agenda

- Hadoop Use Cases
- Big Data Analytics with Hadoop
- Star Schema benchmark
- MySQL and Hadoop Integration

Hadoop: when it makes sense

BIG DATA



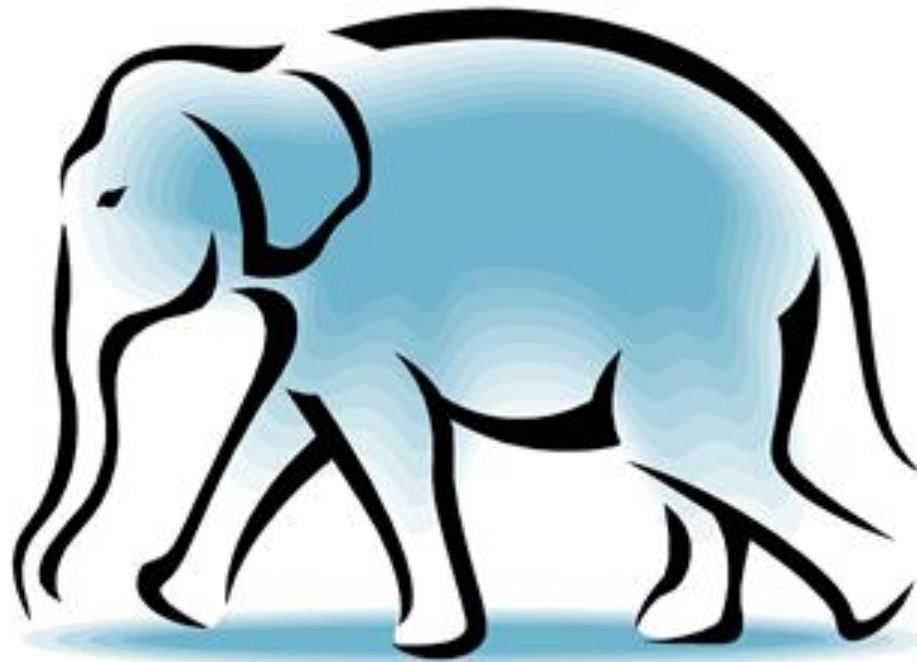
Big Data

“3Vs” of big data

- Volume
 - Petabytes
- Variety
 - Any type of data – usually unstructured/raw data
 - No normalization
- Velocity
 - Data is collected at a high rate

http://en.wikipedia.org/wiki/Big_data#Definition

Inside Hadoop



Where is my data?

Data Nodes

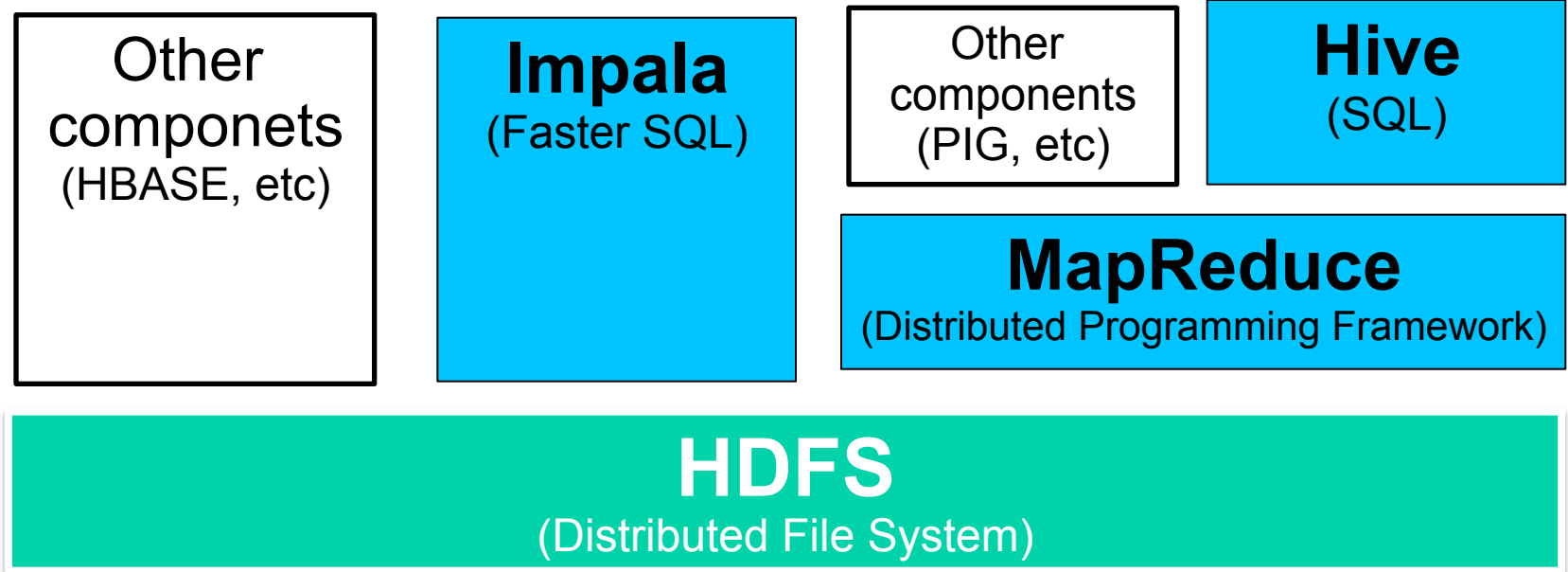


HDFS

(Distributed File System)

- HDFS = Data is spread between MANY machines
- Also Amazon S3 is supported

Inside Hadoop



- HDFS = Data is spread between MANY machines
- Write files, “append-only” mode



Hive

- SQL level for Hadoop
- ***Translates SQL to Map-Reduce jobs***
- Schema on Read – does not check the data on load

Hive Example

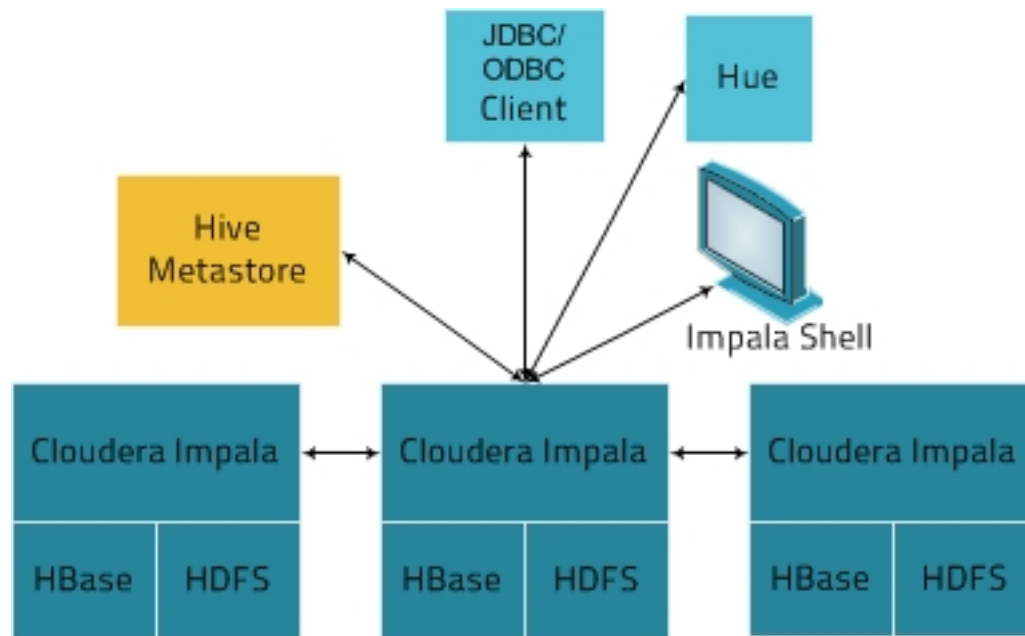
```
hive> create table lineitem (  
  l_orderkey int,  
  l_partkey int,  
  l_suppkey int,  
  l_linenum int,  
  l_quantity double,  
  l_extendedprice double,  
  ...  
  l_shipmode string,  
  l_comment string)  
row format delimited fields terminated by '|';
```

Hive Example

```
hive> create external table lineitem (  
...)  
row format delimited fields terminated by '|'  
location '/ssb/lineorder/';
```

Impala: Faster SQL

Not based on Map-Reduce, directly get data from HDFS



<http://www.cloudera.com/content/cloudera-content/cloudera-docs/Impala/latest/Installing-and-Using-Impala/Installing-and-Using-Impala.html>

Hadoop vs MySQL



Hadoop vs. MySQL for BigData



- ☐ Indexes
- ☐ Partitioning
- ☐ “Sharding”



- ☐ Full table scan
- ☐ Partitioning
- ☐ Map/Reduce

Hadoop (vs. MySQL)

- No indexes
 - All processing is full scan
 - BUT: ***distributed and parallel***
- No transactions
- High latency (usually)

Indexes (BTree) for Big Data challenge

- Creating an index for Petabytes of data?
- Updating an index for Petabytes of data?
- Reading a terabyte index?
- Random read of Petabyte?

Full scan in parallel is better for big data

ETL vs ELT



1. Extract data from external source
2. Transform before loading
3. Load data into MySQL

1. Extract data from external source
2. Load data into Hadoop
3. Transform data/
Analyze data/
Visualize data;

Parallelism

Example:

Loading wikistat into MySQL



<http://dumps.wikimedia.org/other/pagecounts-raw/>

1. Extract data
from external
source

Wikipedia page counts –
download, >10TB

2. Load data
into MySQL
and
Transform

```
load data local infile '$file'
into table wikistats.wikistats_full
CHARACTER SET latin1
FIELDS TERMINATED BY ' '
(project_name, title, num_requests,
content_size)
set request_date =
STR_TO_DATE('$datestr',
'%Y%m%d %H%i%S'),
title_md5=unhex(md5(title));
```

Load timing per hour of wikistat

- InnoDB: 52.34 sec
- MyISAM: 11.08 sec (+ indexes)
- 1 hour of wikistats = 1 minute
- 1 year will load in 6 days
 - (8765.81 hours in 1 year)
- 6 year = > 1 month to load



Loading wikistat into Hadoop

- Just copy files to HDFS...and create hive structure
- How fast to search?
 - Depends upon the number of *nodes*
- **1000 nodes** spark (hadoop fork) cluster
 - 4.5 TB, 104 Billion records
 - **Exec time: 45 sec**
 - **Scanning 4.5TB of data**
- <http://spark-summit.org/wp-content/uploads/2014/07/Building-1000-node-Spark-Cluster-on-EMR.pdf>



Archiving to Hadoop

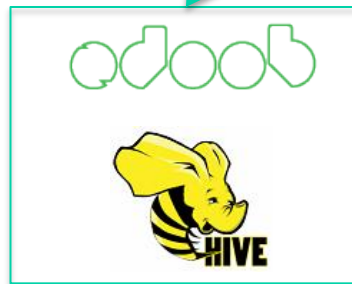
OLTP / Web site

Goal:
keep 100G – 1TB



ELT

BI / Data Analysis



Can store Petabytes
for archiving

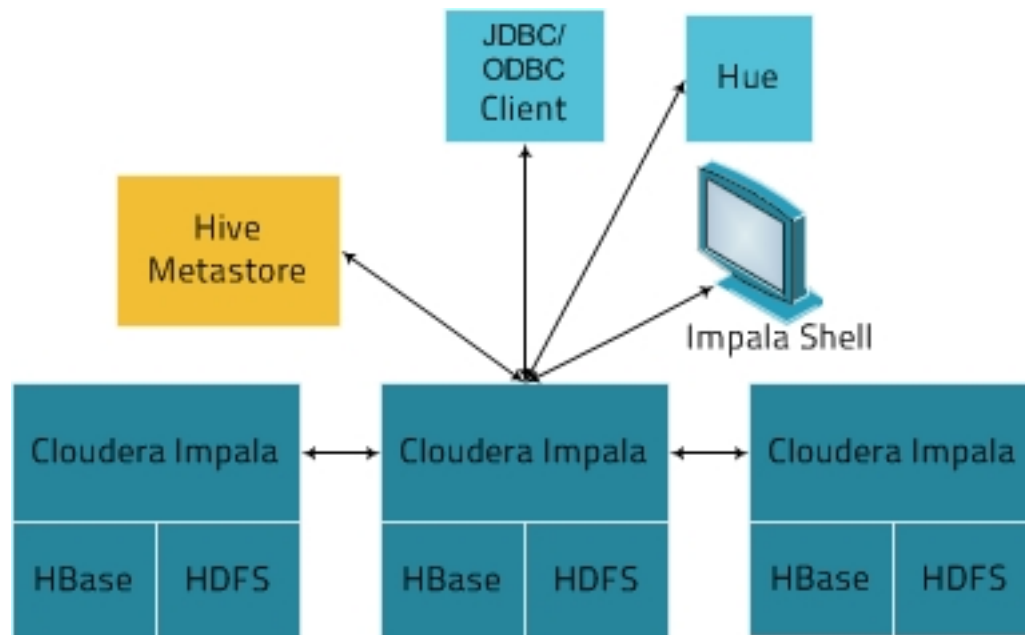
Hadoop Cluster

Star Schema benchmark

Variation of TPC-H with star-schema database design
<http://www.percona.com/docs/wiki/benchmark:ssb:start>



Cloudera Impala



Impala and Columnar Storage

Parquet: Columnar Storage for Hadoop

- Column-oriented storage
- Supports compression
- Star Schema Benchmark example
 - RAW data: 990GB
 - Parquet data: 240GB

<http://parquet.io/>

http://www.cloudera.com/content/cloudera-content/cloudera-docs/Impala/latest/Installing-and-Using-Impala/ciiu_parquet.html

Impala Benchmark: Table

1. Load data into HDFS

```
$ hdfs dfs -put /data/ssb/lineorder* /ssb/lineorder/
```

<http://www.percona.com/docs/wiki/benchmark:ssb:start>

Impala Benchmark: Table

2. Create external table

```
CREATE EXTERNAL TABLE lineorder_text
(
    LO_OrderKey bigint ,
    LO_LineNumber tinyint ,
    LO_CustKey int,
    LO_PartKey int,
    ...
) ROW FORMAT DELIMITED FIELDS TERMINATED BY '|' LINES TERMINATED
BY '\n' STORED AS TEXTFILE LOCATION '/ssb/lineorder';
```

<http://www.percona.com/docs/wiki/benchmark:ssb:start>

Impala Benchmark: Table

3. Convert to Parquet

```
Impala-shell> create table lineorder  
like lineorder_text
```

```
STORED AS PARQUETFILE;
```

```
Impala-shell> insert into lineorder  
select * from lineorder_text;
```

<http://www.percona.com/docs/wiki/benchmark:ssb:start>

Impala Benchmark: SQLs

Q1.1

<http://www.percona.com/docs/wiki/benchmark:ssb:start>

```
select
    sum(lo_extendedprice*lo_discount) as
revenue
from
    lineorder, dates
where
    lo_orderdate = d_datekey
    and d_year = 1993
    and lo_discount between 1 and 3
    and lo_quantity < 25;
```

Impala Benchmark: SQLs

Q3.1

```
select  c_nation, s_nation, d_year, sum(lo_revenue) as revenue
from
    customer, lineorder, supplier, dates
where
    lo_custkey = c_custkey and lo_suppkey = s_suppkey
    and lo_orderdate = d_datekey
    and c_region = 'ASIA' and s_region = 'ASIA'
    and d_year >= 1992 and d_year <= 1997
group by c_nation, s_nation, d_year
order by d_year asc, revenue desc;
```

Partitioning for Impala / Hive

Partitioning by Year

```
CREATE TABLE lineorder
```

```
(
```

```
    LO_OrderKey bigint ,
```

```
    LO_LineNumber tinyint ,
```

```
    LO_CustKey int,
```

```
    LO_PartKey int,
```

```
    ...
```

```
) partitioned by (year int);
```

Will create a “virtual” field on “year”



```
select * from lineorder where ... year = '2013'
```

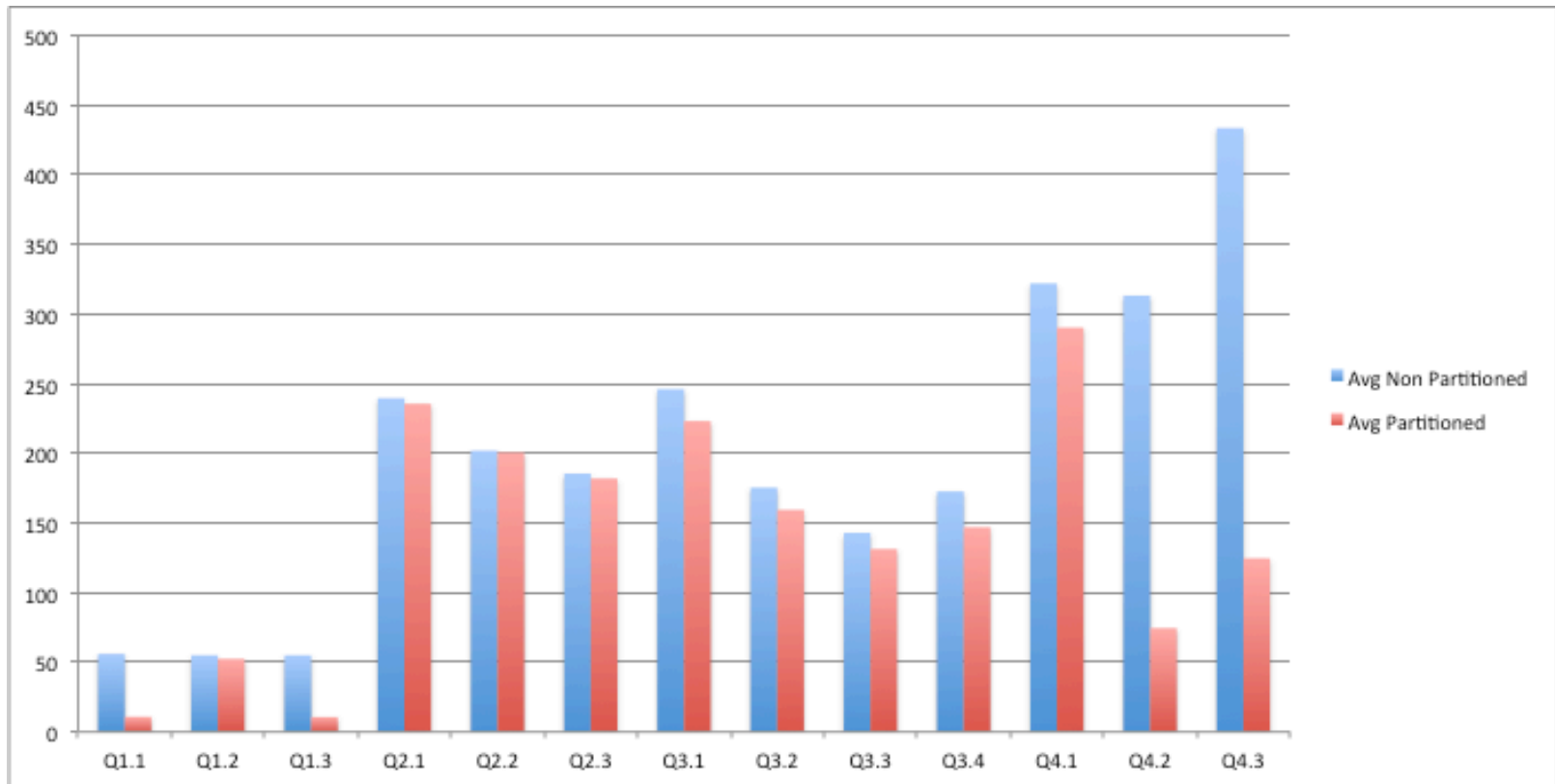
... will only read the “2013” partition (much less disk IOs!)

Impala Benchmark

CDP Hadoop with Impala

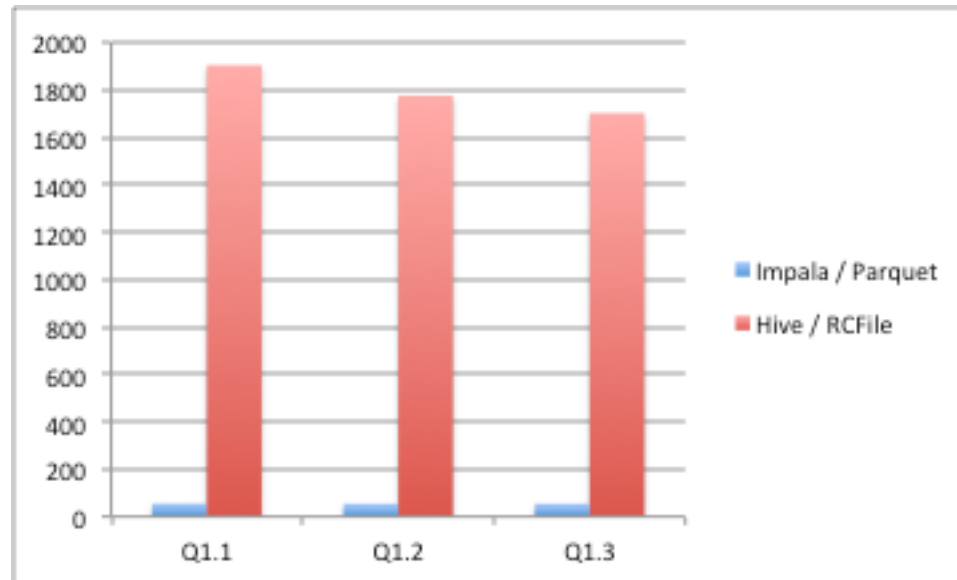
- 10 m1.xlarge Amazon EC2
- RAW data: 990GB
- Parquet data: 240GB

Impala Benchmark



Scanning 1Tb of row data 50-200 seconds
10 xlarge nodes on Amazon EC2

Impala vs Hive Benchmark



10 xlarge nodes on Amazon EC2
Impala and Hive

Amazon Elastic Map Reduce



Hive on Amazon S3

```
hive> create external table lineitem (  
  l_orderkey int, l_partkey int, l_suppkey int,  
  l_linenumbr int, l_quantity double,  
  l_extendedprice double, l_discount double, l_tax  
  double, l_returnflag string,  
  l_linestatus string, l_shipdate string,  
  l_commitdate string, l_receiptdate string,  
  l_shipinstruct string, l_shipmode string,  
  l_comment string)  
row format delimited fields terminated by '|'  
location 's3n://data.s3ndemo.hive/tpch/lineitem';
```

Amazon Elastic Map Reduce

- Store data on S3
- Prepare SQL file (create table, select, etc)
- Run Elastic Map Reduce
 - Will start N boxes then stop them
- Results loaded to S3

Hadoop and MySQL Together

Integrating MySQL and Hadoop



Integration: Hadoop -> MySQL

Hadoop Cluster



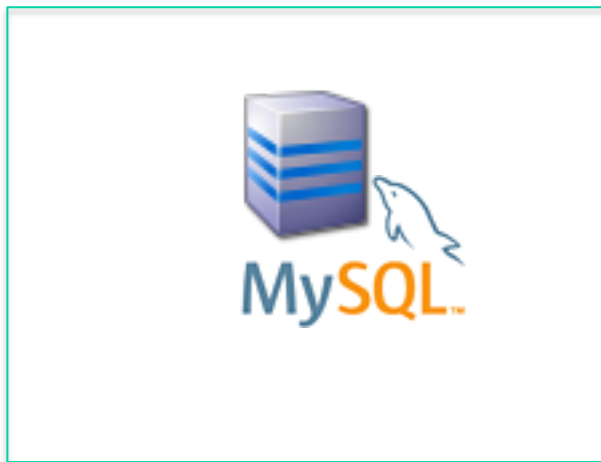
MySQL Server



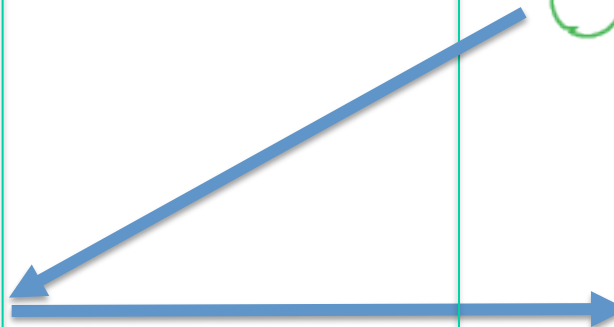
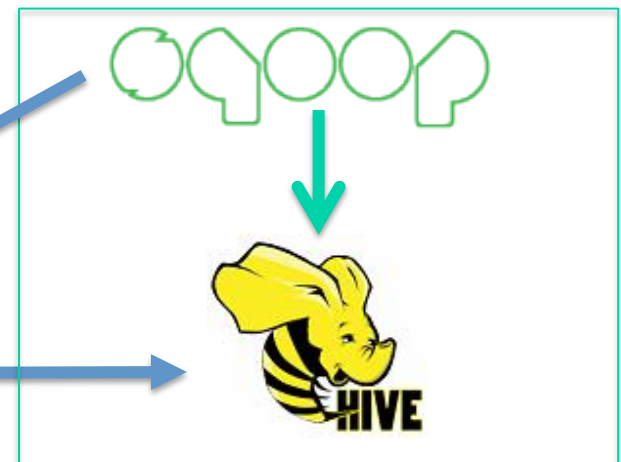
MySQL -> Hadoop: Sqoop

```
$ sqoop import  
--connect jdbc:mysql://mysql_host/db_name  
--table ORDERS  
--hive-import
```

MySQL Server

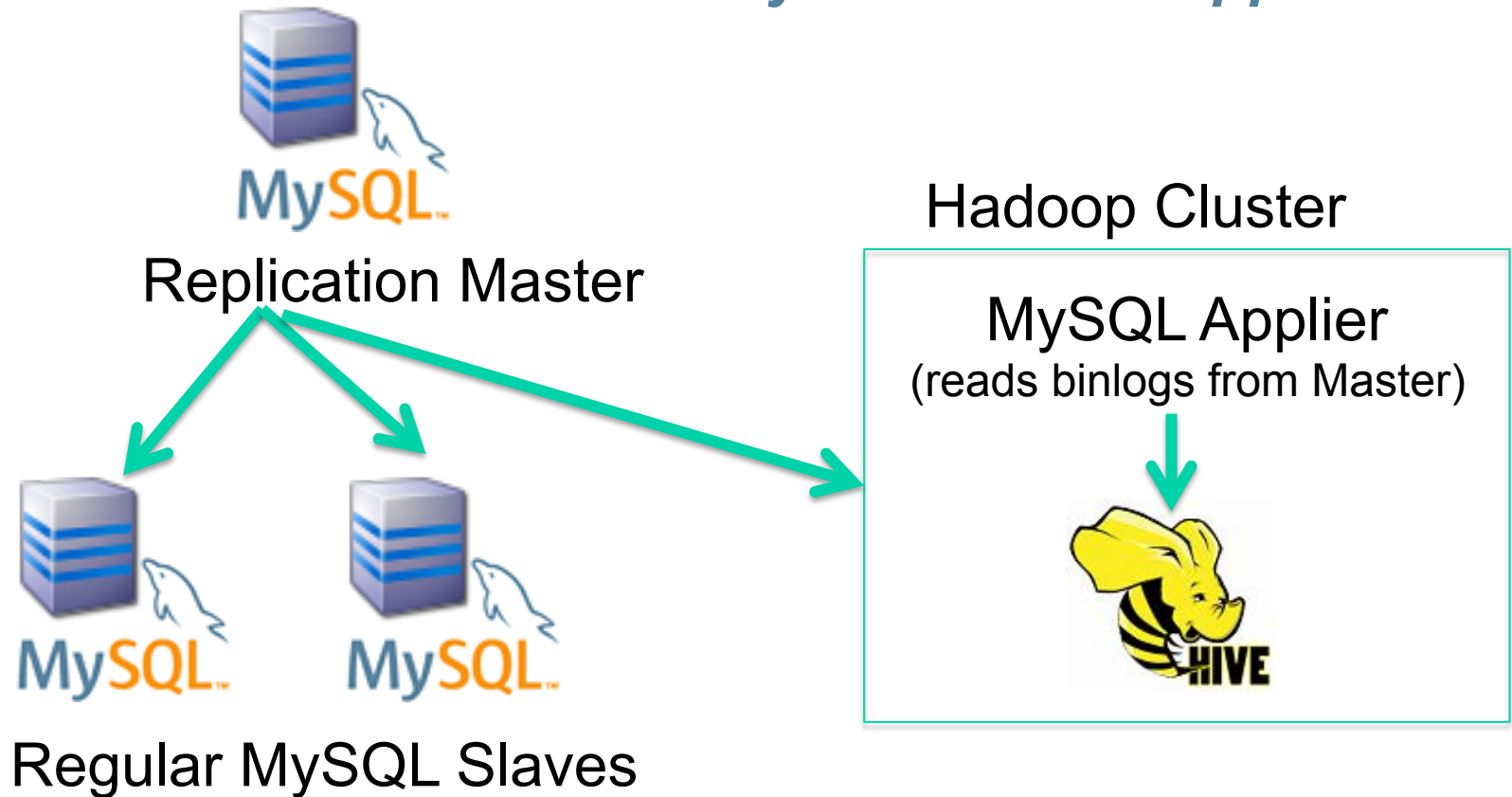


Hadoop Cluster



MySQL to Hadoop: Hadoop Applier

Only inserts are supported



MySQL to Hadoop: Hadoop Applier

- Download from: <http://labs.mysql.com/>
- Alpha version right now
- We need to write code how to process data

Questions?



Thank you!

Blog: <http://www.arubin.org>