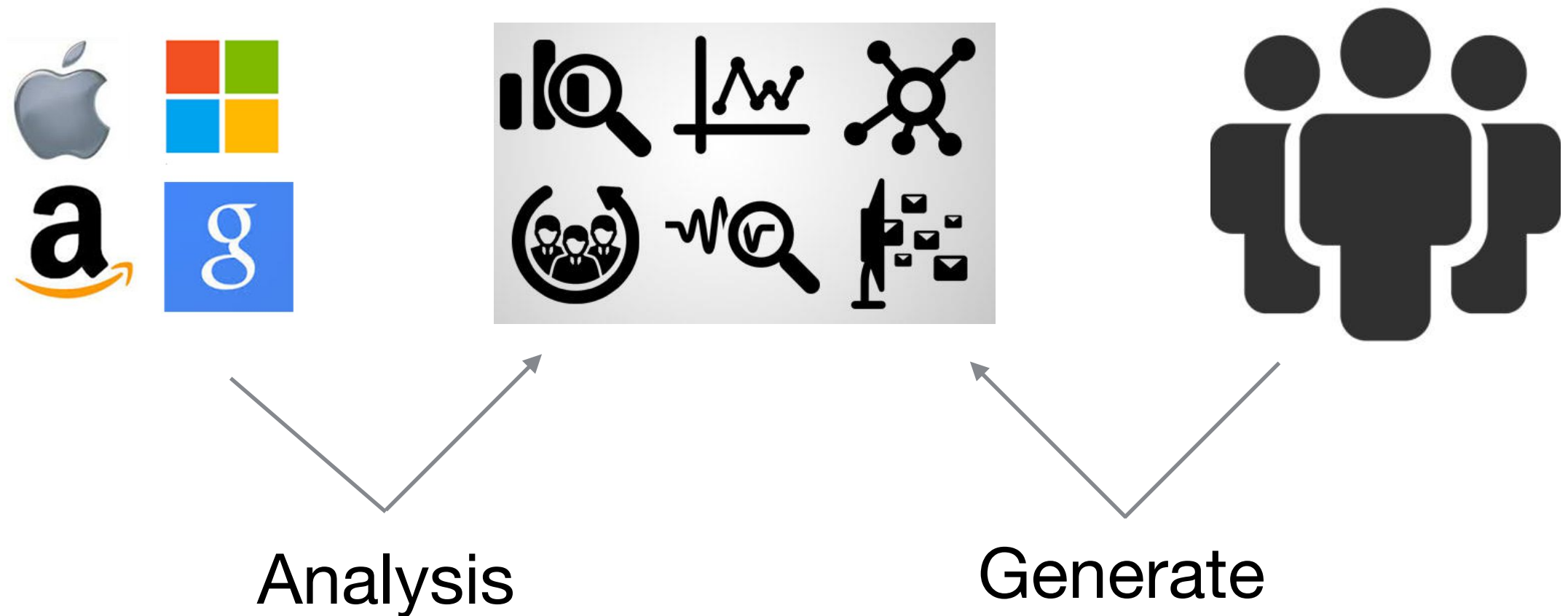


Bohr: Similarity Aware Geo-distributed Data Analytics

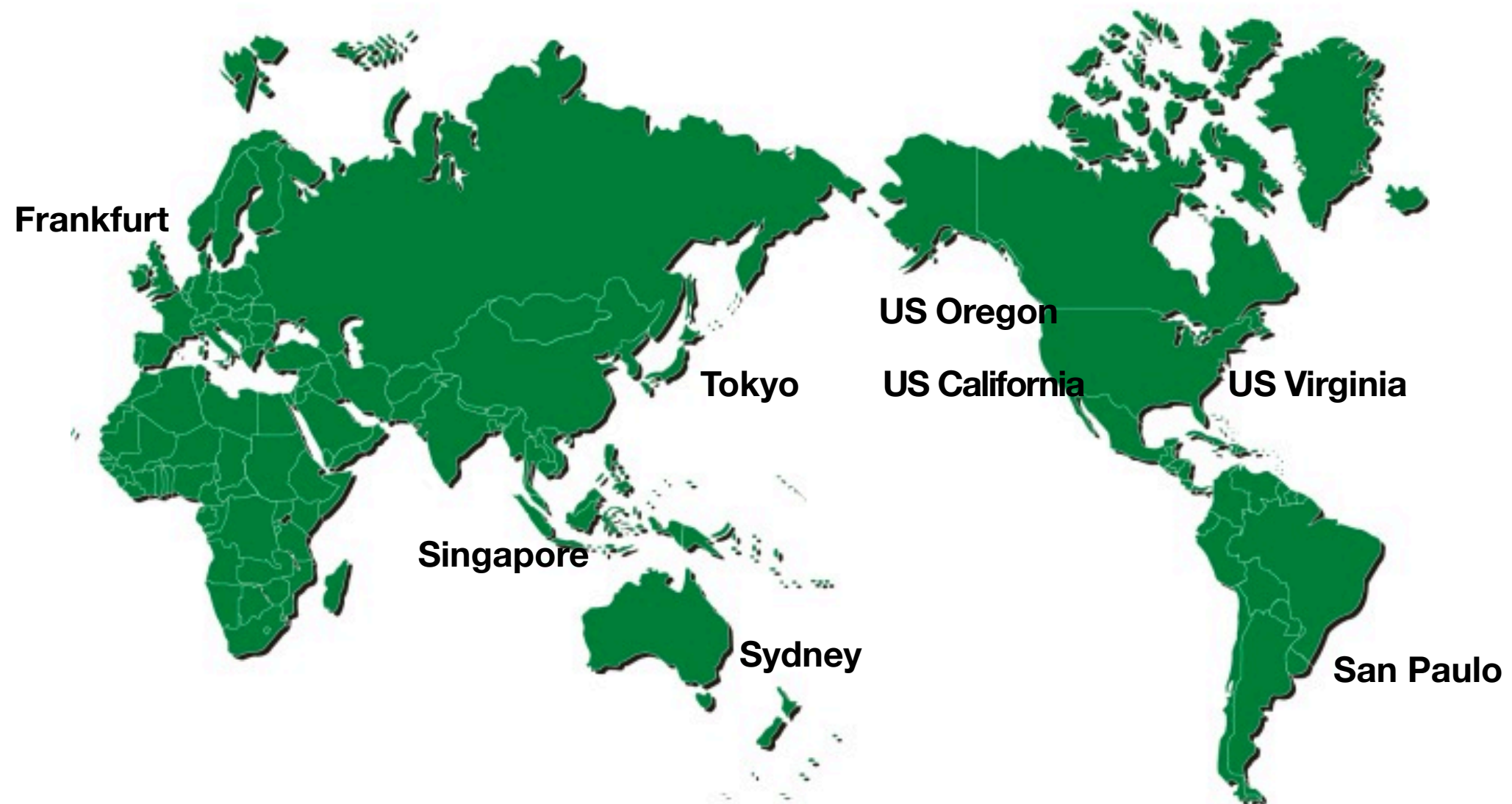
Hangyu Li, Hong Xu, Sarana Nutanong
City University of Hong Kong



Big Data Analytics

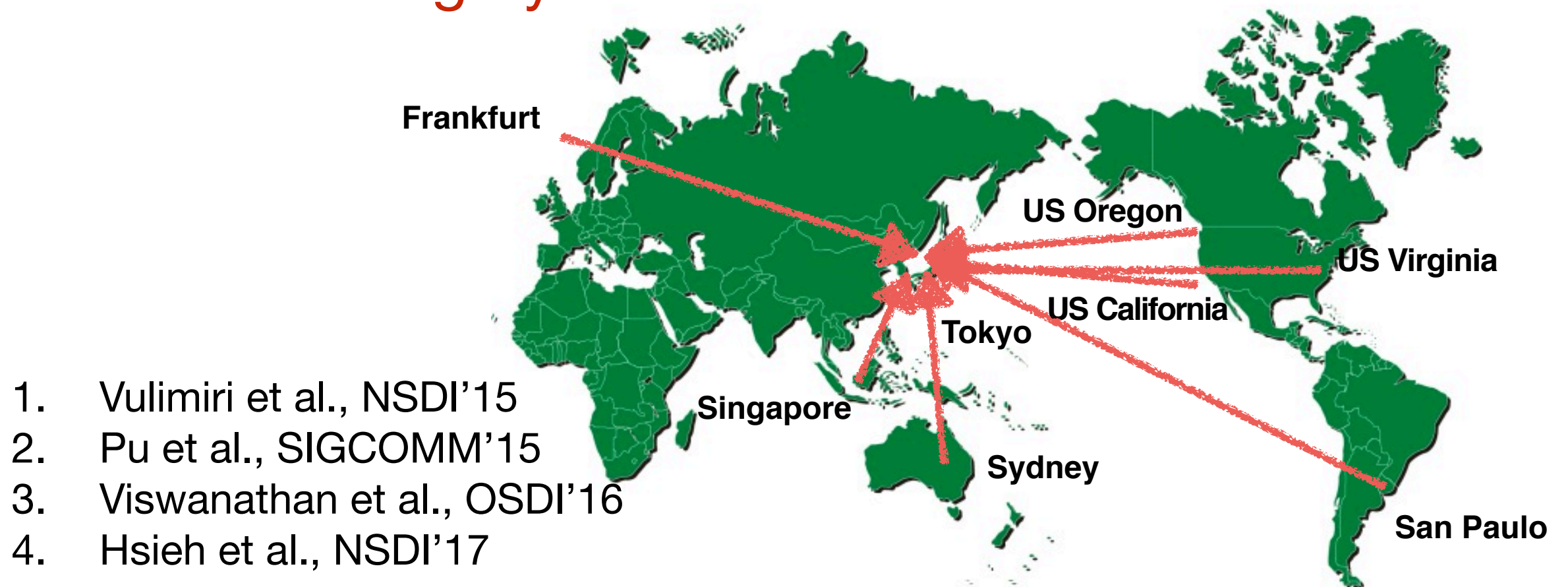


Data are geo-distributed



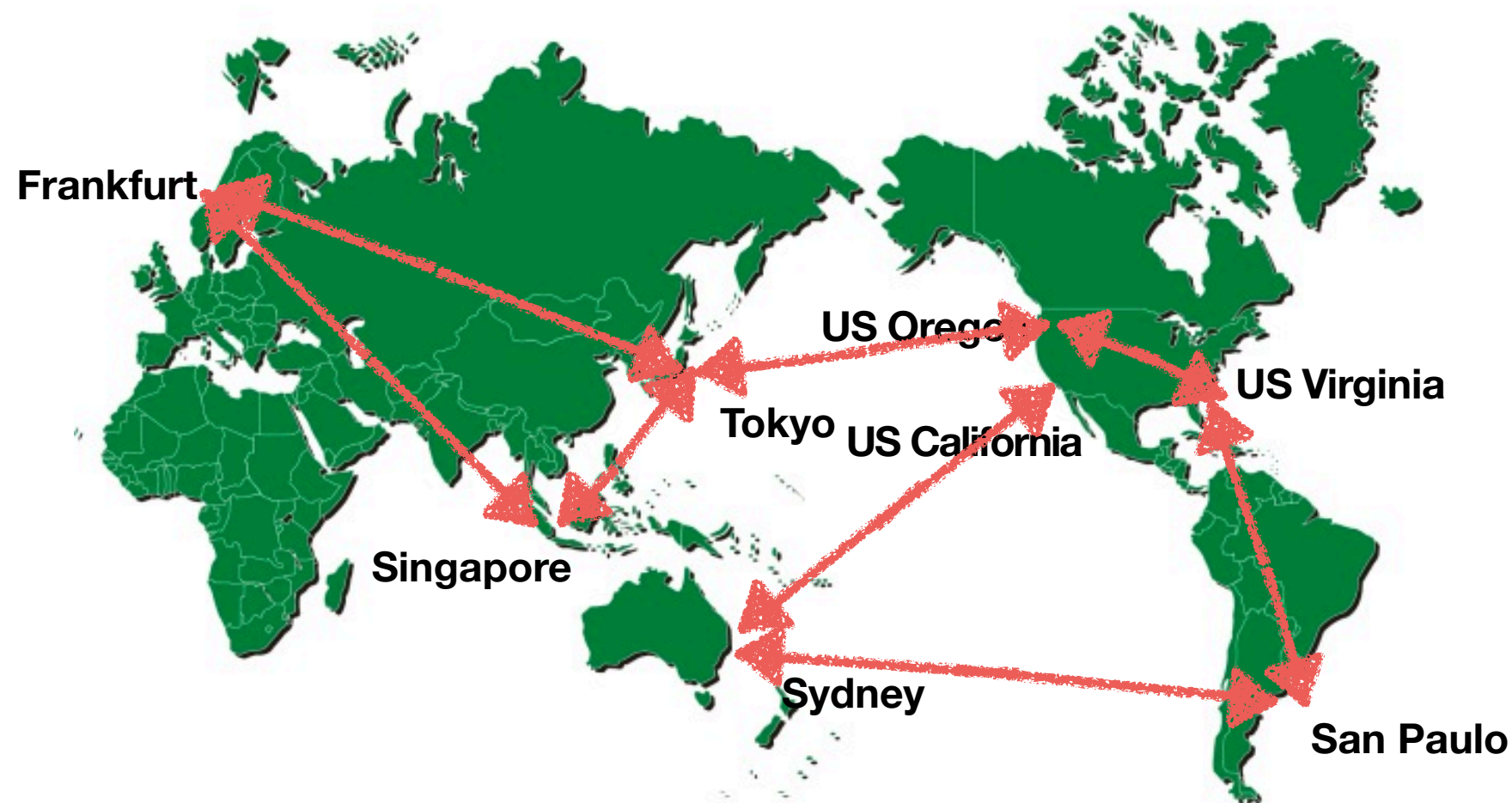
Centralizing data is infeasible

- Moving data through wide area networks(WANS) can be **extremely slow**.
- **Data sovereignty laws**.



Processing queries in-place

Bottleneck site exists.



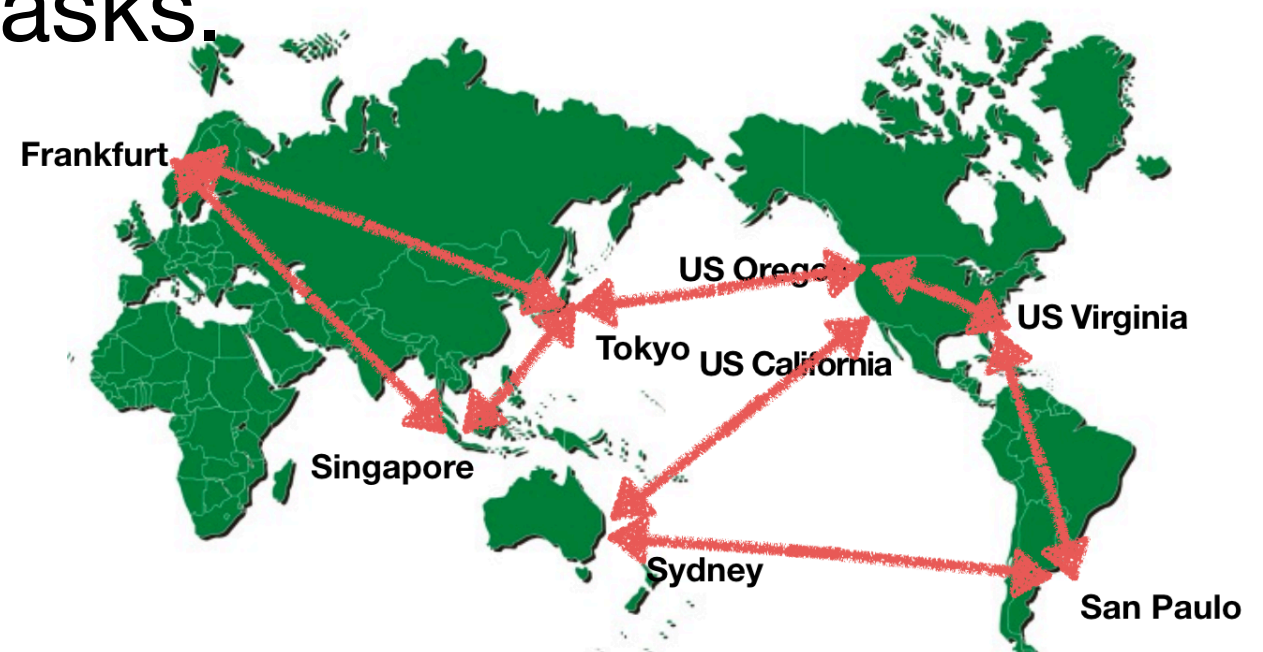
Processing queries in-place

System	Workload	Task / Data Placement	Optimize aspect
JetStream	streaming	no	WAN bandwidth usage
Geode	batch	data	WAN bandwidth usage
Iridium	batch	both	query response time

Processing queries in-place

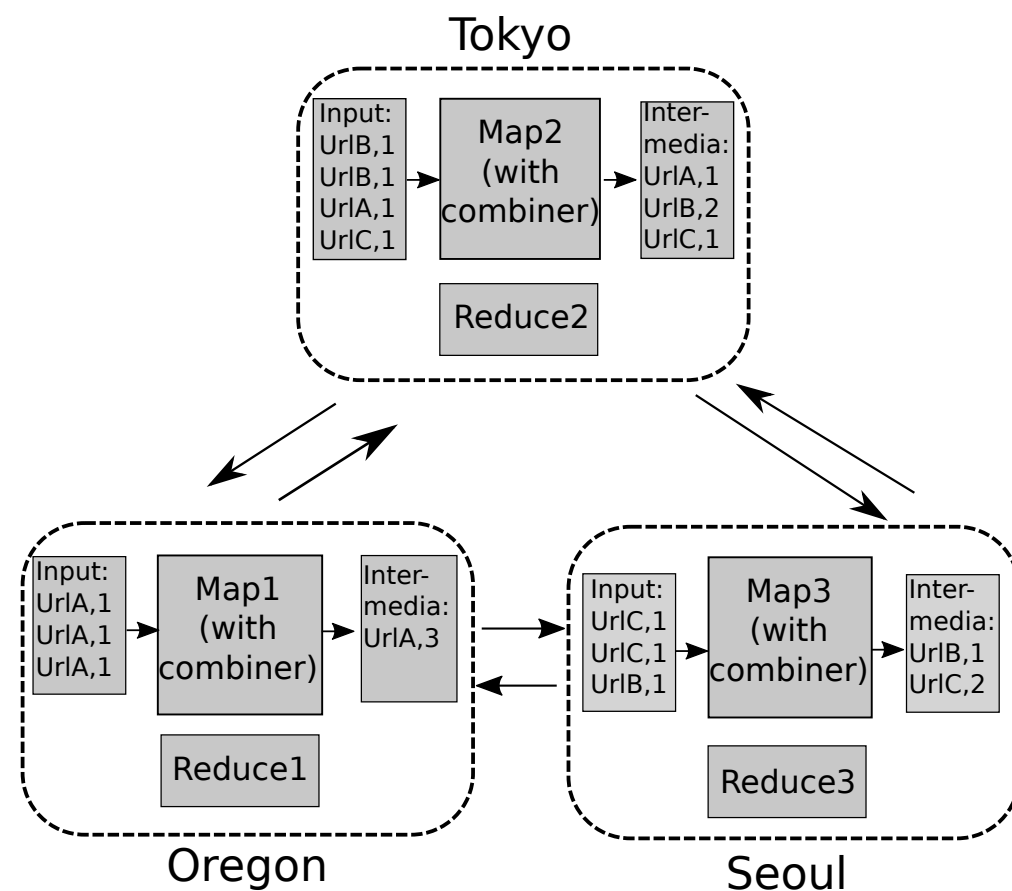
Iridium (Pu et al. SIGCOMM'15)

- **Recurring** queries, abundant computation resource.
- Transfer data out from the bottleneck site.
- Re-schedule reducer tasks.



Iridium can be improved

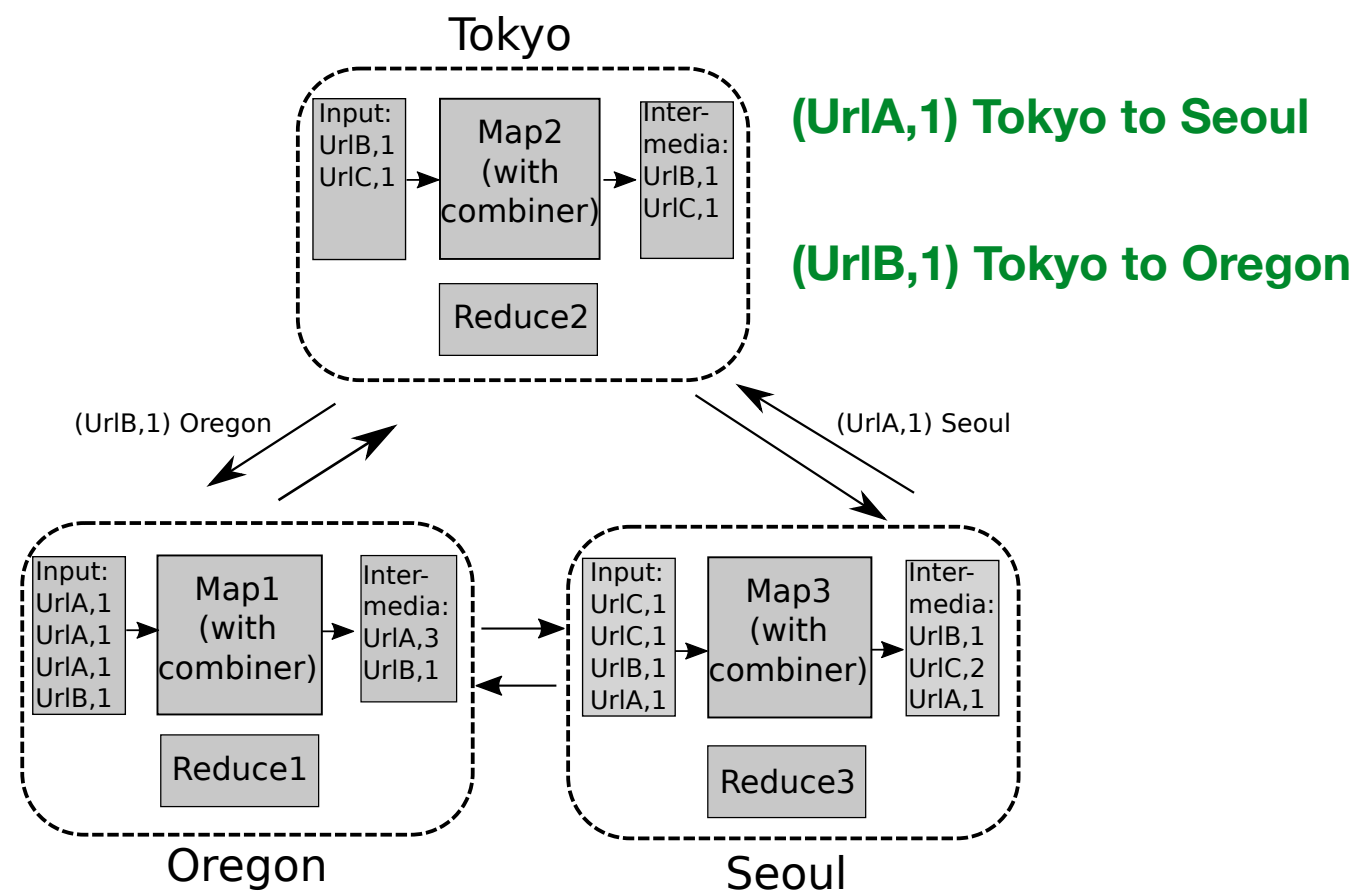
- **similarity oblivious**
- **data reduction ratio is constant**



Site	Intermed-iate record
Tokyo	3
Oregon	1
Seoul	2
Total	6

Iridium can be improved

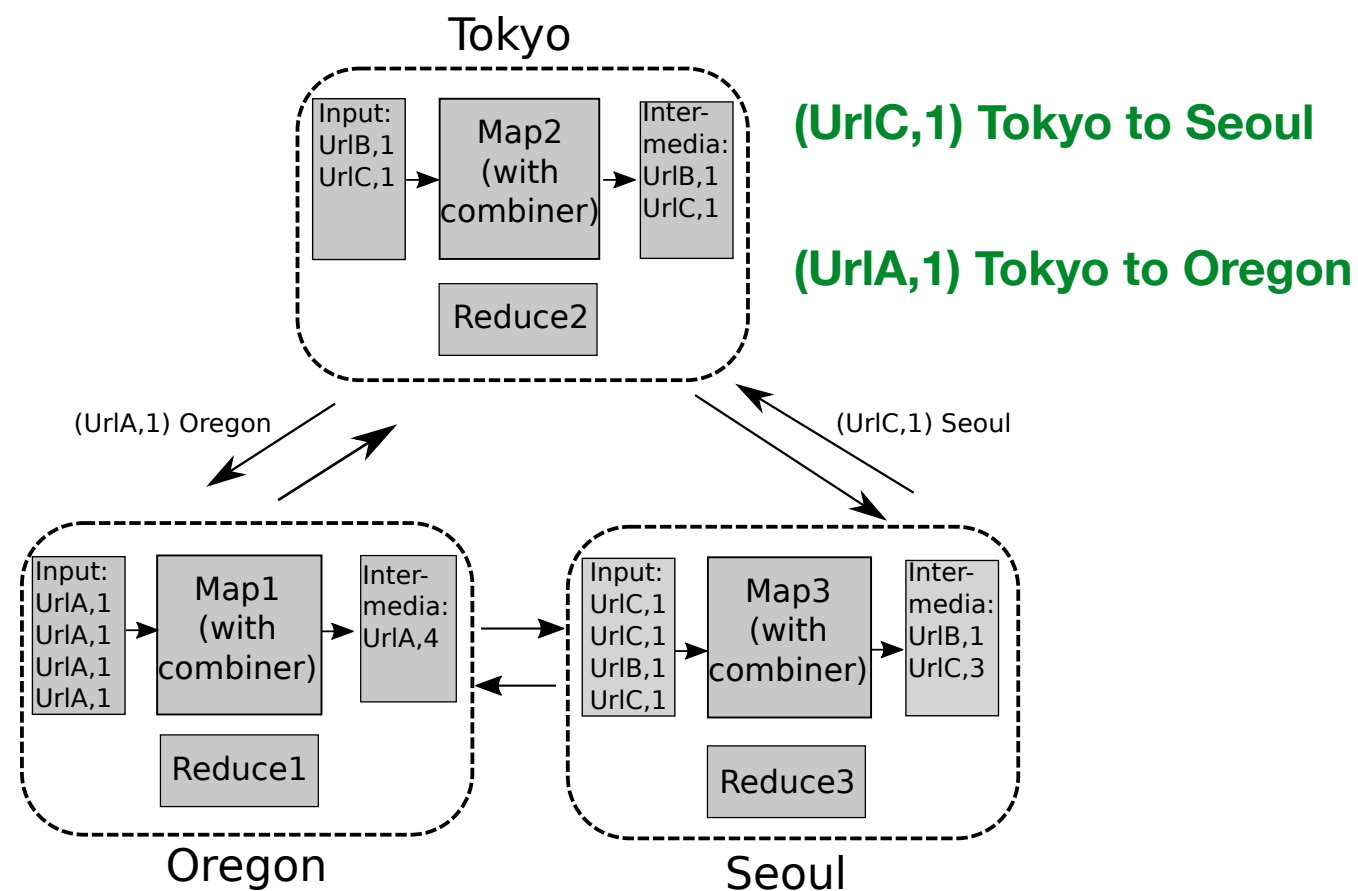
- **similarity oblivious**
- **data reduction ratio is constant**



Site	Intermediate record
Tokyo	2
Oregon	2
Seoul	3
Total	7

Iridium can be improved

- **similarity oblivious**
- **data reduction ratio is constant**



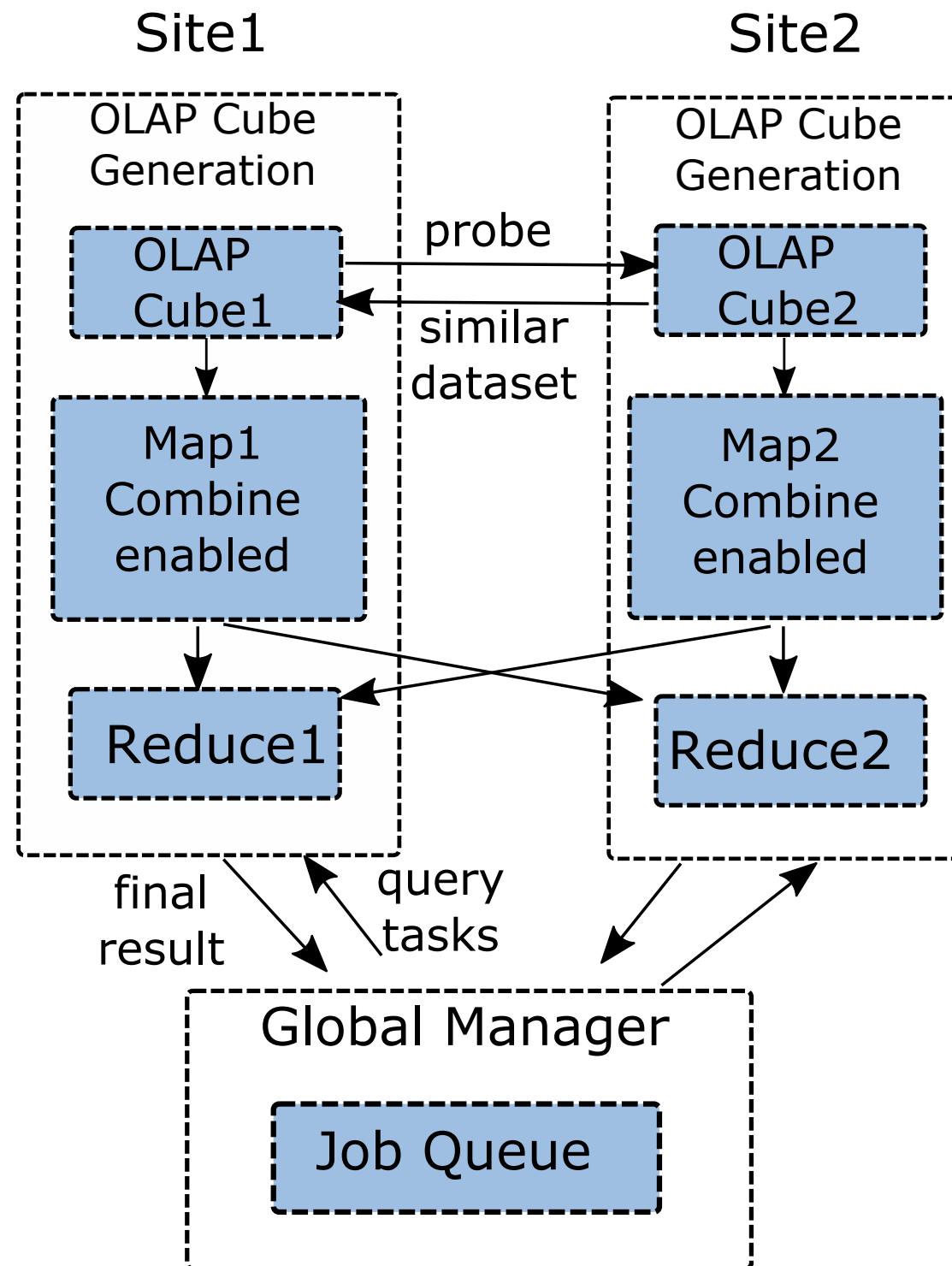
Site	Intermediate record
Tokyo	2
Oregon	1
Seoul	2
Total	5

Bohr Design

Key ideas:

- **Recurring** queries, abundant computation resource.
- **Similarity aware** data transfer.
- **Measuring and predicting** the data reduction ratio.

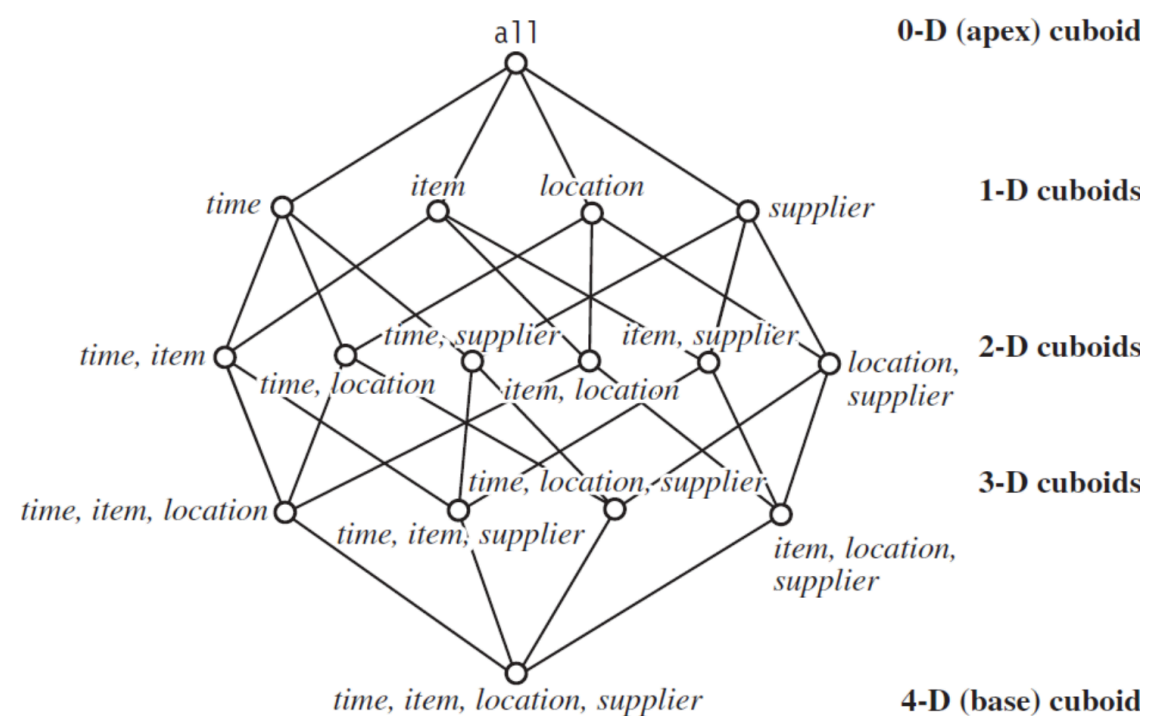
Bohr Design



Why using OLAP Cube?

- Sorting dataset according to the attribute.
- Using record level similarity score.

<i>location (cities)</i>		Chicago	854	882	89	623
		New York	1087	968	38	872
		Toronto	818	746	43	591
		Vancouver				
<i>time (quarters)</i>	Q1	605	825	14	400	
	Q2	680	952	31	512	
	Q3	812	1023	30	501	
	Q4	927	1038	38	580	
		computer	home	phone	entertainment	
		security				



Why using OLAP Cube?

Advantages:

- Format the unformatted raw data.
- Multiple dimension cube for one dataset.

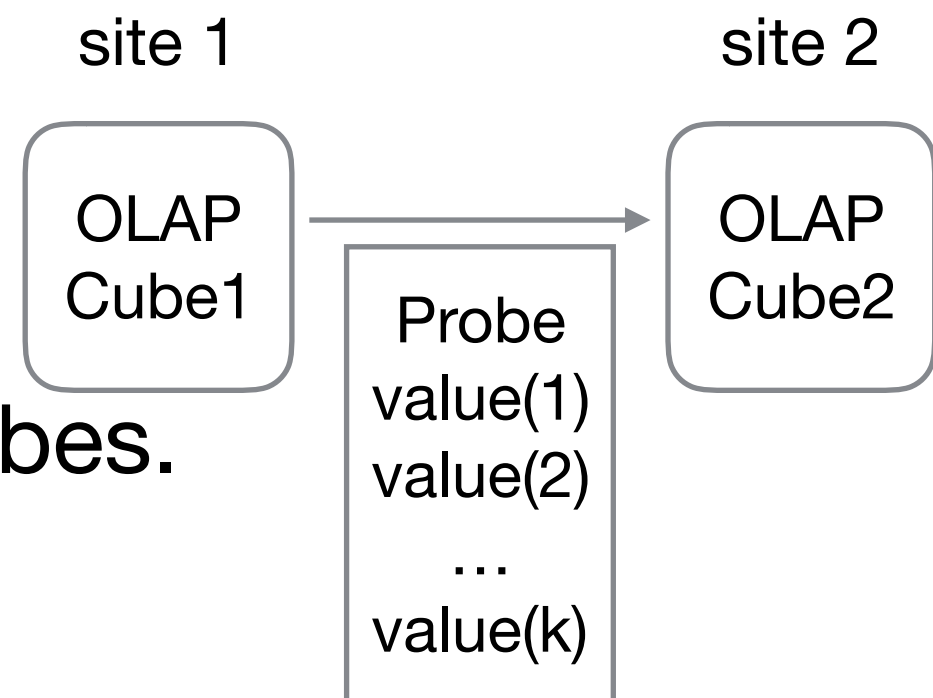
Similarity Search & Check

Search:

- OLAP instruction(eg. dice, slice, roll up) to retrieve the needed attributes.

Check:

- Cross site check: simple probes.
- The probe components.



Reduce tasks scheduling

t : data transfer time.

I : input data size.

R: data reduction ratio.

U: uplink bandwidth.

D: downlink bandwidth.

r : reduce tasks percentage.

$$\min t \quad (1)$$

Goal : Minimize the data transfer time

$$\text{s.t.} \quad (1 - r_i)I_i R_i / U_i \leq t, \forall i, \quad (2)$$

Time to upload data from site i.

$$r_i(\sum_{j=1}^N I_j R_j - I_i R_i) / D_i \leq t, \forall i, \quad (3)$$

Time to download data from other sites.

$$\sum_i r_i = 1, r_i \geq 0, \forall i. \quad (4)$$

System Implementation

Based on Spark v2.1.0 and Iridium.

Utilize Apache Kylin OLAP data cube to deal with all the OLAP operation.

Use Gurobi Optimizer for optimize the LP function.



Evaluation Methodology

- **Workloads:**
 - TPC-DS.
 - Amplab big data benchmark.
- **Hardware platform:**

A real EC2 deployment with 10 sites(Singapore,Tokyo...).

Evaluation Methodology

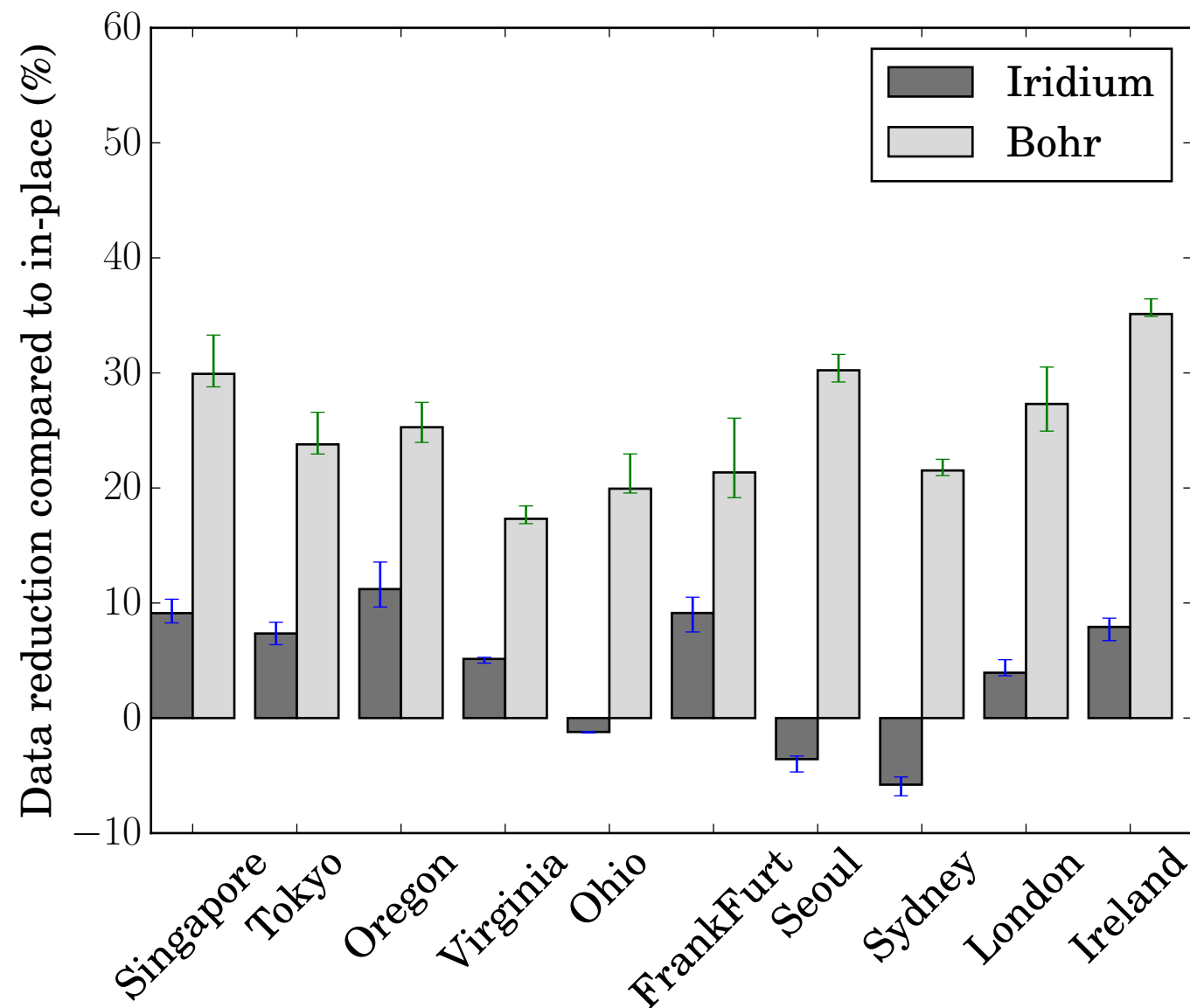
- **Baseline:**

Iridium(Pu et al. Sigcomm' 15).

- **Performance metrics:**

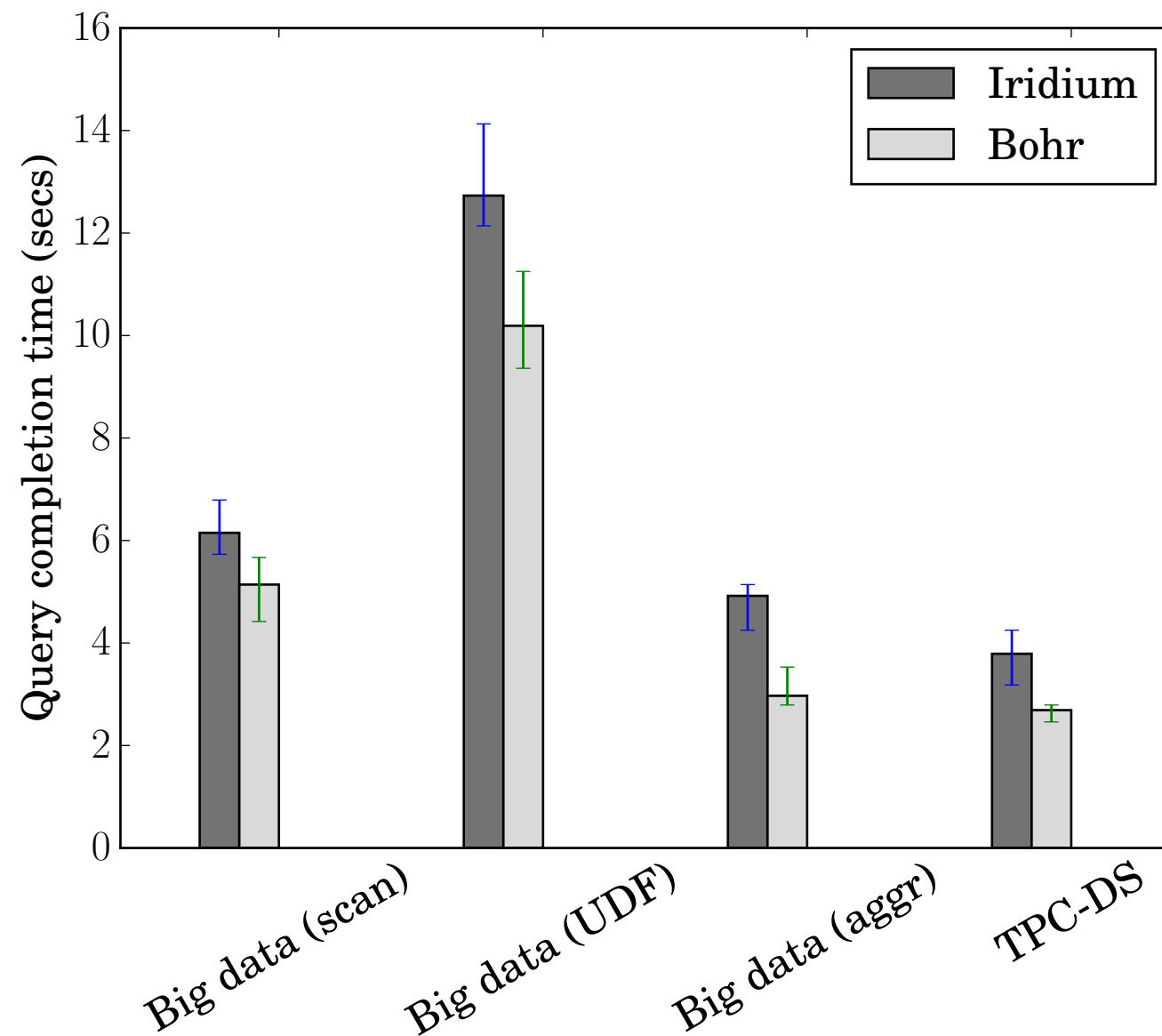
- Query completion time.
- Data reduction ratio.

Data Reduction Ratio



Bohr can increase the **data reduction ratio** in all sites.

Query Completion Time



Bohr saves **30%** QCT compare to Iridium.

Conclusion

We propose Bohr:

- **Similarity aware** data transfer.
- **Predicting and measuring** data reduction ratio for each site.

Bohr is **30%** faster than Iridium.

Thank you!

Future work

1. Generalize the basic idea to more workloads.
2. Dealing with the overhead bring by operate OLAP cubes.
3. Break recurring batch workload condition
4. Multiple types of queries accessing one dataset.