



# SEARCH FOR JOINABLE TABLES IN DATA LAKES

Erkang (Eric) Zhu, PhD Candidate  
University of Toronto

# ABOUT ME



# ABOUT ME

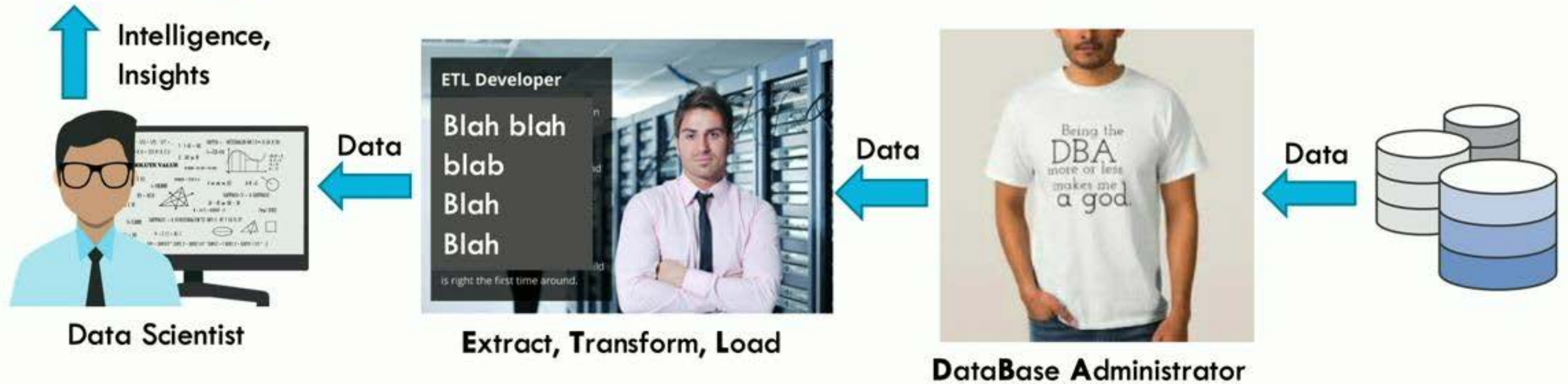


# ABOUT ME



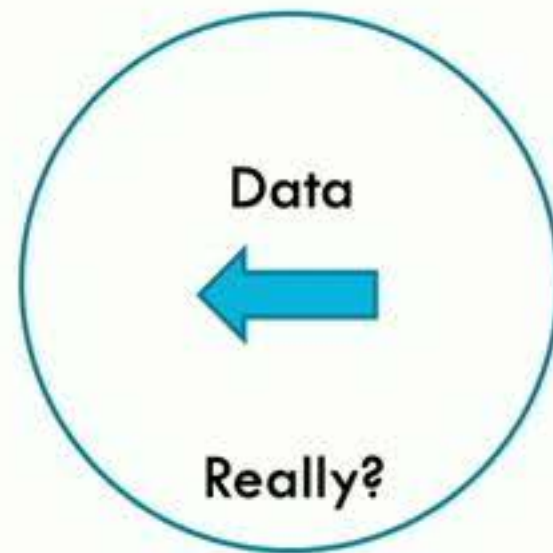
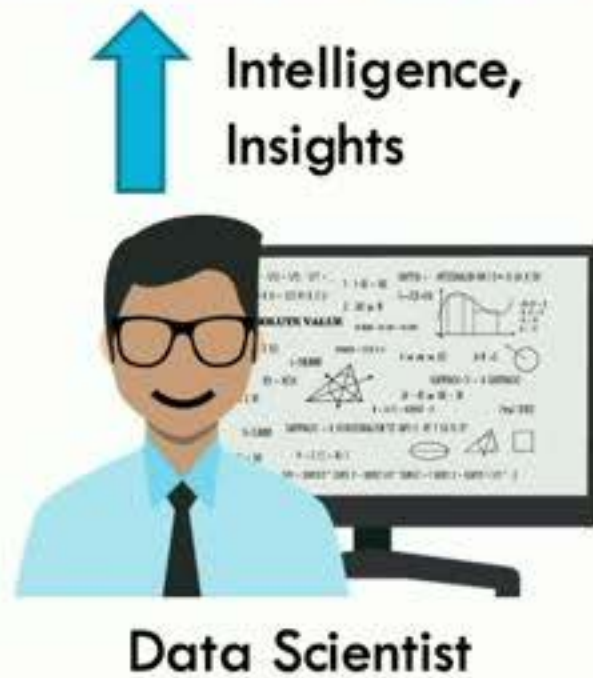


## Traditional Enterprise Analytics — Heavy IT

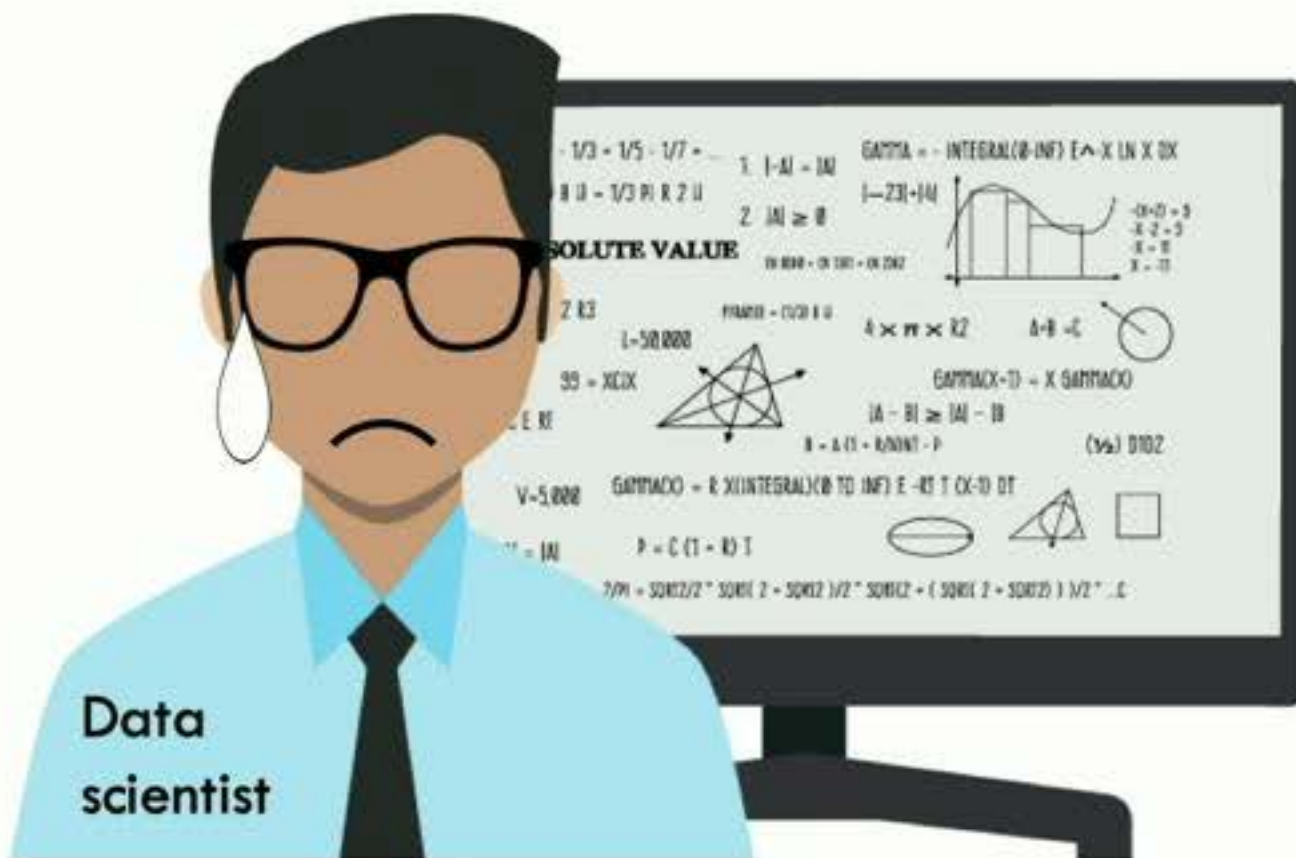




## Data Lake Solution — Minimal IT







Data  
scientist

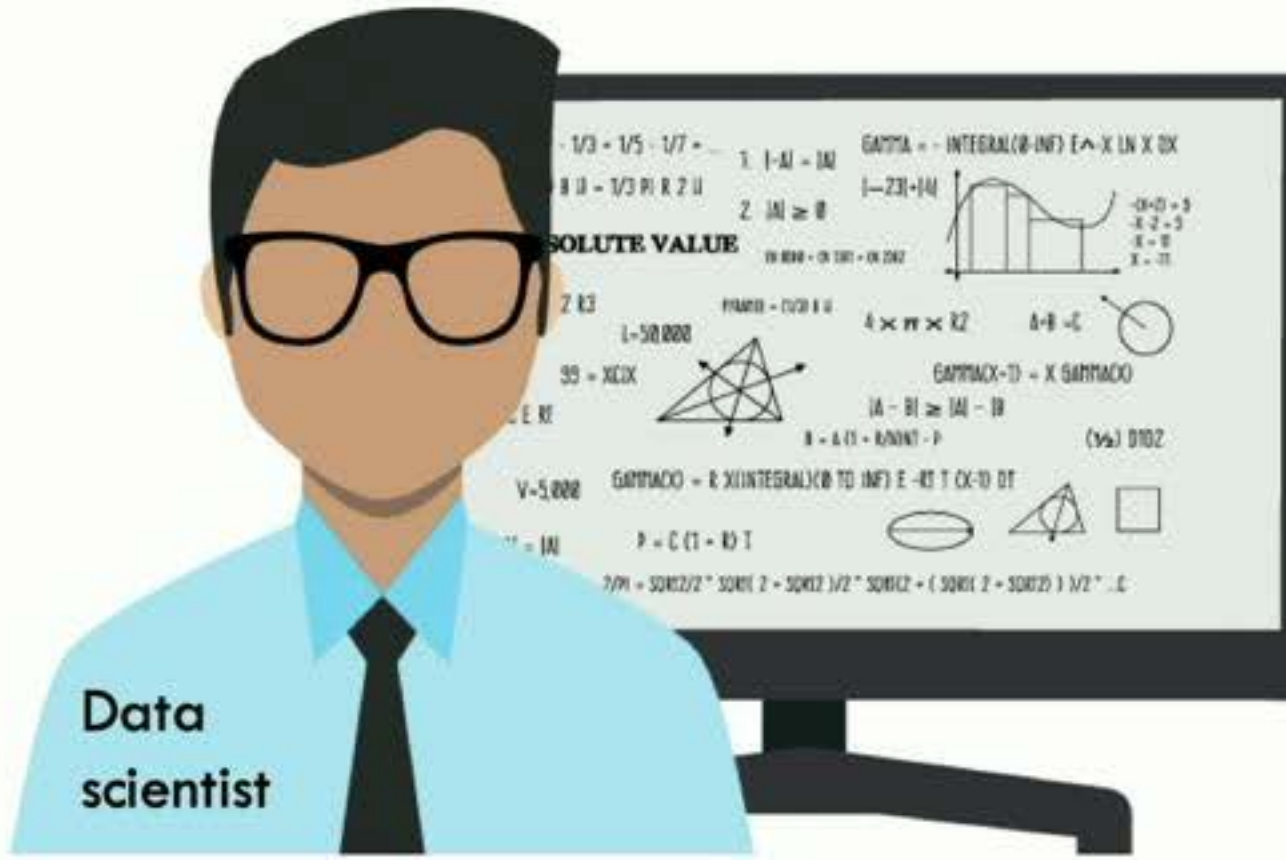
| Company       | Headquarter   | Ticker         |
|---------------|---------------|----------------|
| Alphabet Inc. | Mountain View | GOOGL (NASDAQ) |
| Apple Inc.    | Palo Alto     | AAPL (NASDAQ)  |
| Alibaba       | Hangzhou      | BABA (NYSE)    |
| ...           | ...           | ...            |

Question: More information about these companies?



Query: Tables that can be joined with my own table?





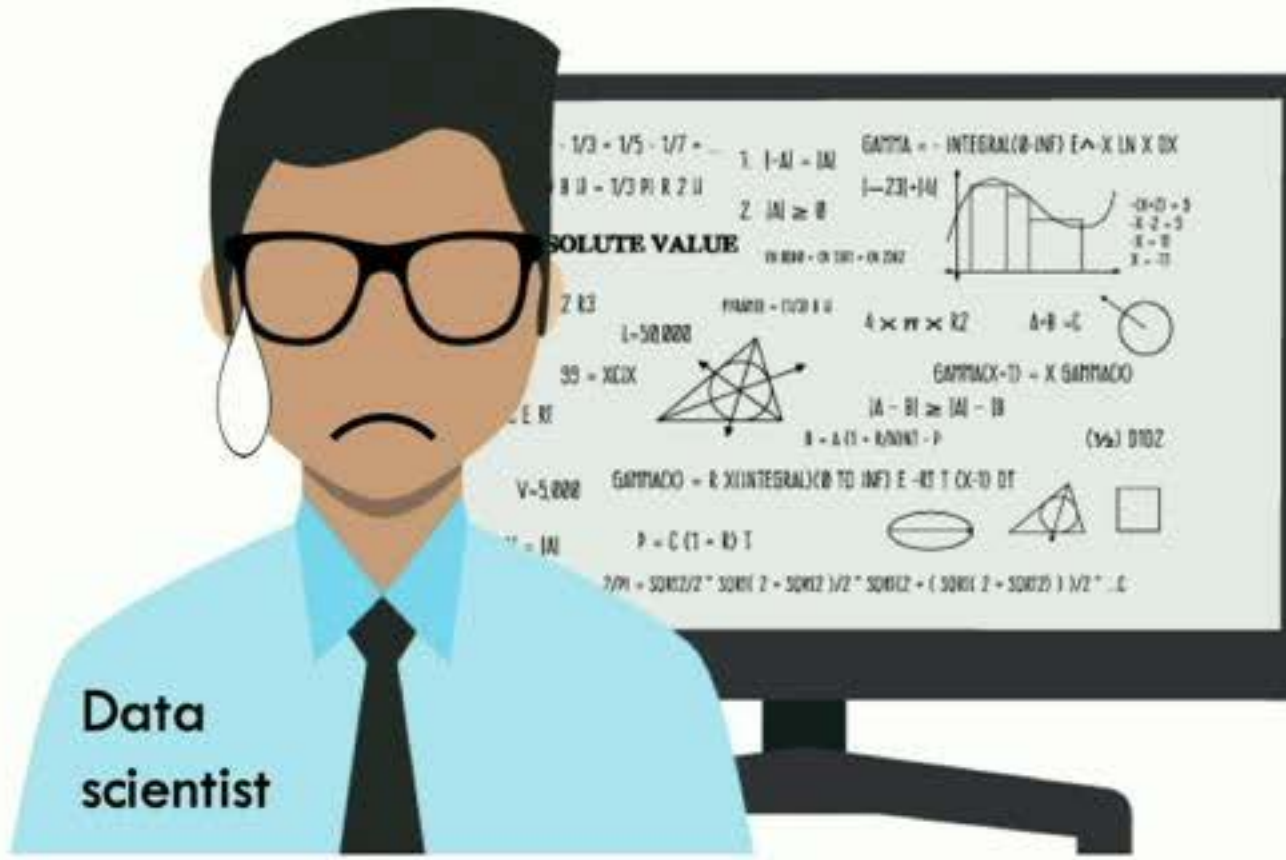
Task: write *ad hoc* data transformation script.



15 min

| Company       | Headquarter   | Ticker         |
|---------------|---------------|----------------|
| Alphabet Inc. | Mountain View | GOOGL (NASDAQ) |
| Apple Inc.    | Palo Alto     | AAPL (NASDAQ)  |
| Alibaba       | Hangzhou      | BABA (NYSE)    |
| ...           | ...           | ...            |

| Exchange | Ticker | Open     | Close    | Market Cap |
|----------|--------|----------|----------|------------|
| NASDAQ   | GOOGL  | 1,032.50 | 1,042.32 | 766 B      |
| NASDAQ   | AAPL   | 171.10   | 175.12   | 892 B      |
| NYSE     | BABA   | 179.37   | 174.13   | 427 B      |
| ...      | ...    | ...      |          |            |



Task: write *ad hoc* data transformation script.



15 min **X 100**

| Company       | Headquarter   | Ticker         |
|---------------|---------------|----------------|
| Alphabet Inc. | Mountain View | GOOGL (NASDAQ) |
| Apple Inc.    | Palo Alto     | AAPL (NASDAQ)  |
| Alibaba       | Hangzhou      | BABA (NYSE)    |
| ...           | ...           | ...            |

| Exchange | Ticker | Open     | Close    | Market Cap |
|----------|--------|----------|----------|------------|
| NASDAQ   | GOOGL  | 1,032.50 | 1,042.32 | 766 B      |
| NASDAQ   | AAPL   | 171.10   | 175.12   | 892 B      |
| NYSE     | BABA   | 179.37   | 174.13   | 427 B      |
| ...      | ...    | ...      |          |            |



“Jack of all trades”

↑ Intelligence,  
Insights



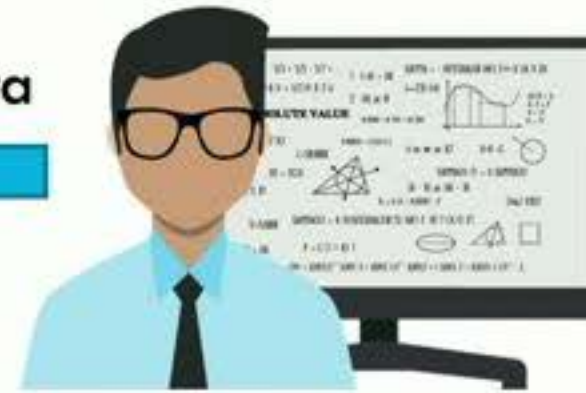
Data Analytics

Data



Transformation Scripts

Data



Locate Raw Files

Data



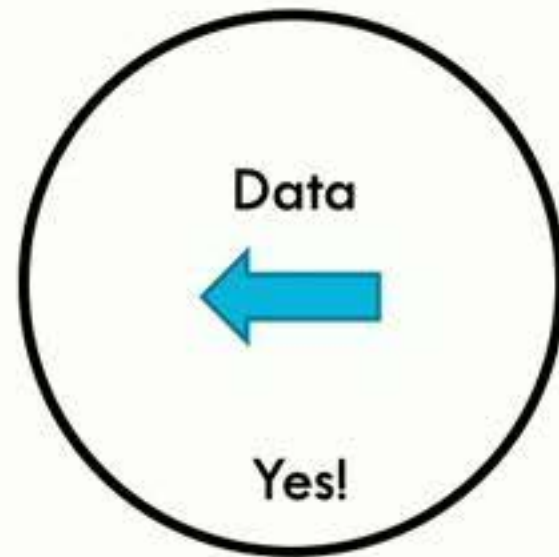
Data Lake

# CONTRIBUTIONS

Algorithmic solutions that enable data scientists to **search** for joinable tables and automatically **join** them – making data scientists more powerful.



Data Scientist



Data Lake

# JOINABLE TABLE SEARCH

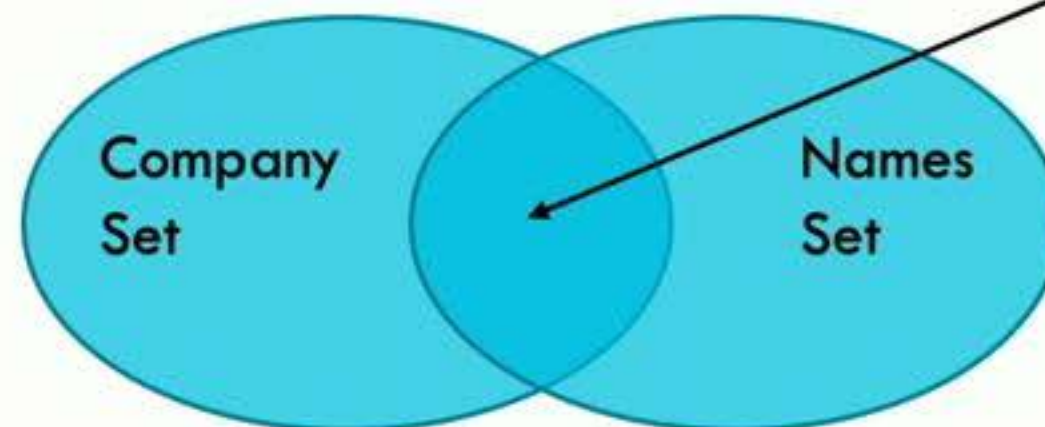
**Query Table**

| Company               | Headquarter   | CEO           |
|-----------------------|---------------|---------------|
| Google Inc.           | Mountain View | Sundar Pichai |
| Databricks            | San Francisco | Ali Ghodsi    |
| Microsoft Corporation | Redmond       | Satya Nadella |
| ...                   | ...           | ...           |

**Query Column**

**Data Lake**

| Names                 | CIK        | Q4 Revenue     |
|-----------------------|------------|----------------|
| Google Inc.           | 0001288776 | 26.06 Billions |
| Microsoft Corporation | 0000789019 | 23.3 Billions  |
| ...                   | ...        | ...            |



Fraction of distinct values in the query that are joinable

# WHY SEARCHING FOR JOINABLE TABLES?



[1] <https://open.canada.ca/data/en/dataset/be073ee2-a302-4d32-af20-a48f5fbe2e63>

# WHY SEARCHING FOR JOINABLE TABLES?



Discovery new/unexpected correlations

[1] <https://open.canada.ca/data/en/dataset/be073ee2-a302-4d32-af20-a48f5fbe2e63>

[2] <https://open.canada.ca/data/en/dataset/07dcaaf5-6a58-4ce3-8876-e30feefef8dd>

# WHY SEARCHING FOR JOINABLE TABLES?

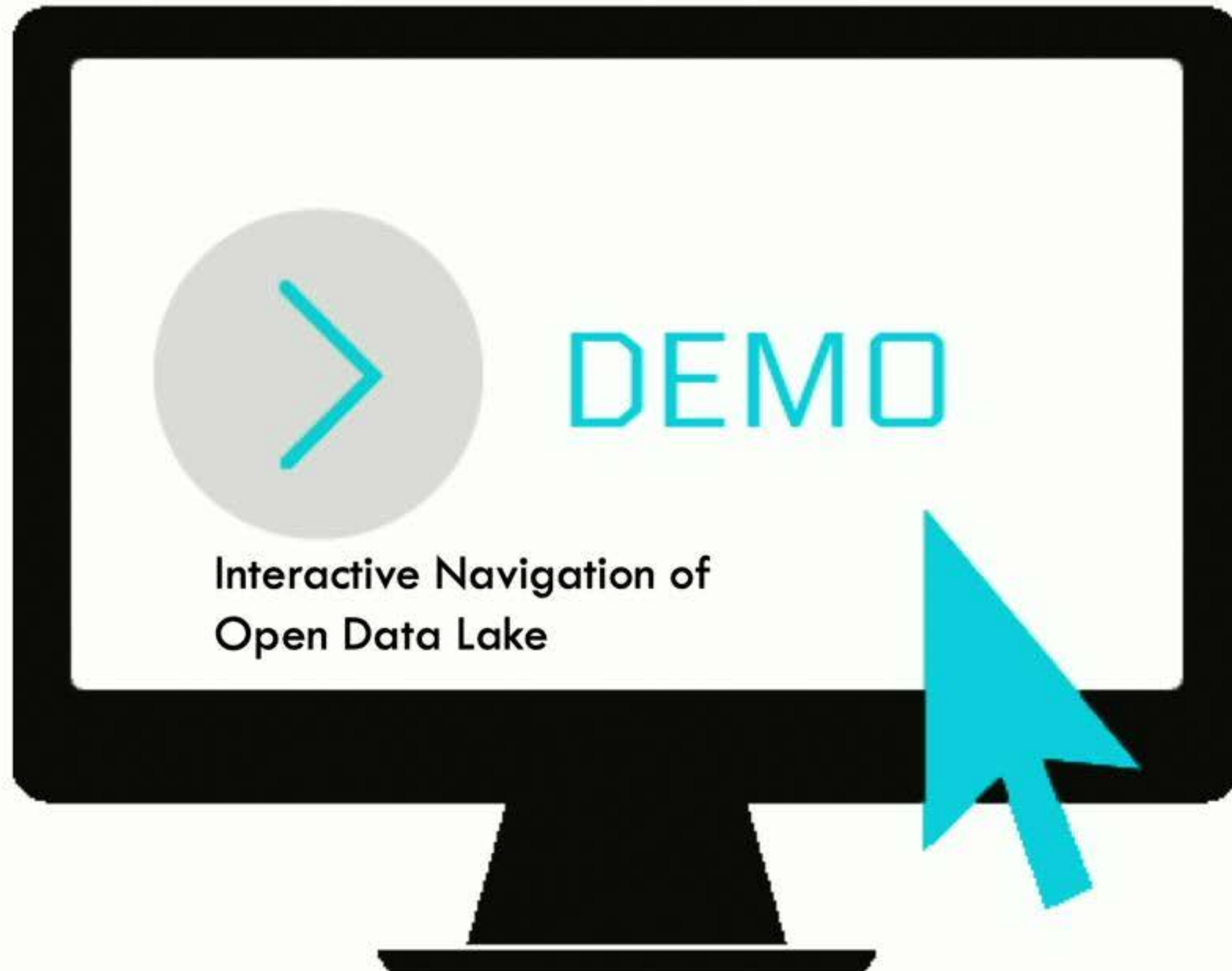
| NO <sub>x</sub> Emission | Facility         | Postal | ... |
|--------------------------|------------------|--------|-----|
| 9430.81                  | Oil Sands        | T9H3E3 |     |
| 50.7                     | Carberry Factory | R0K0H0 |     |
| 60.6                     | Grand Falls      | E3Y4A5 |     |
| 99                       | Ingleside        | K0C1M0 |     |

[1] <https://open.canada.ca/data/en/dataset/f6660f68-5e78-47b3-9306-9805e1d3d6e9>

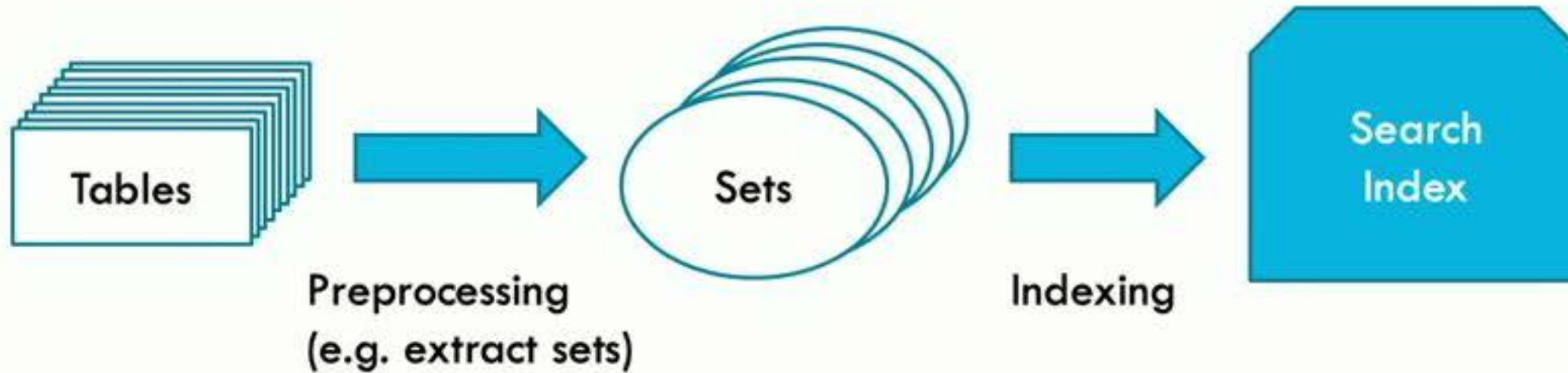
| Contributor | Postal Code | Recipient Party | Amount | ... |
|-------------|-------------|-----------------|--------|-----|
| E*,B*       | T9H3E3      | Liberal         | 397.46 |     |
| J*,S*       | K0K0H0      | Conservative    | 500.00 |     |
| J*,S*       | K0K0H0      | Conservative    | 250.00 |     |
| M*,S*       | K0C1M0      | Liberal         | 400.00 |     |

[2] <https://open.canada.ca/data/en/dataset/ef1e3528-b570-4a42-92ef-18a9749af8f2>

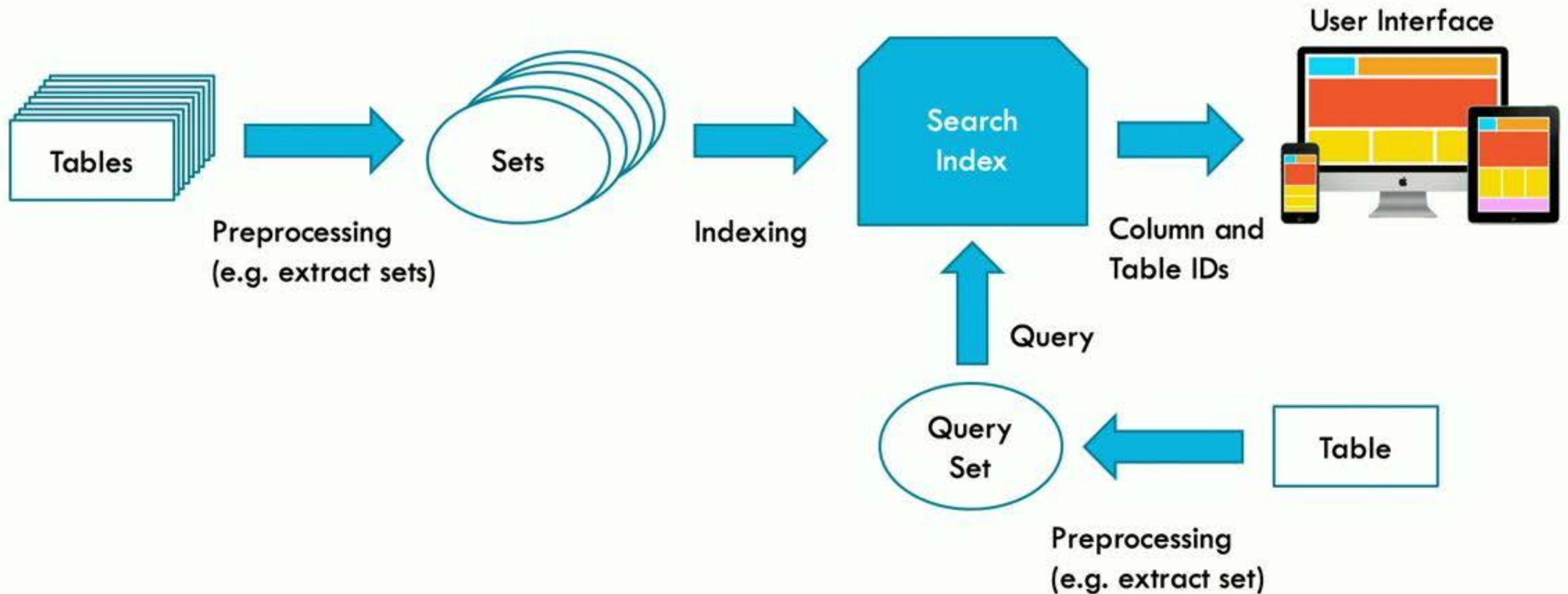
Discover new features for modeling, prediction, or insights



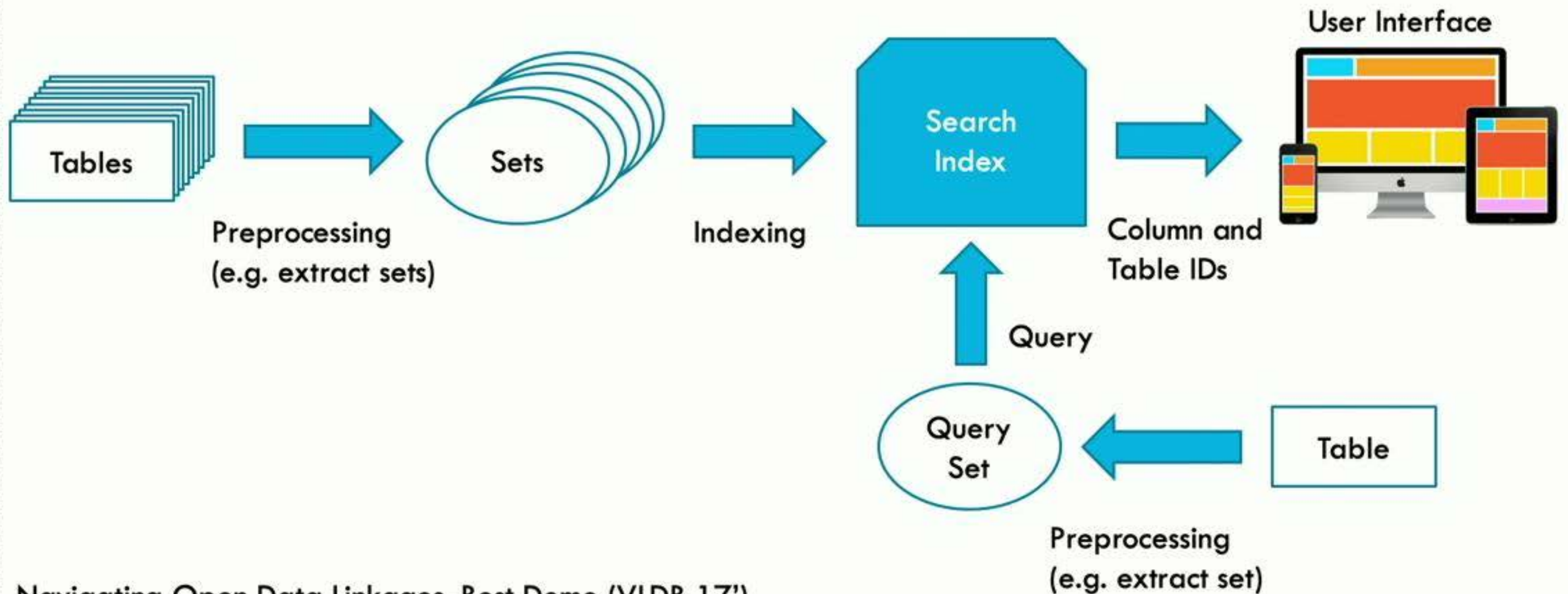
# A JOINABLE TABLE SEARCH SYSTEM



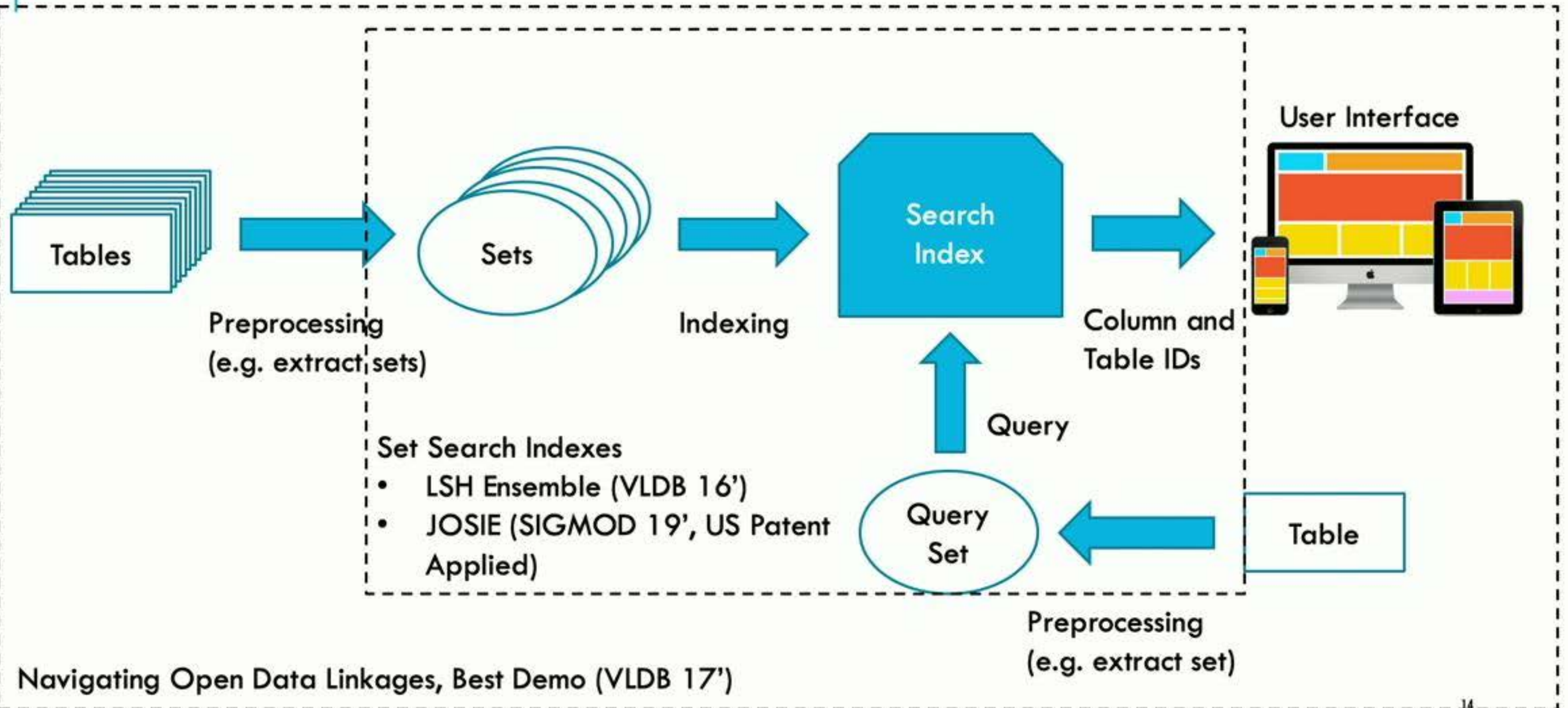
# A JOINABLE TABLE SEARCH SYSTEM



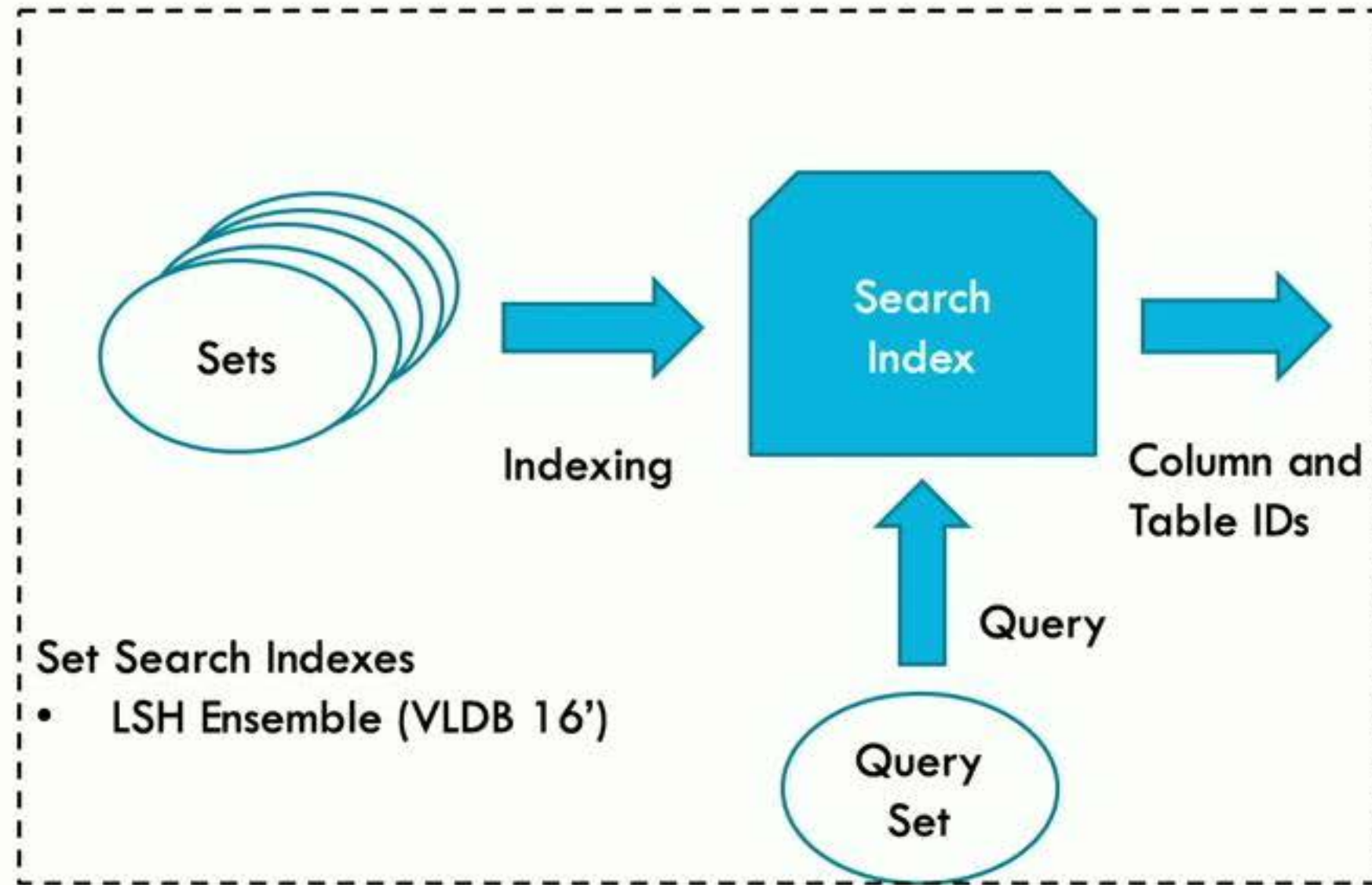
# A JOINABLE TABLE SEARCH SYSTEM



# A JOINABLE TABLE SEARCH SYSTEM



# A JOINABLE TABLE SEARCH SYSTEM



# DATA SKETCHES

## Challenges:

- A set from a table column can have tens of thousands of values
- A data lake can have hundreds of thousands of tables
- An Open Data lake: 200K tables, and Max set size: 22M

## Data sketches for sets:

- Small, fixed-size summaries (a few KBs)
- Probabilistic guarantee (more on this later)
- $O(|X|)$  time to build,  $|X|$  is set size

We use MinHash<sup>1</sup> data sketches for joinable table search

[1] A. Broder. On the resemblance and containment of documents. In Compression and Complexity of Sequences, pages 21–28, 1997.

# MINHASH

Random, independent  
hash functions (e.g.,  
MurmurHash3)

$h_1$  {"ATL", "LAX", "ORD", "DFW", "DEN", "JFK", "SFO", "LAS", "SEA", "CLT"}

| Hash Function | Min Hash Value |
|---------------|----------------|
| $h_1$         |                |
| $h_2$         |                |
| $h_3$         |                |
| $h_4$         |                |
| $h_5$         |                |
| $h_6$         |                |
| $h_7$         |                |
| $h_8$         |                |

# MINHASH

Random, independent  
hash functions (e.g.,  
MurmurHash3)

{“ATL”, “LAX”, “ORD”, “DFW”, “DEN”, “JFK”, “SFO”, “LAS”, “SEA”, “CHT”}

0x123, 0x3f1, 0xa1c, 0x312, 0x456, 0x42a, 0xb12, 0x98a, 0x547, 0x132

| Hash Function | Min Hash Value |
|---------------|----------------|
| $h_1$         | 0x123          |
| $h_2$         |                |
| $h_3$         |                |
| $h_4$         |                |
| $h_5$         |                |
| $h_6$         |                |
| $h_7$         |                |
| $h_8$         |                |

# MINHASH

Random, independent  
hash functions (e.g.,  
MurmurHash3)

{“ATL”, “LAX”, “ORD”, “DFW”, “DEN”, “JFK”, “SFO”, “LAS”, “SEA”, “CHT”}

0x25b, 0xa21, 0x81a, 0x514, 0xb51, 0x431, 0x14a, 0x981, 0x1c7, 0x16b

| Hash Function | Min Hash Value |
|---------------|----------------|
| $h_1$         | 0x123          |
| $h_2$         | 0x14a          |
| $h_3$         |                |
| $h_4$         |                |
| $h_5$         |                |
| $h_6$         |                |
| $h_7$         |                |
| $h_8$         |                |

# MINHASH

Random, independent  
hash functions (e.g.,  
MurmurHash3)

{“ATL”, “LAX”, “ORD”, “DFW”, “DEN”, “JFK”, “SFO”, “LAS”, “SEA”, “CLT”}

| Hash Function | Min Hash Value |
|---------------|----------------|
| $h_1$         | 0x123          |
| $h_2$         | 0x14a          |
| $h_3$         | 0x41c          |
| $h_4$         | 0x1bc          |
| $h_5$         | 0x2c3          |
| $h_6$         | 0x5d1          |
| $h_7$         | 0x613          |
| $h_8$         | 0x312          |

[Optional]: speed up using only one hash function and universal hashing

# MINHASH

For two sets  $X$  and  $Y$ , the probability of a minimum hash value collision:

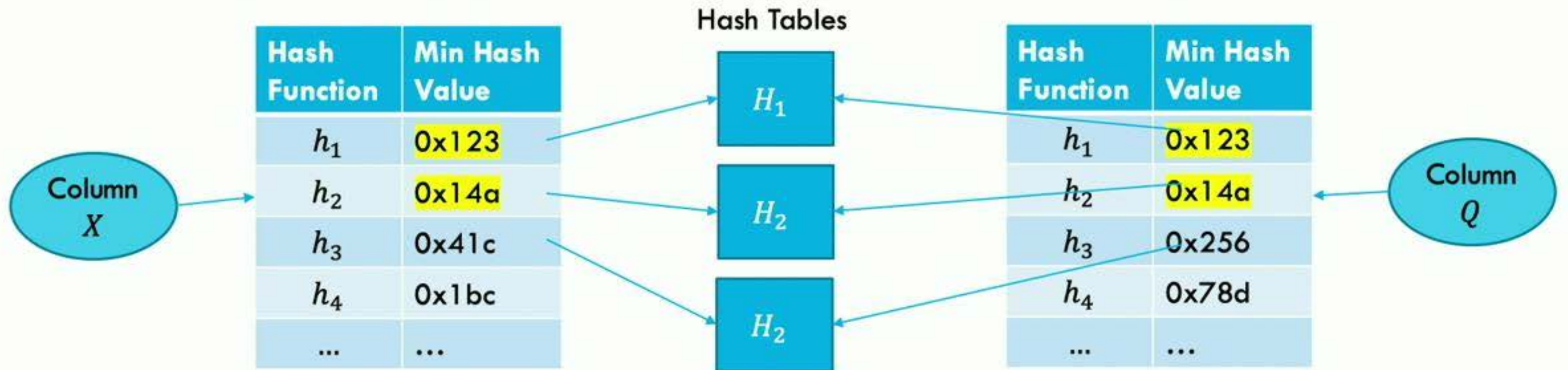
$$P[\min(h_1(X)) = \min(h_1(Y))] = \frac{|X \cap Y|}{|X \cup Y|} \leftarrow \text{This is Jaccard similarity}$$

| $X$ | Hash Function | Min Hash Value |
|-----|---------------|----------------|
|     | $h_1$         | 0x123          |
|     | $h_2$         | 0x14a          |
|     | $h_3$         | 0x41c          |
|     | $h_4$         | 0x1bc          |
|     | ...           | ...            |

| $Y$ | Hash Function | Min Hash Value |
|-----|---------------|----------------|
|     | $h_1$         | 0x123          |
|     | $h_2$         | 0x14a          |
|     | $h_3$         | 0x256          |
|     | $h_4$         | 0x78d          |
|     | ...           | ...            |

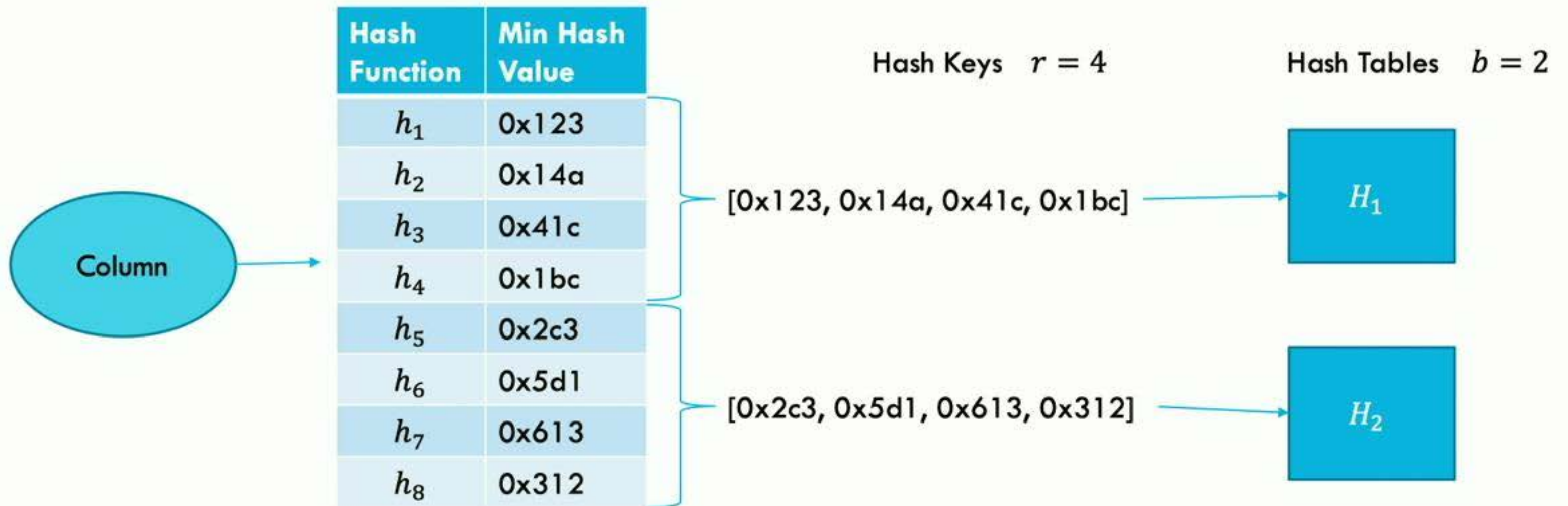
# MINHASH

We can use hash tables to find “approximately” joinable columns under Jaccard similarity

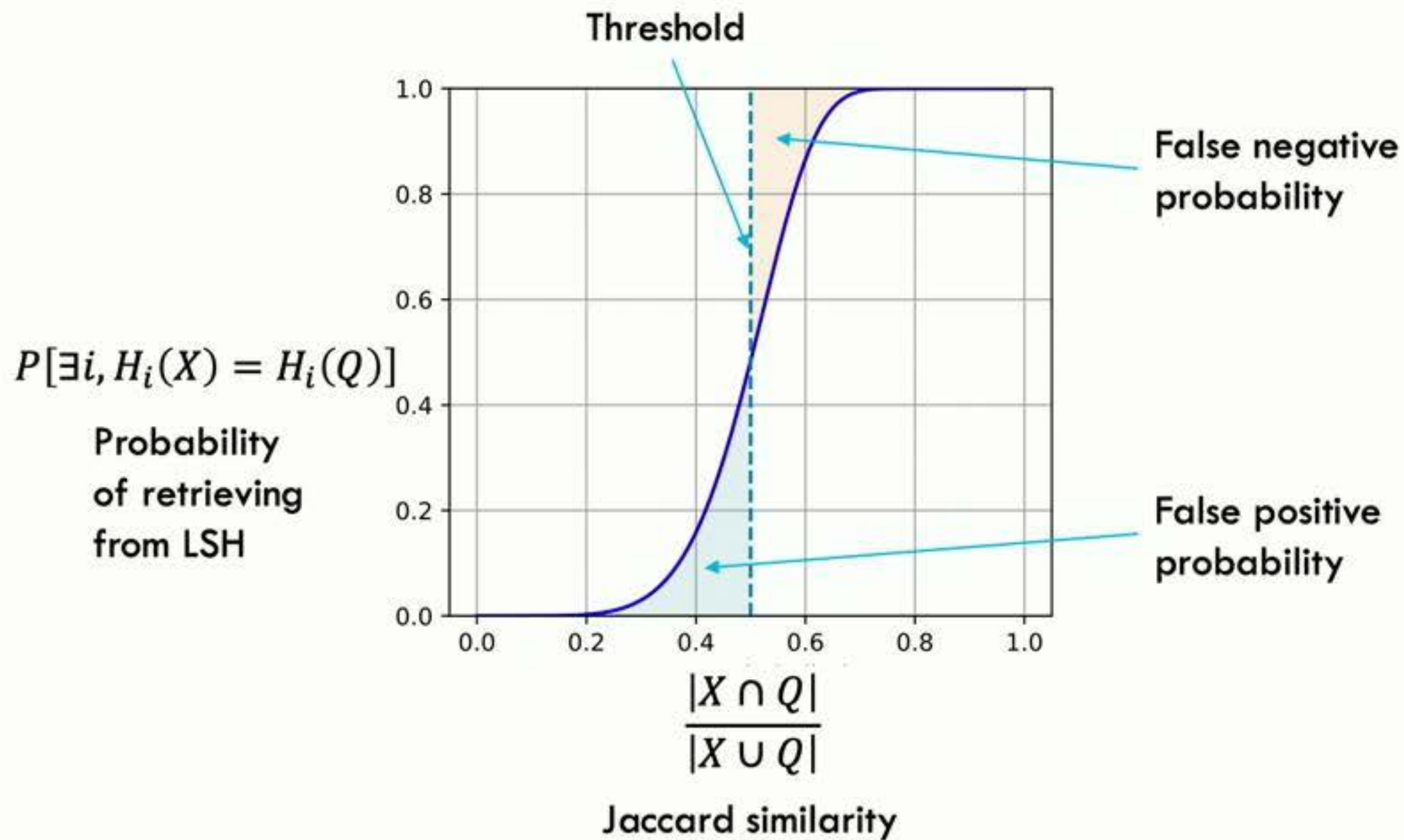


# MINHASH LSH

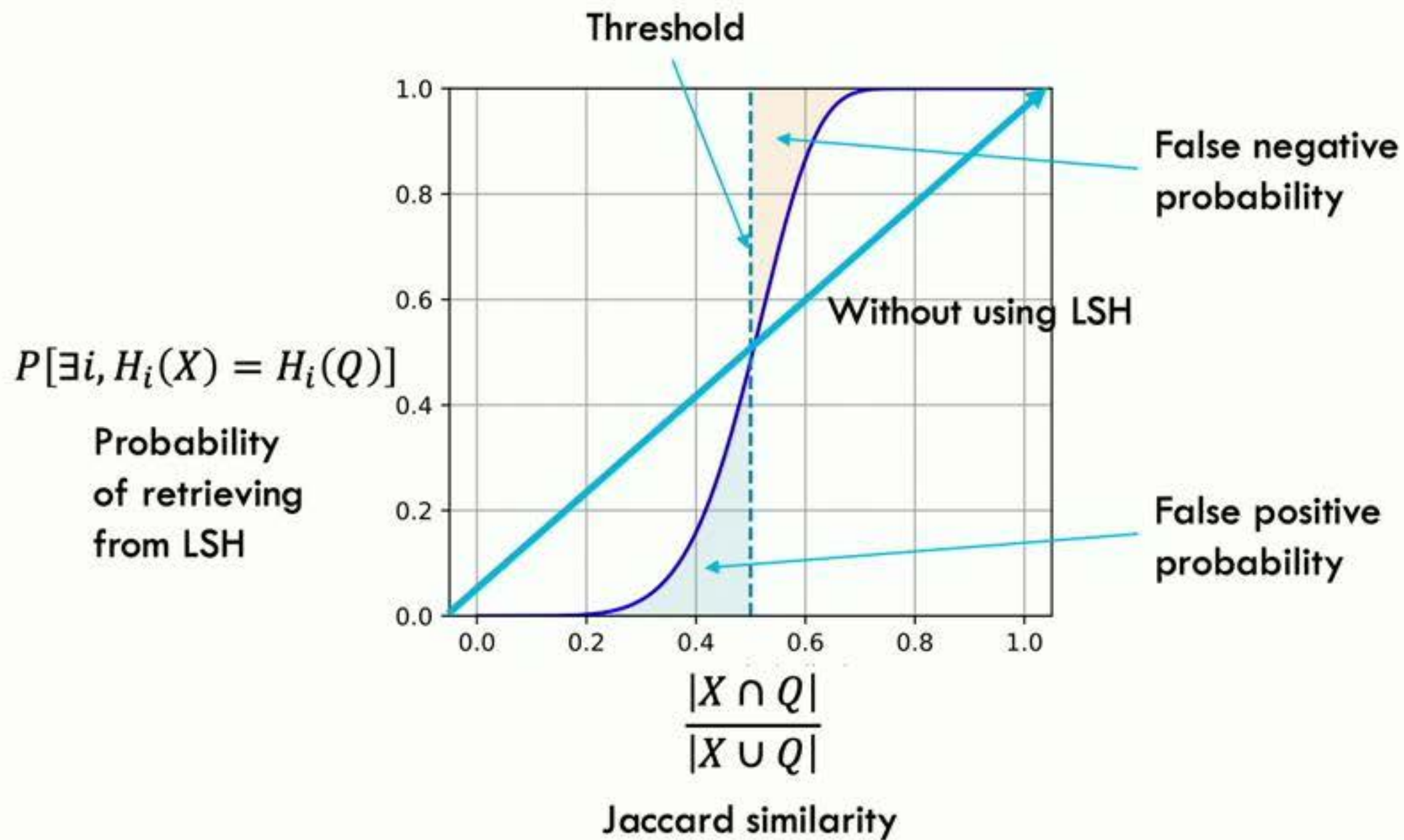
Locality Sensitive Hashing<sup>1</sup> (LSH) index provides better accuracy through “boosting”



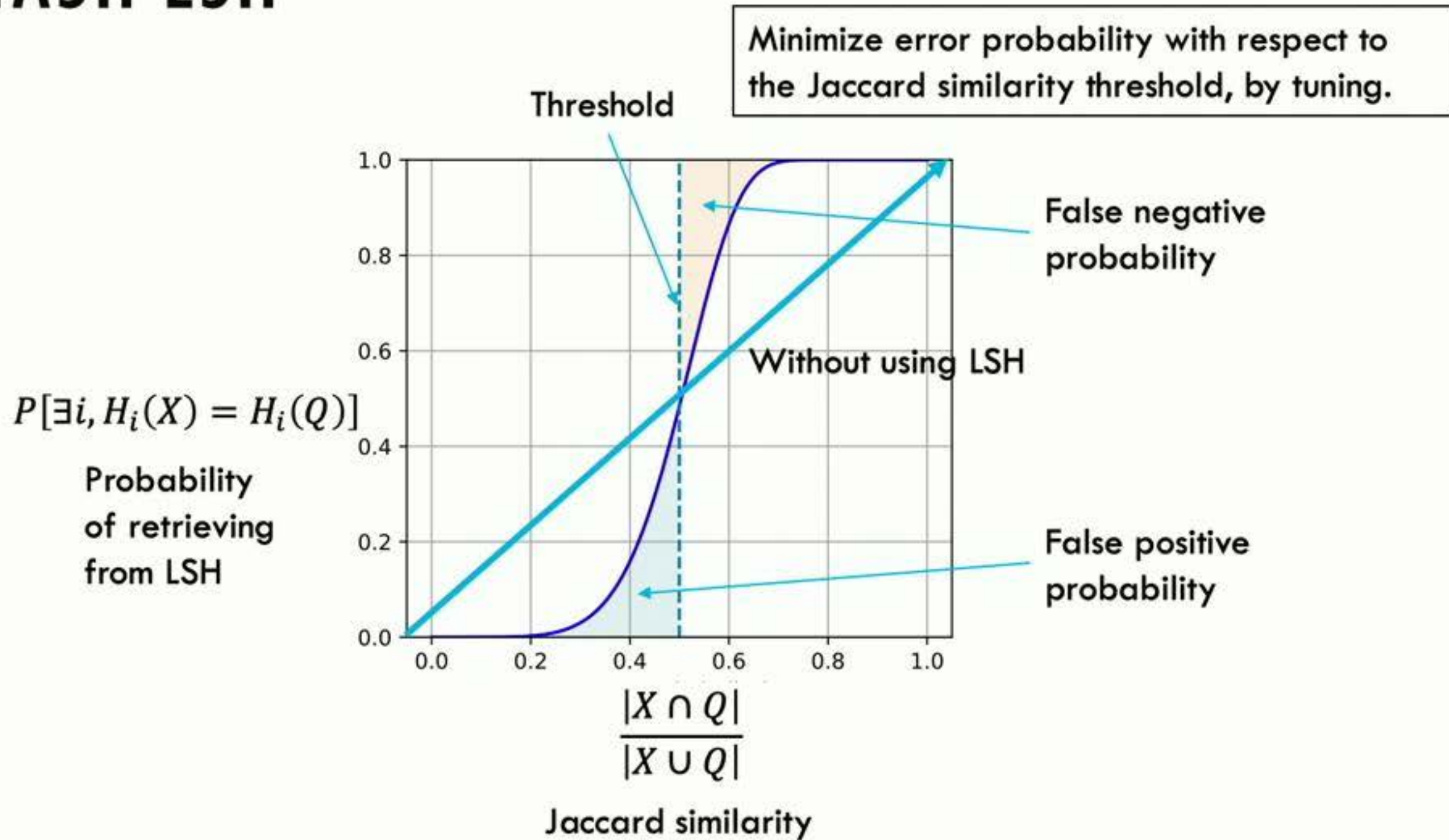
# MINHASH LSH



# MINHASH LSH

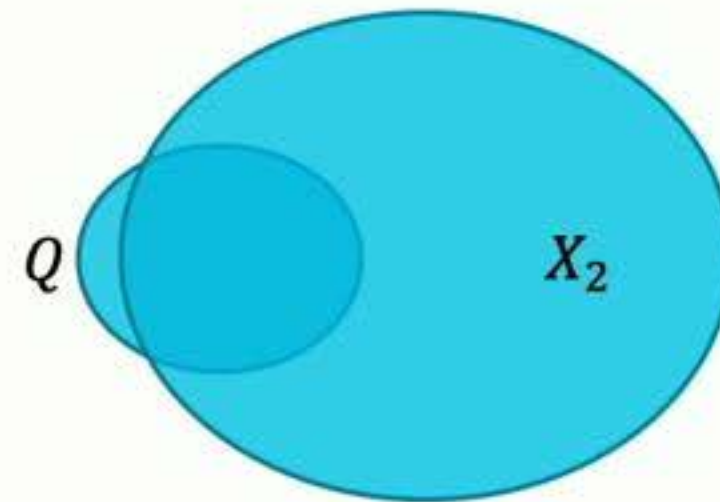
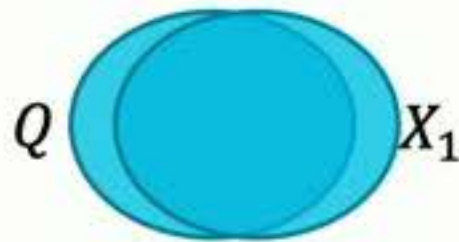


# MINHASH LSH



# IS JACCARD SIMILARITY THE RIGHT MEASURE?

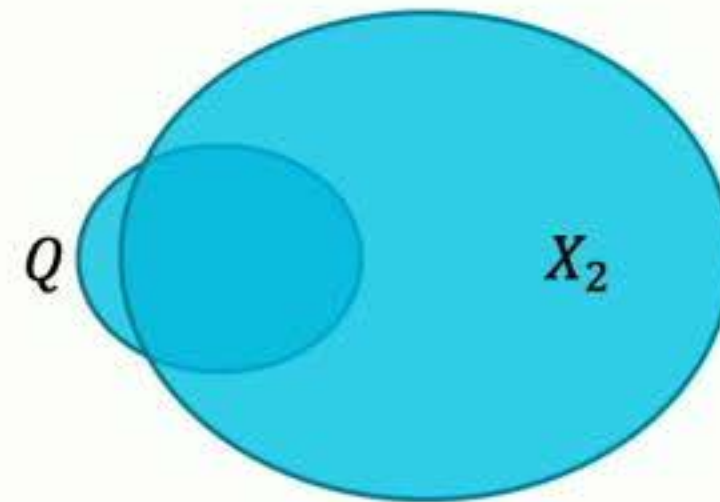
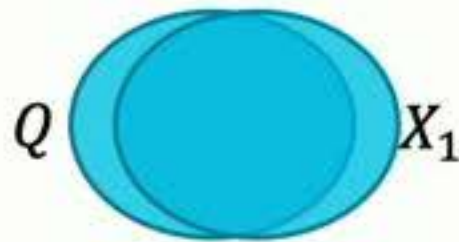
As  $|X_1 \cap Q| = |X_2 \cap Q|$ , both  $X_1$  and  $X_2$  are equally good for joining with  $Q$



# IS JACCARD SIMILARITY THE RIGHT MEASURE?

As  $|X_1 \cap Q| = |X_2 \cap Q|$ , both  $X_1$  and  $X_2$  are equally good for joining with  $Q$

But  $\frac{|X_1 \cap Q|}{|X_1 \cup Q|} > \frac{|X_2 \cap Q|}{|X_2 \cup Q|}$



Alternatively,  $\frac{|X_1 \cap Q|}{|Q|} = \frac{|X_2 \cap Q|}{|Q|}$

# CONTAINMENT SEARCH

We use containment,  $\frac{|X \cap Q|}{|Q|}$ , as a measure for joinable tables

We can “fix” MinHash LSH for containment

# CONTAINMENT SEARCH

We use containment,  $\frac{|X \cap Q|}{|Q|}$ , as a measure for joinable tables

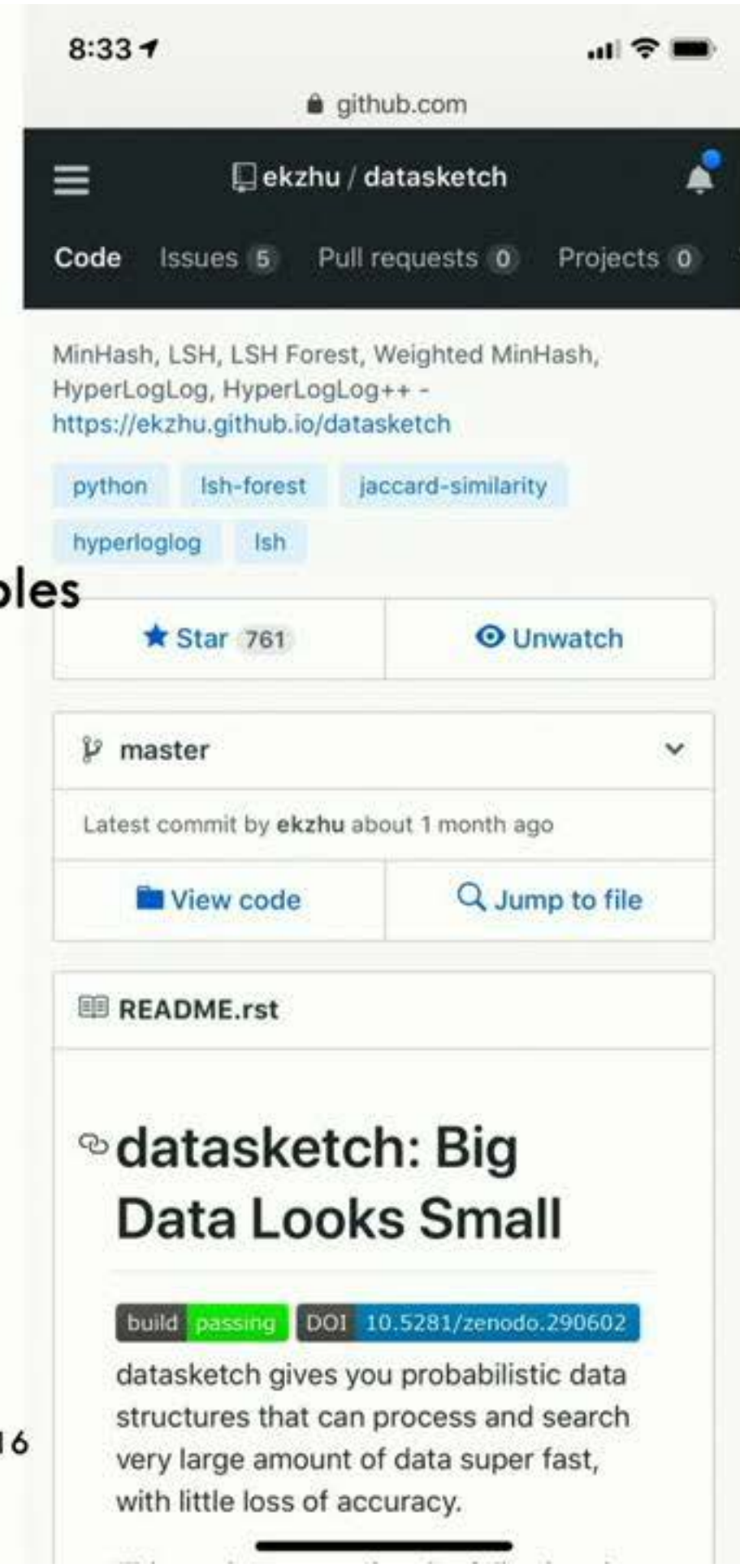
We can “fix” MinHash LSH for containment

## LSH Ensemble<sup>1</sup>

- Part of *datasketch* Python library<sup>2</sup>, together with MinHash LSH
- The library is used by Google (TimeSketch), MIT (Aurum Data Discovery) and Stanford (NLP)
- Over 760 stars on Github

[1] Erkang Zhu, Fatemeh Nargesian, Ken Pu, Renée J. Miller, “LSH Ensemble: L Internet-Scale Domain Search”, VLDB 2016

[2] <https://github.com/ekzhu/datasketch>



# FROM CONTAINMENT TO JACCARD

Apply Inclusion-Exclusion Principle:

$$\underset{\text{Jaccard}}{\frac{|X \cap Q|}{|X \cup Q|}} = \frac{|X \cap Q|}{|X| + |Q| - |X \cap Q|} = \frac{\frac{|X \cap Q|}{|Q|}}{1 + \frac{|X|}{|Q|} - \frac{|X \cap Q|}{|Q|}}$$

Containment

# FROM CONTAINMENT TO JACCARD

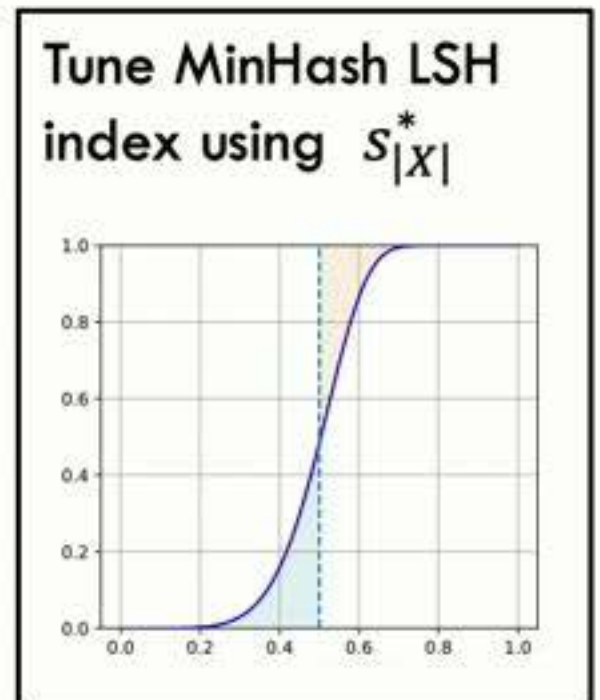
A containment threshold ( $t^*$ ) can be translated into a Jaccard threshold ( $s^*$ )

$$s_{|X|}^* = \frac{t^*}{1 + \frac{|X|}{|Q|} - t^*}$$

# FROM CONTAINMENT TO JACCARD

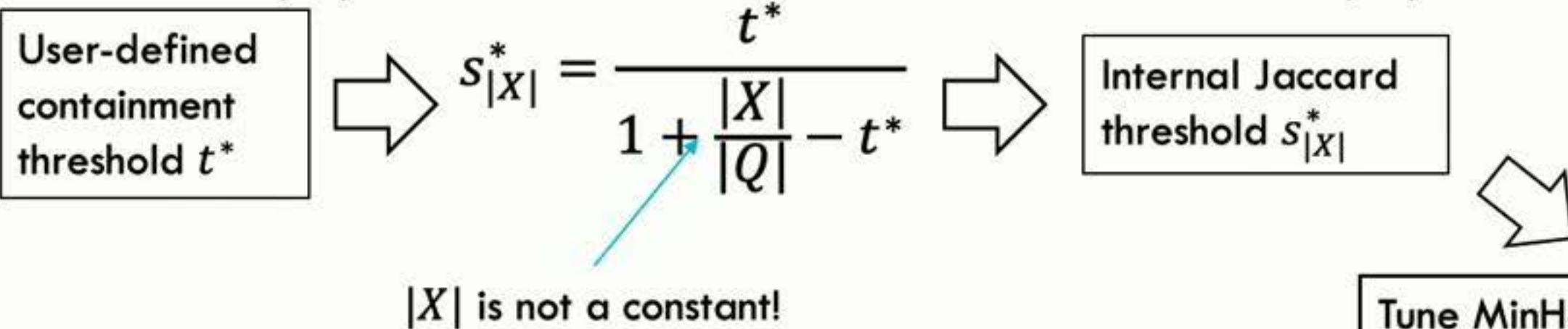
A containment threshold ( $t^*$ ) can be translated into a Jaccard threshold ( $s^*$ )

User-defined containment threshold  $t^*$   $\Rightarrow s_{|X|}^* = \frac{t^*}{1 + \frac{|X|}{|Q|} - t^*}$   $\Rightarrow$  Internal Jaccard threshold  $s_{|X|}^*$



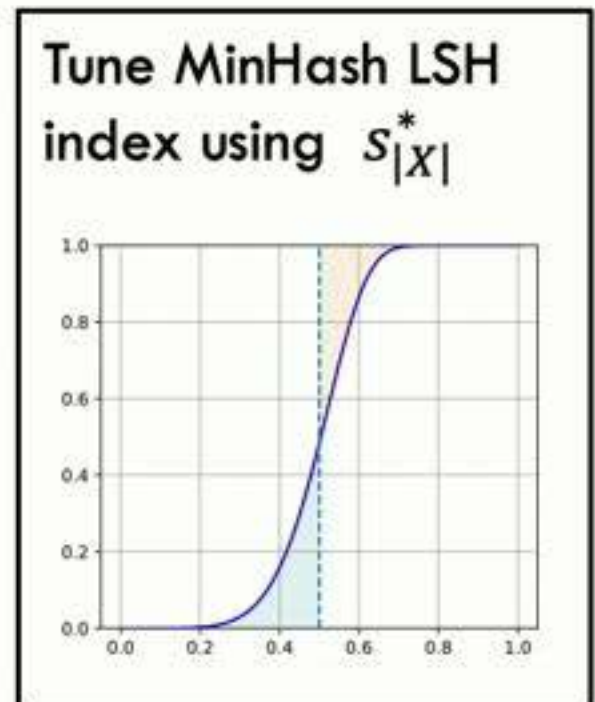
# FROM CONTAINMENT TO JACCARD

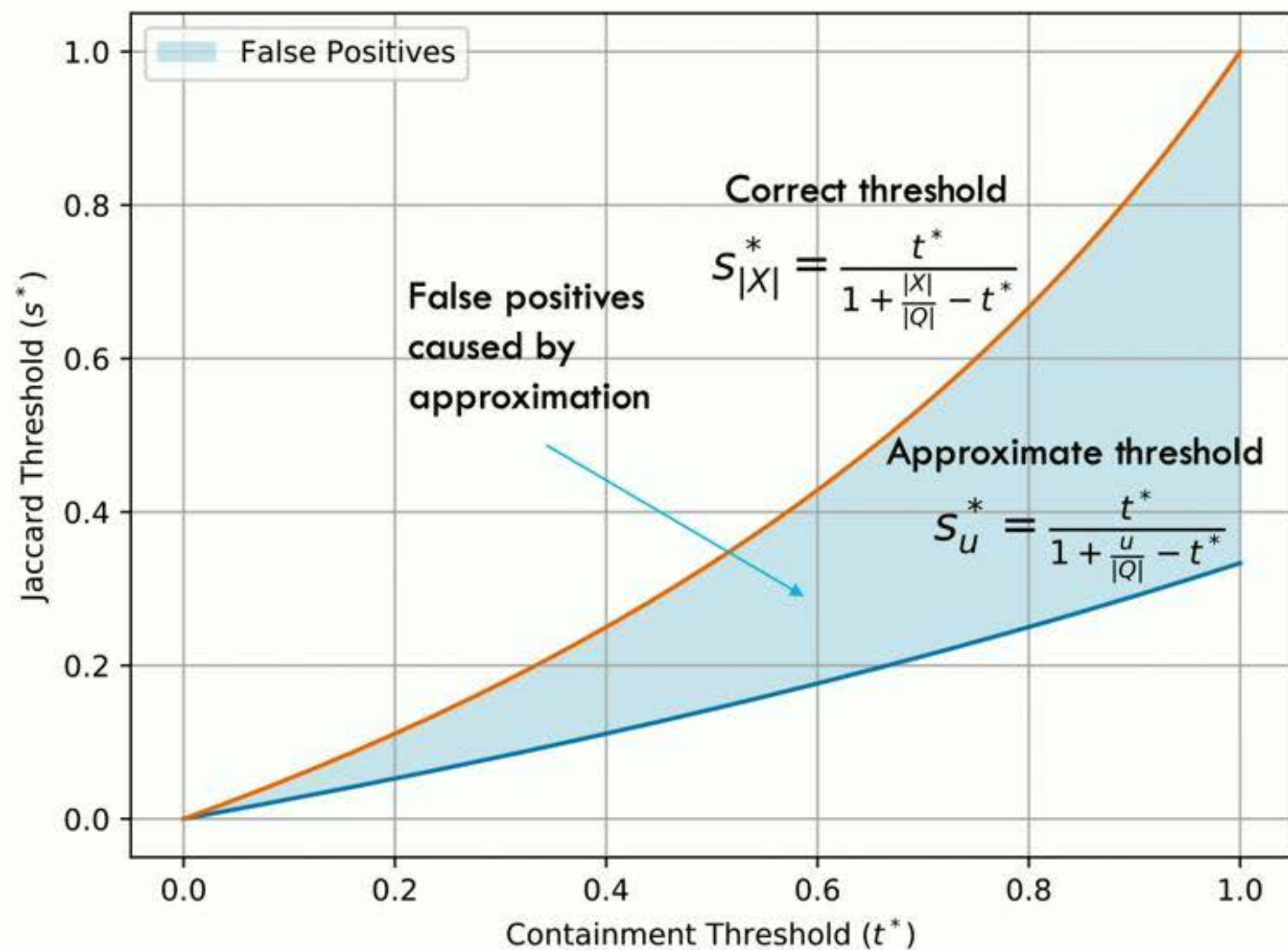
A containment threshold ( $t^*$ ) can be translated into a Jaccard threshold ( $s^*$ )

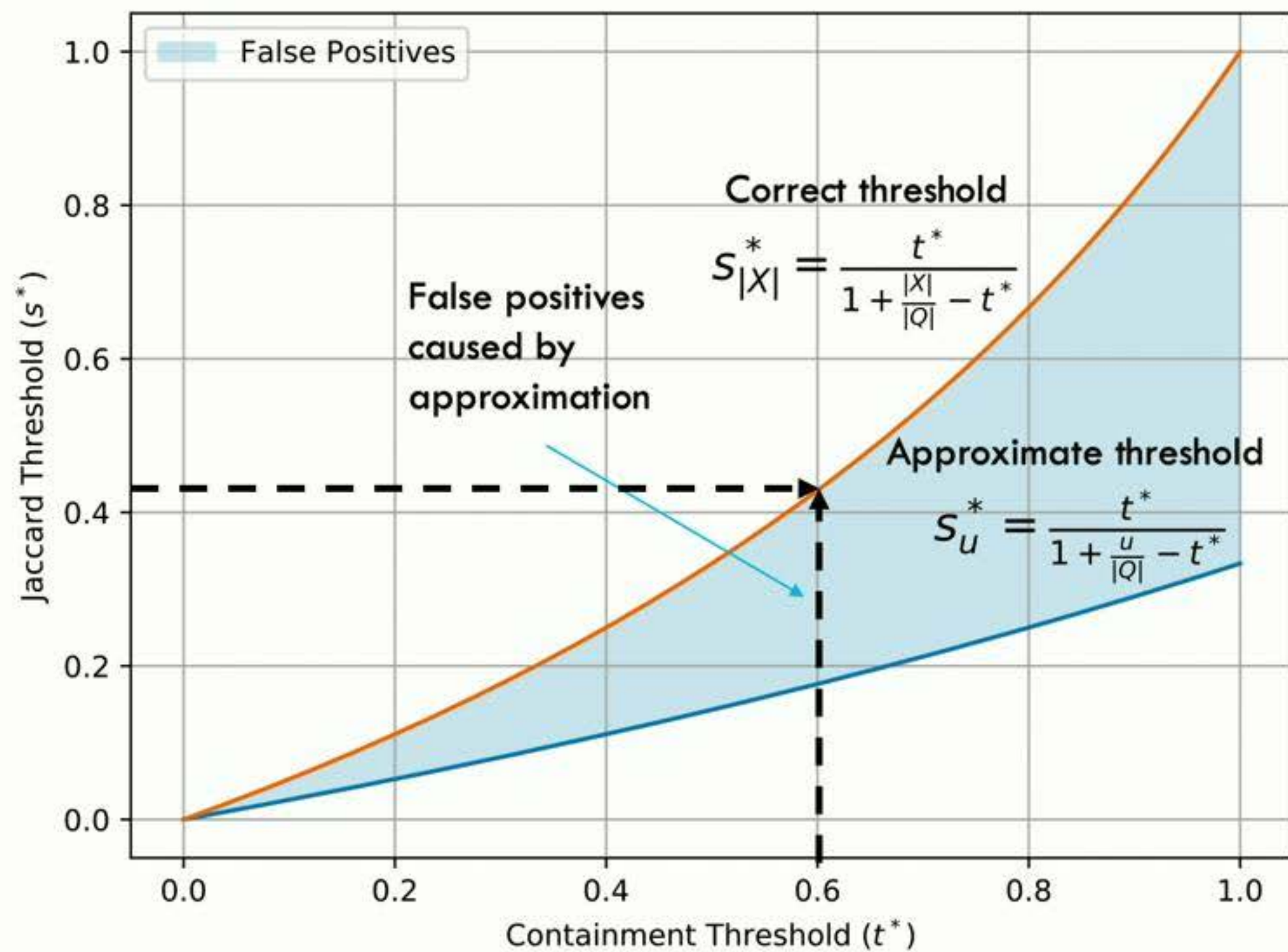


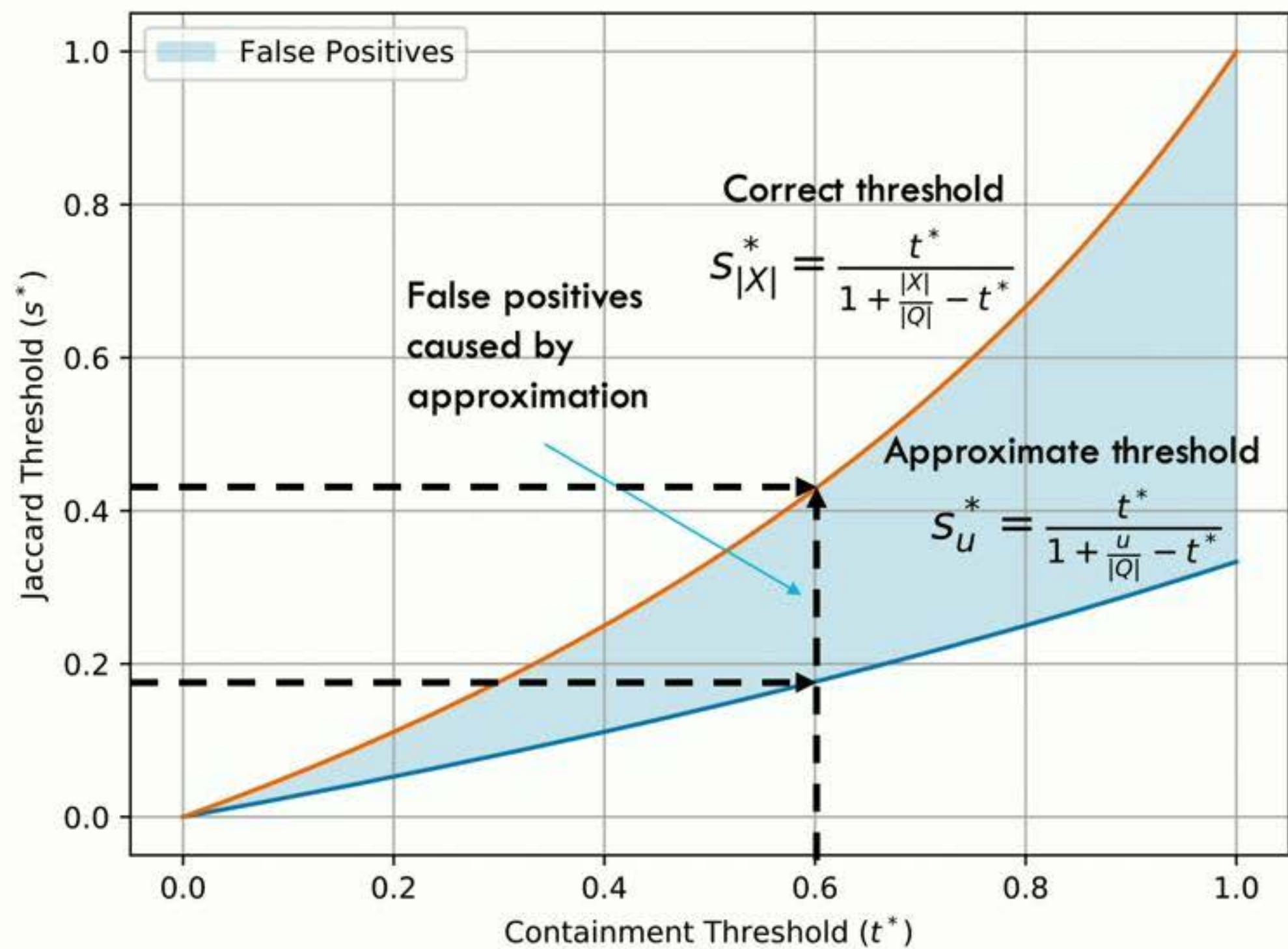
Alternative: use  $u$  (the largest set size) to approximate  $|X|$

$$s_u^* = \frac{t^*}{1 + \frac{u}{|Q|} - t^*}$$









# REDUCE FALSE POSITIVES — BRUTE FORCE

We can verify the exact containment of all sets in the output — “posting processing”

# REDUCE FALSE POSITIVES — BRUTE FORCE

We can verify the exact containment of all sets in the output — “posting processing”



Large set

<https://www.theguardian.com/lifeandstyle/2015/dec/02/eat-drink-bad-for-us-evolution>

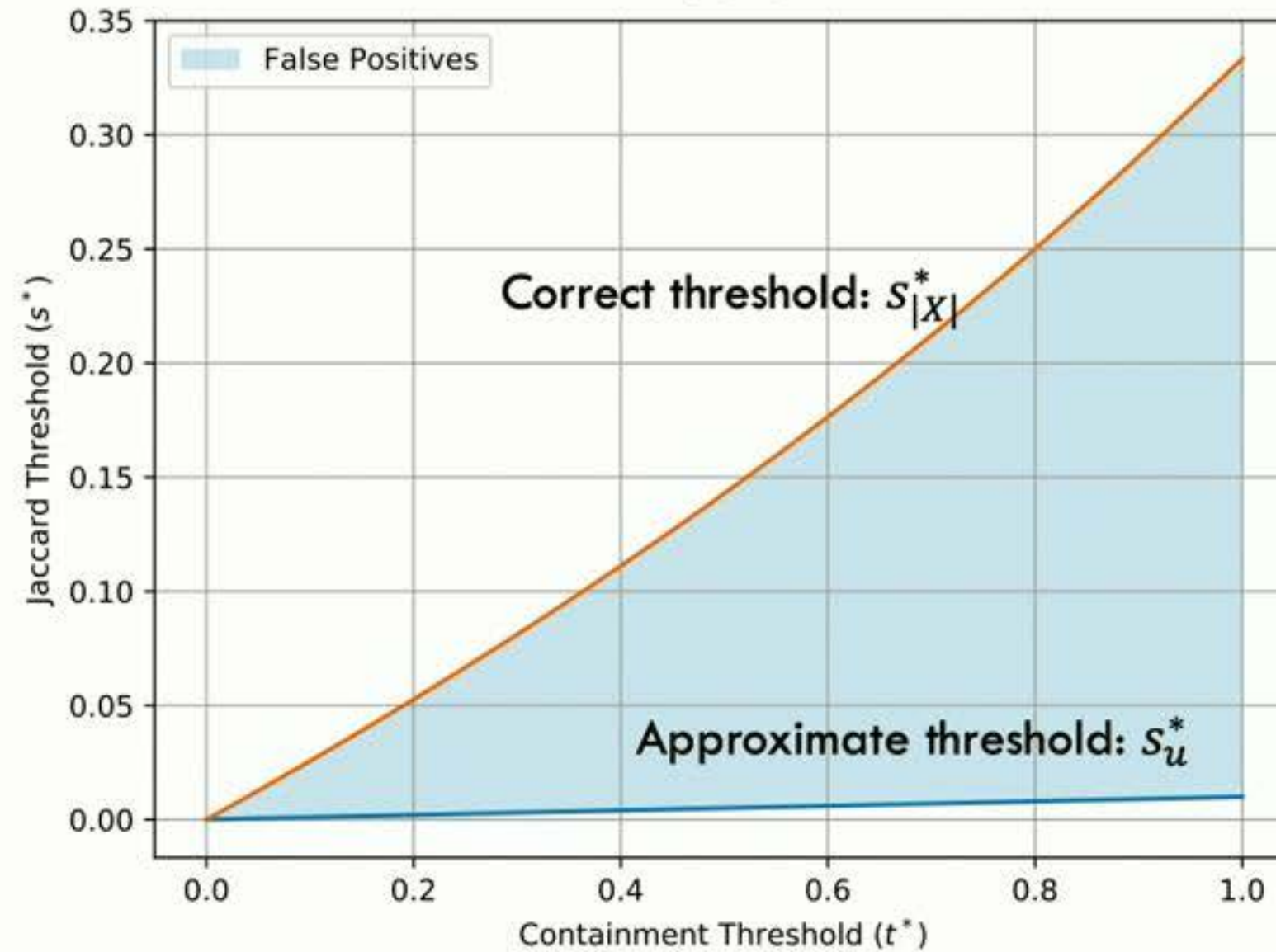
Set sizes in the index:

{3, 3, 3, 3, 4, 4, 4, 4, 99, 99, 99, 100, 100, 100}

Set sizes in the index:

$\{3, 3, 3, 3, 4, 4, 4, 4, 99, 99, 99, 100, 100, 100\}$

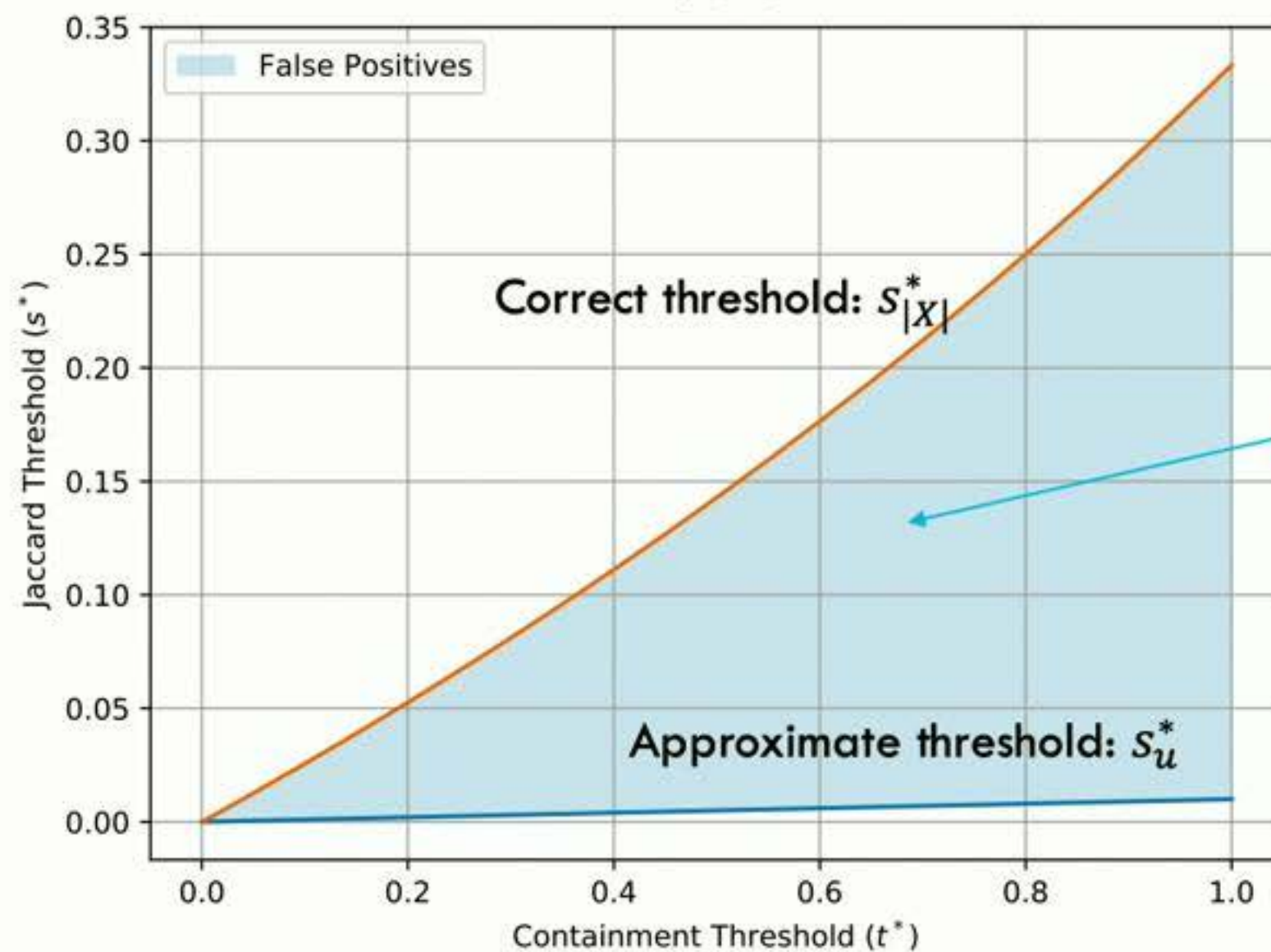
$u = 100, |X| = 3$



Set sizes in the index:

$\{3, 3, 3, 3, 4, 4, 4, 4, 99, 99, 99, 100, 100, 100\}$

$u = 100, |X| = 3$



Too many false positives to remove by post-processing



Use two indexes:

Set sizes in index 1

{3, 3, 3, 3, 4, 4, 4, 4}

Set sizes in index 2

{99, 99, 99, 100, 100, 100}

Use two indexes:

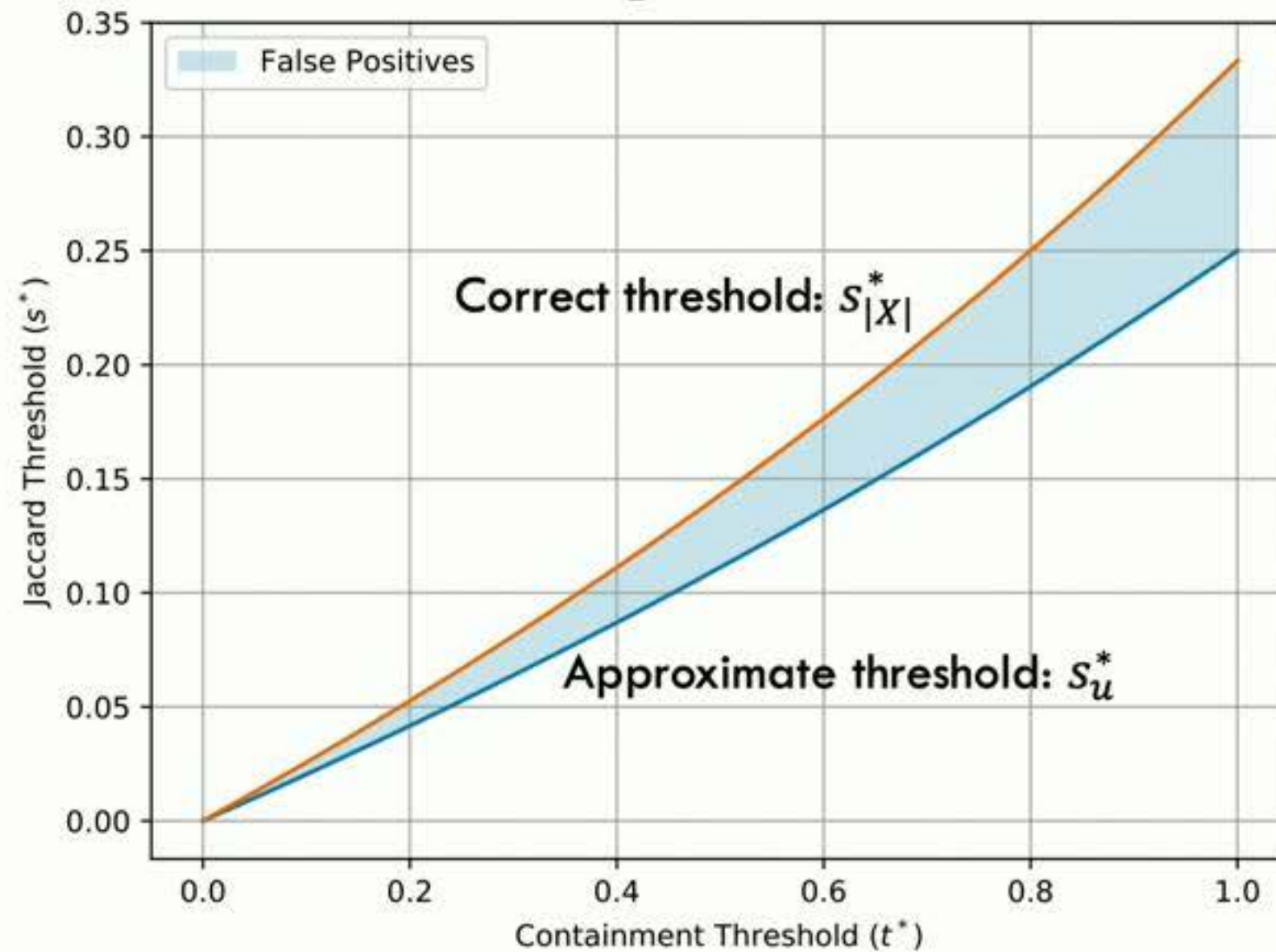
Set sizes in index 1

$\{3, 3, 3, 3, 4, 4, 4, 4\}$

Set sizes in index 2

$\{99, 99, 99, 100, 100, 100\}$

$u_1 = 4, |X| = 3$

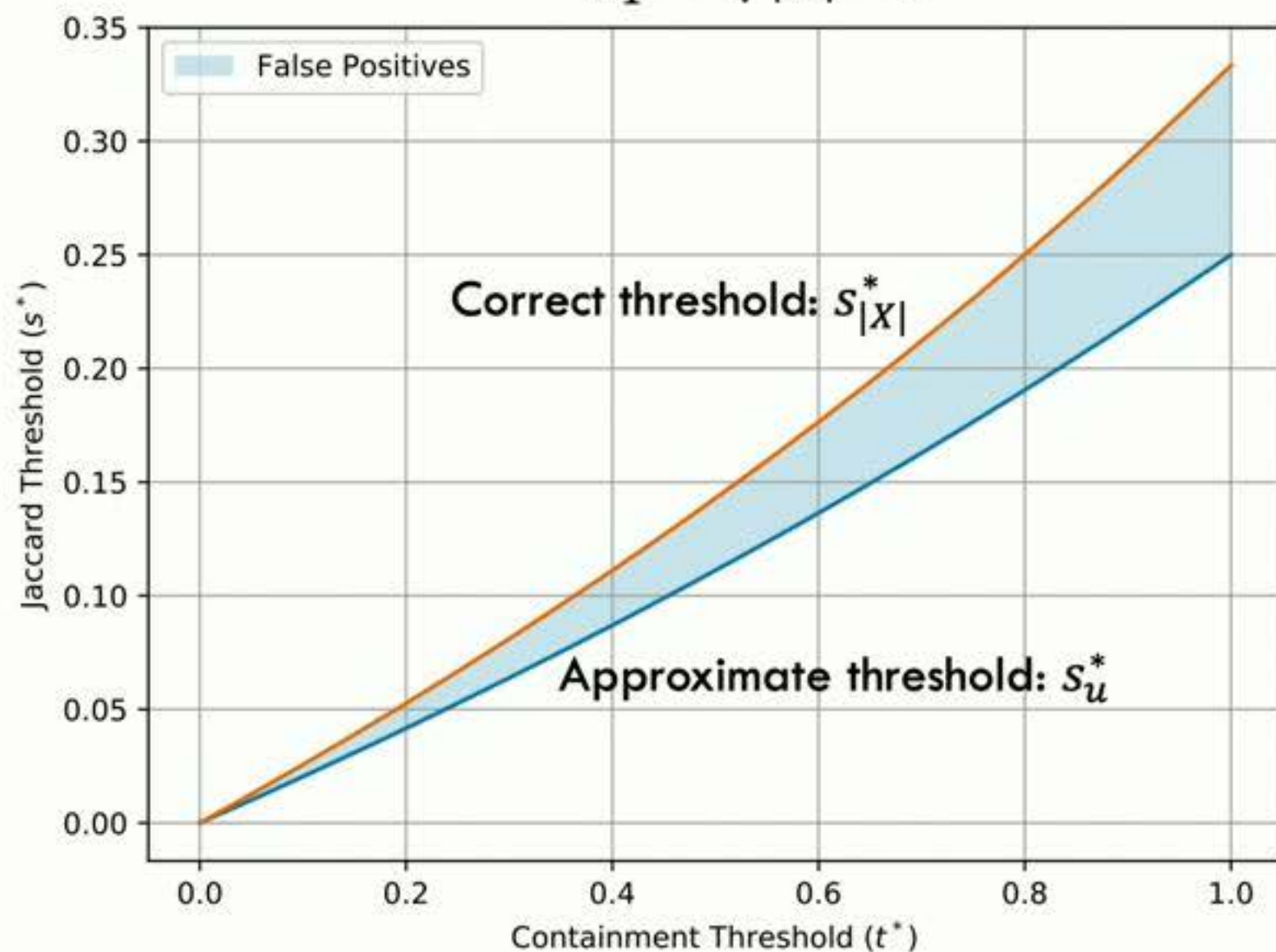


Use two indexes:

Set sizes in index 1

$\{3, 3, 3, 3, 4, 4, 4, 4\}$

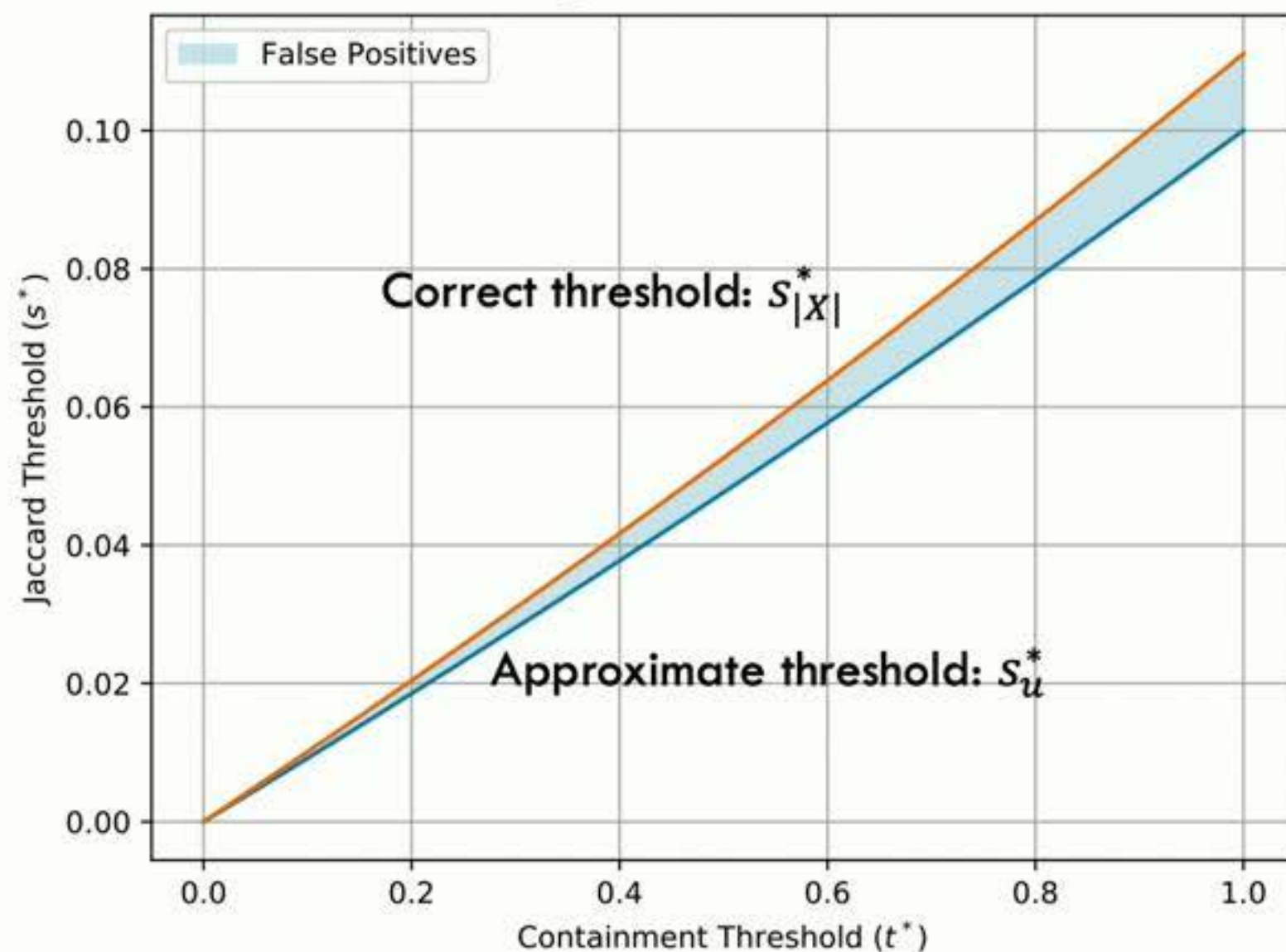
$u_1 = 4, |X| = 3$

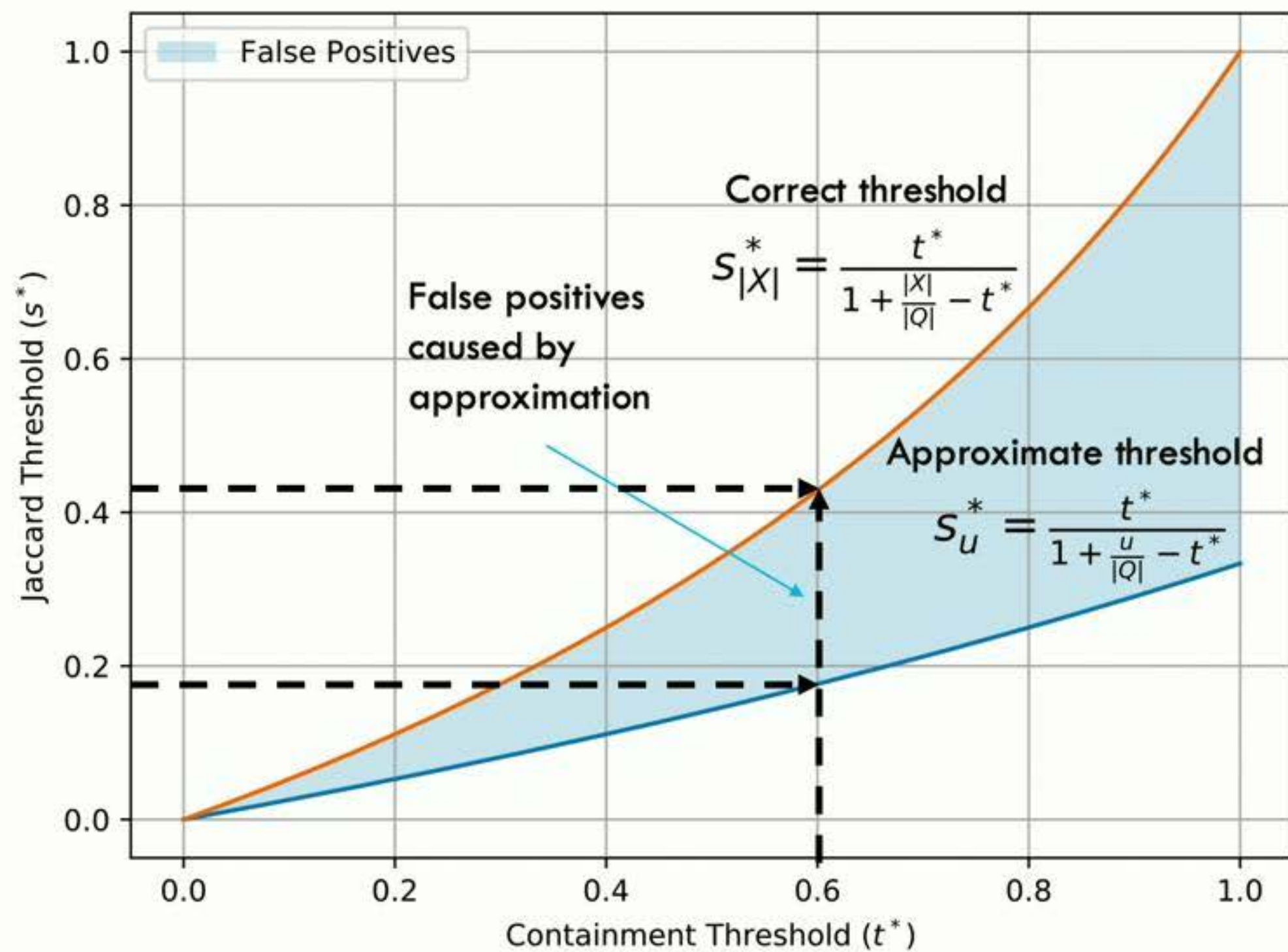


Set sizes in index 2

$\{99, 99, 99, 100, 100, 100\}$

$u_2 = 100, |X| = 99$





# REDUCE FALSE POSITIVES — BRUTE FORCE

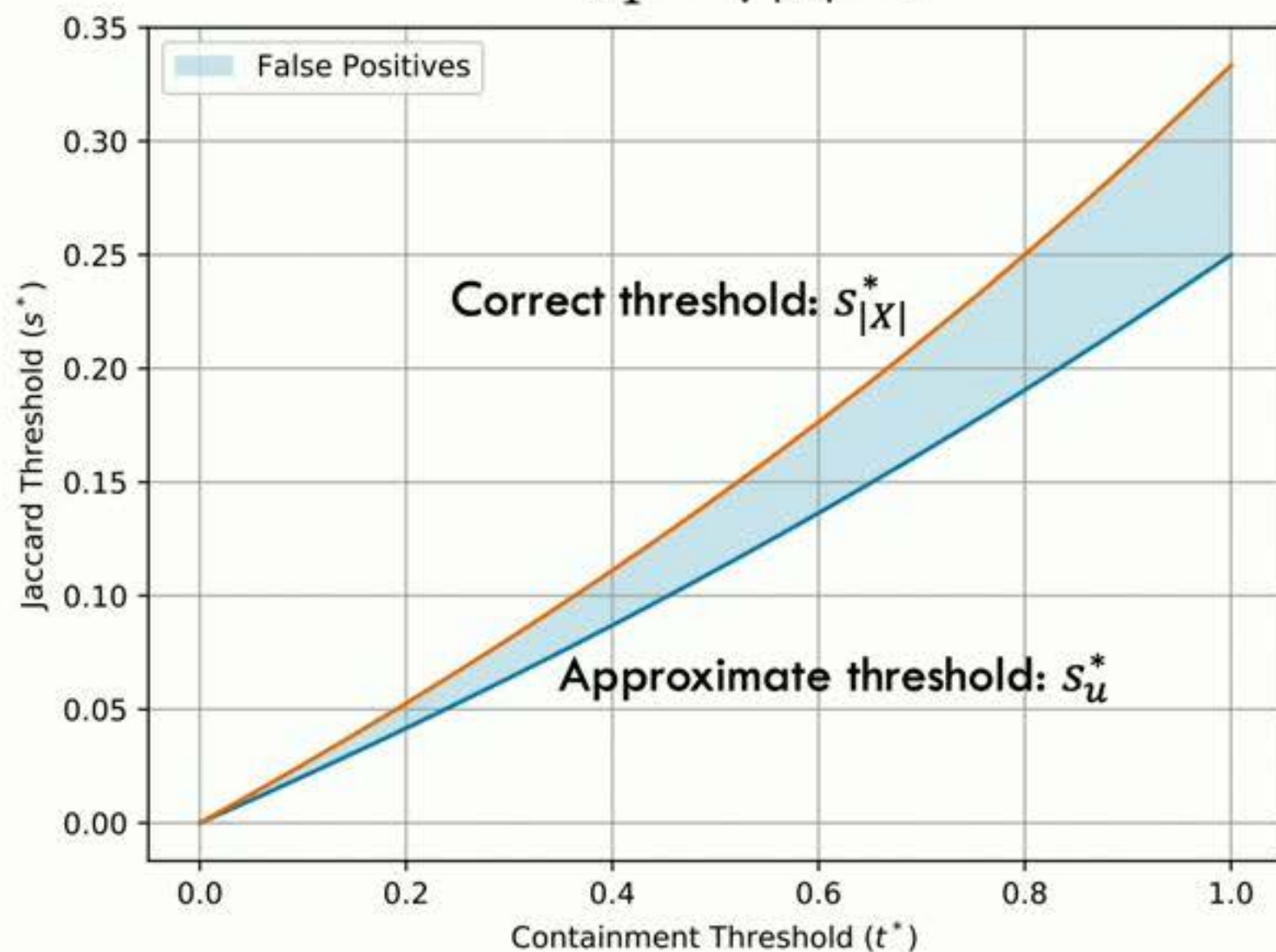
We can verify the exact containment of all sets in the output — “posting processing”

Use two indexes:

Set sizes in index 1

$\{3, 3, 3, 3, 4, 4, 4, 4\}$

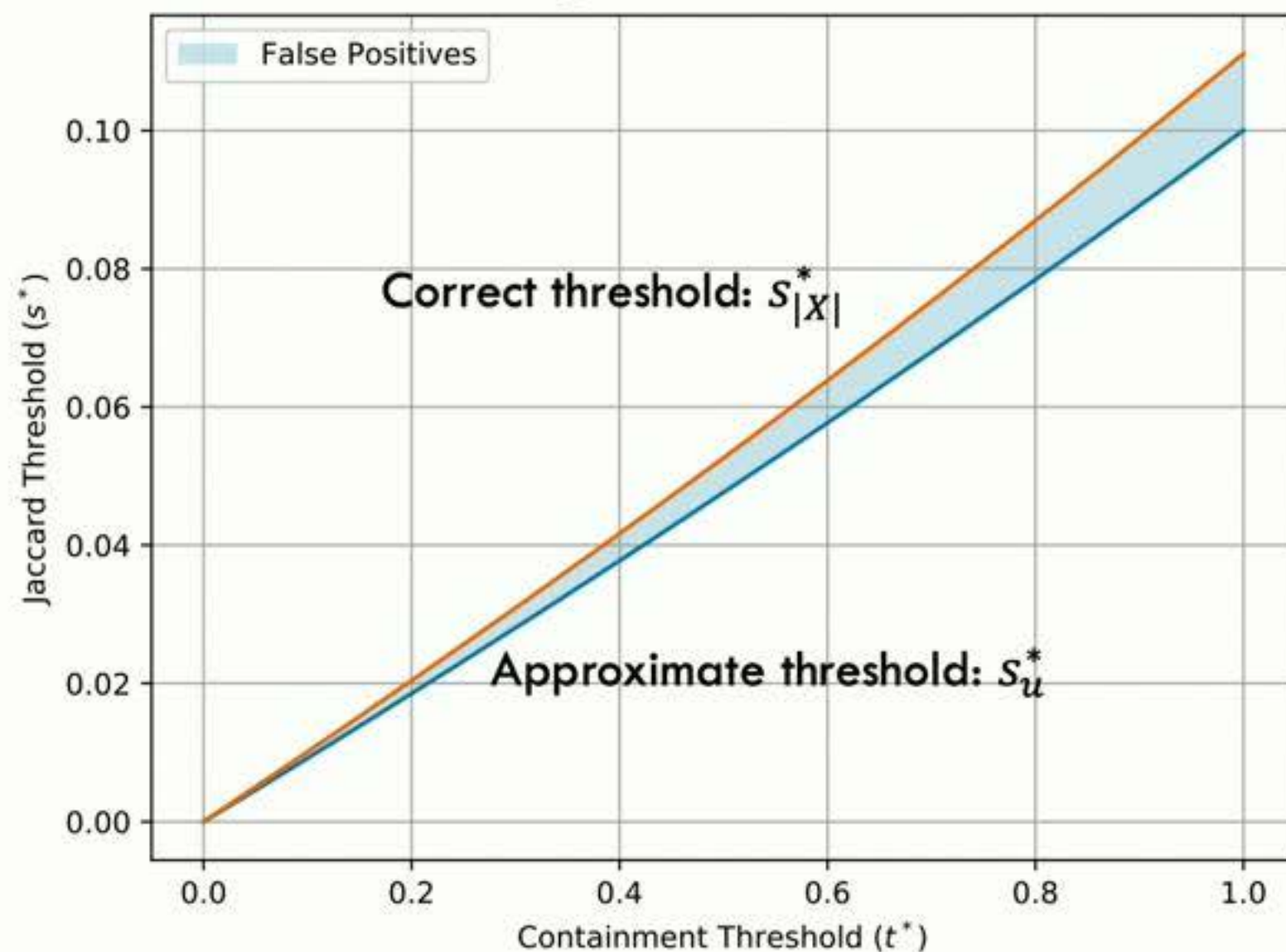
$u_1 = 4, |X| = 3$



Set sizes in index 2

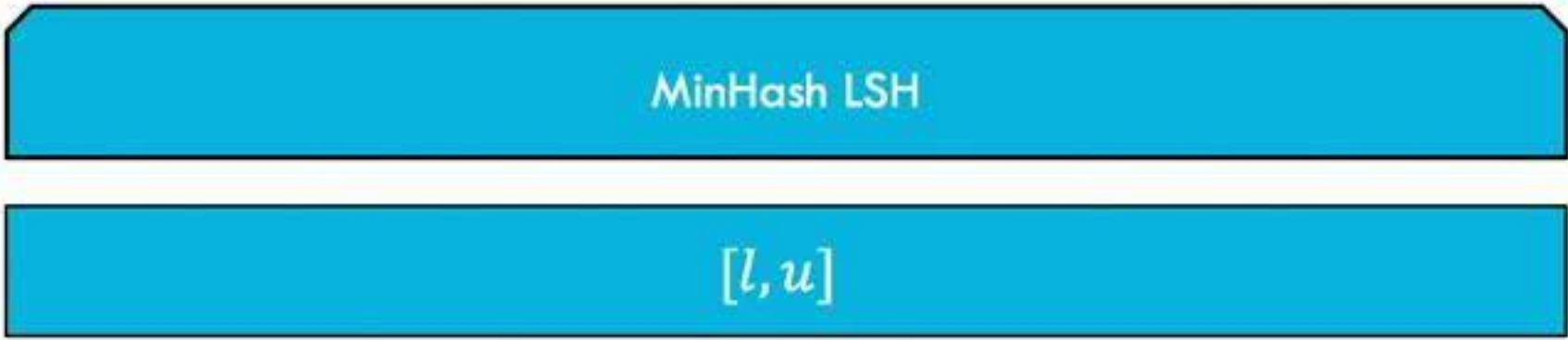
$\{99, 99, 99, 100, 100, 100\}$

$u_2 = 100, |X| = 99$



# REDUCE FALSE POSITIVES - PARTITIONS

Partition the sets into  $n$  disjoint partitions with increasing set sizes:

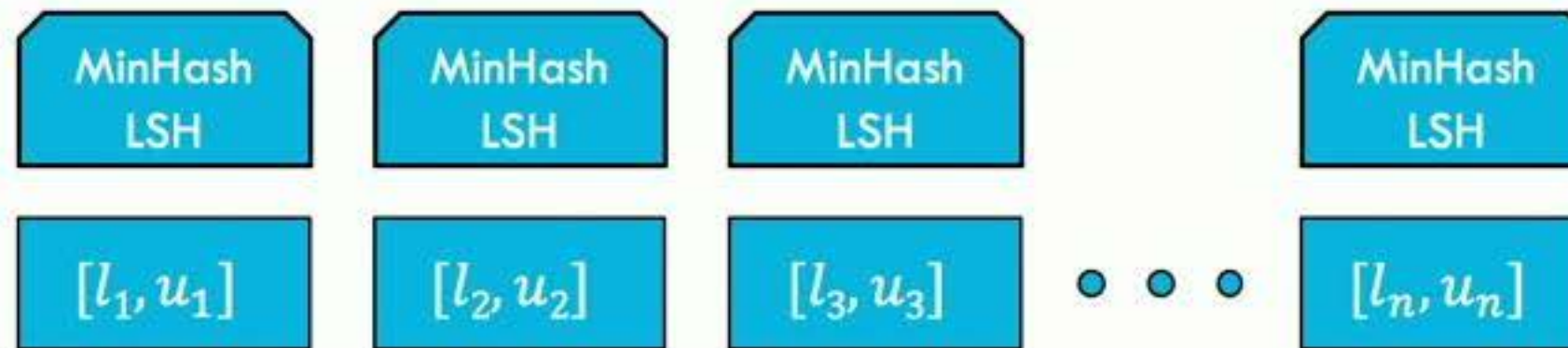


MinHash LSH

$[l, u]$

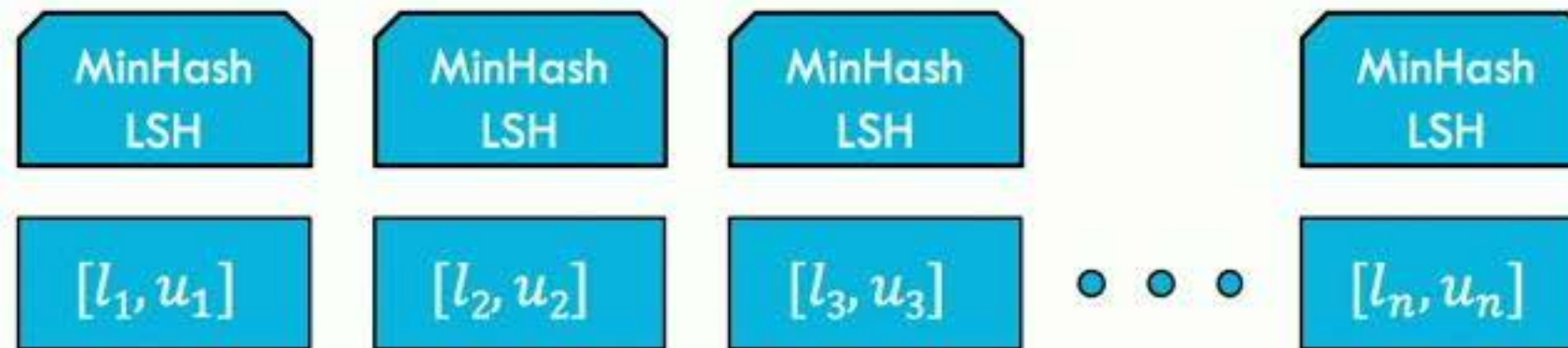
# REDUCE FALSE POSITIVES - PARTITIONS

Partition the sets into  $n$  disjoint partitions with increasing set sizes:



# REDUCE FALSE POSITIVES - PARTITIONS

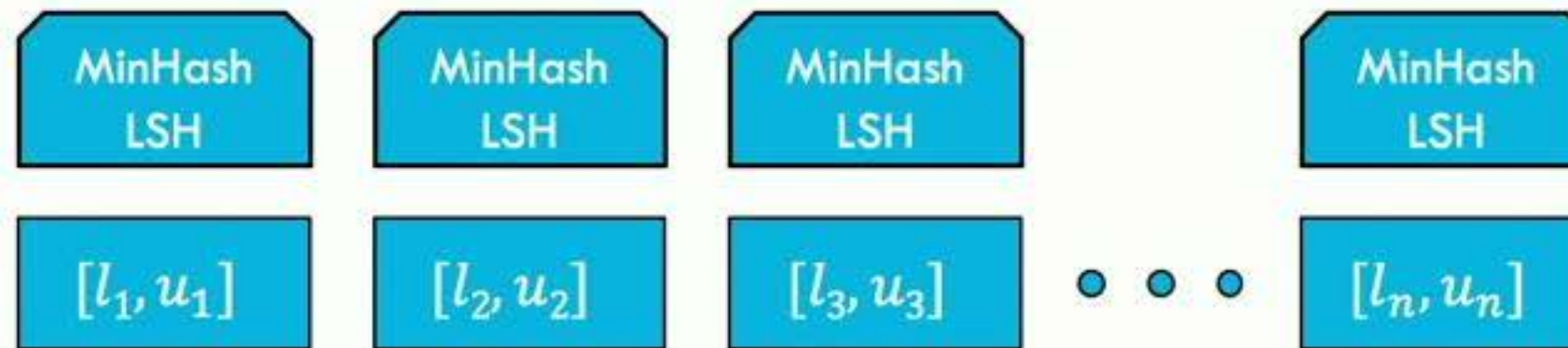
Partition the sets into  $n$  disjoint partitions with increasing set sizes:



Reduce false positive at the rate of  $\frac{1}{(u-u_i)^2}$  by using  $u_i$  instead of  $u$  for each partition!

# REDUCE FALSE POSITIVES - PARTITIONS

Partition the sets into  $n$  disjoint partitions with increasing set sizes:



Reduce false positive at the rate of  $\frac{1}{(u-u_i)^2}$  by using  $u_i$  instead of  $u$  for each partition!

Infinite partitions? Query time grows linearly with  $n$ !

# UPPER BOUND NUM. FALSE POSITIVES

Upper bound of number of false positives over a set size interval  $[l, u]$

$$N_{l,u}^{FP} \leq \sum_{x \in [l,u]} N_x \cdot \left(1 - \frac{x}{u}\right)$$

Number of sets with size  $x$

$[l, u]$



# OPTIMAL PARTITIONING — UNIFORM DIST.

Example:  $u = 100$ ,  $l = 1$ ,  $N_{l,u} = 100$ ,  $n = 2$ , assume uniform distribution



# OPTIMAL PARTITIONING — UNIFORM DIST.

Example:  $u = 100$ ,  $l = 1$ ,  $N_{l,u} = 100$ ,  $n = 2$ , assume uniform distribution



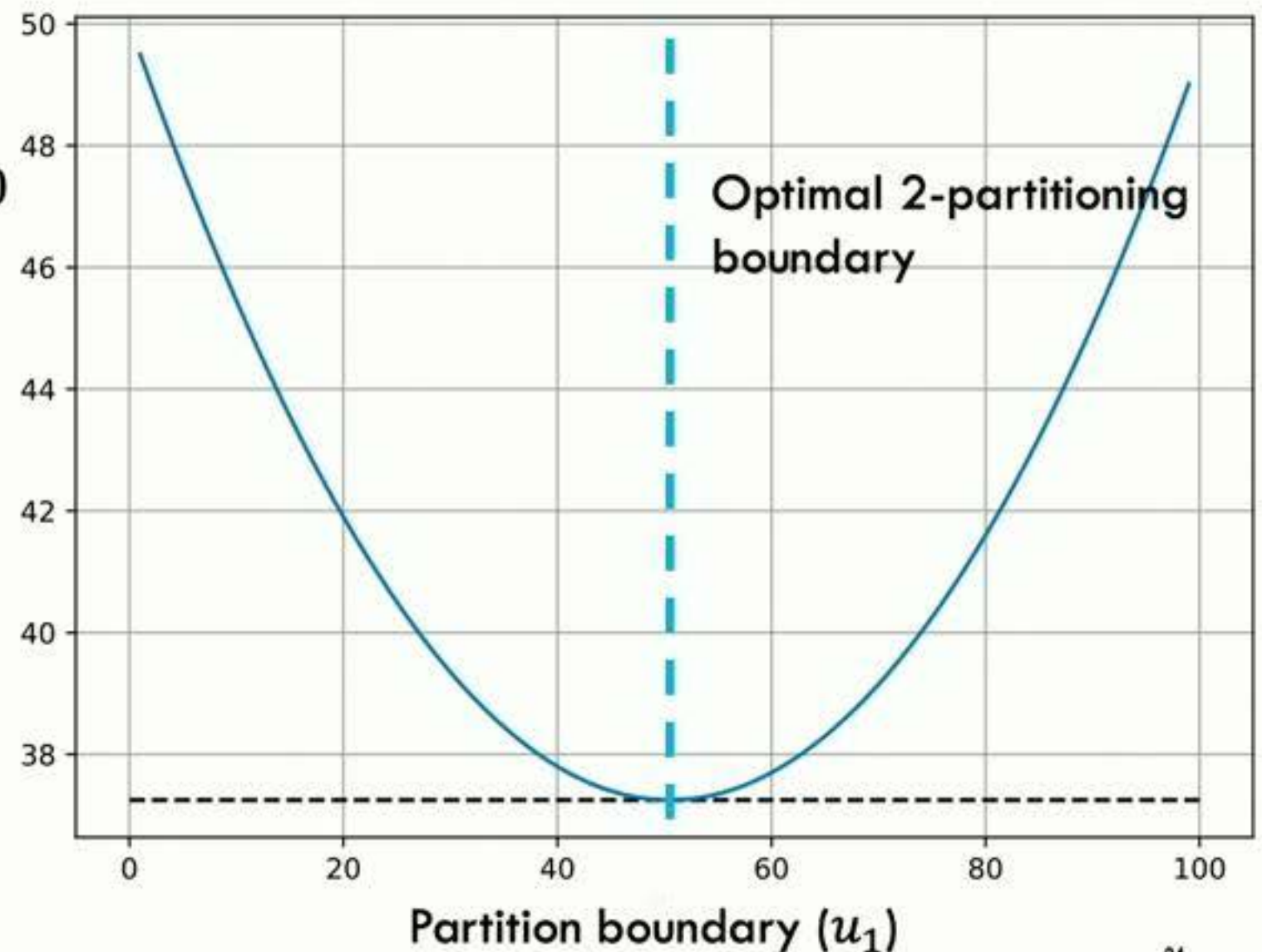
# OPTIMAL PARTITIONING – UNIFORM DIST.

Example:  $u = 100$ ,  $l = 1$ ,  $N_{l,u} = 100$ ,  $n = 2$ , assume uniform distribution

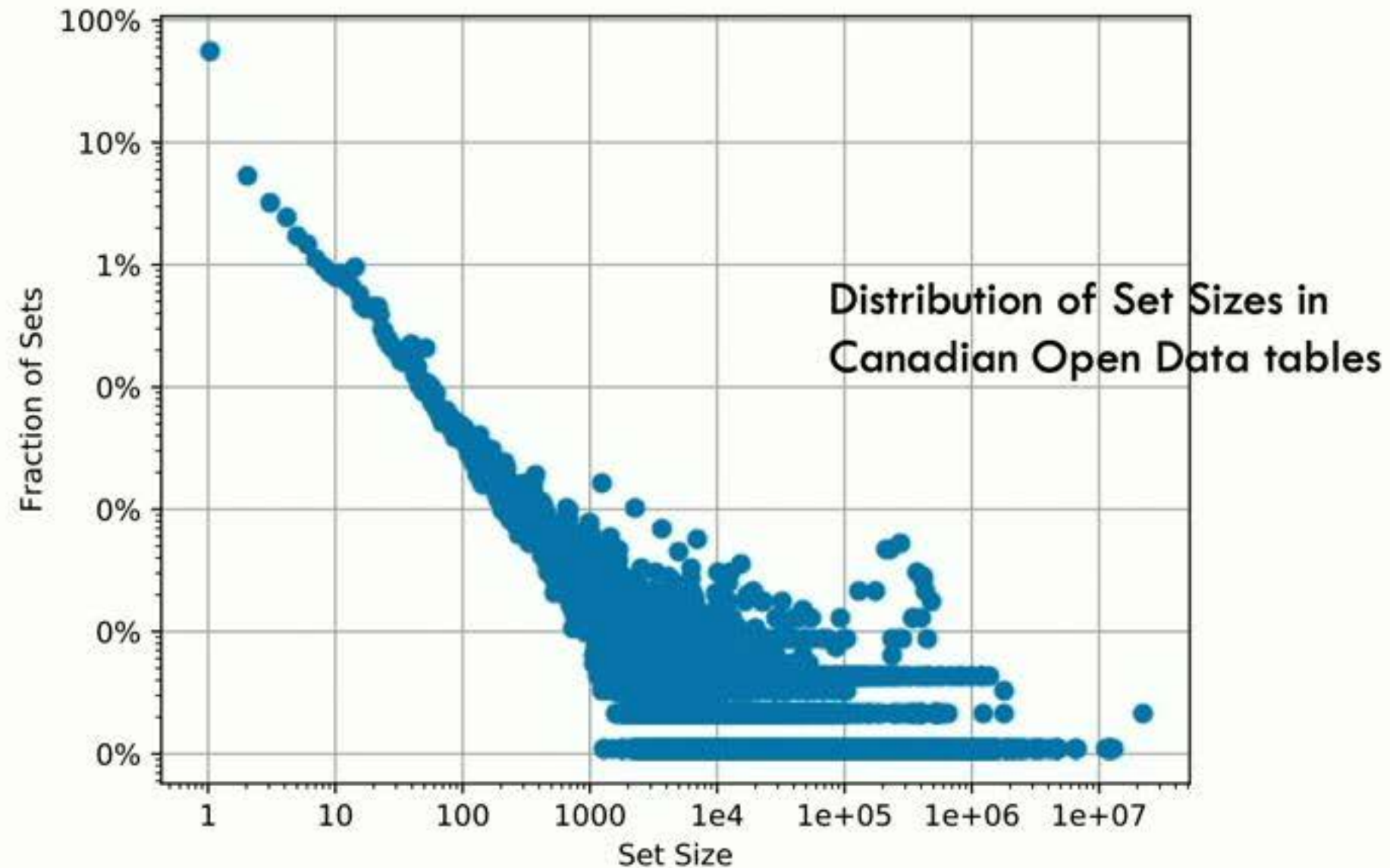


Number of false positives  
(upper bound):

$$N_{l_1, u_1}^{FP} + N_{l_2, u_2}^{FP}$$



# REAL SIZE DISTRIBUTIONS ARE NOT UNIFORM



# OPTIMAL PARTITIONING FOR ANY DISTRIBUTION

Setting  $u_1, u_2, \dots, u_{n-1}$  to minimize the sum of upper bounds of  $n$  partitions:

$$\sum_{i=1}^n \left( \sum_{x \in [l_i, u_i]} N_x \cdot \left( 1 - \frac{x}{u_i} \right) \right)$$

# OPTIMAL PARTITIONING FOR ANY DISTRIBUTION

Setting  $u_1, u_2, \dots, u_{n-1}$  to minimize the sum of upper bounds of  $n$  partitions:

$$\sum_{i=1}^n \left( \sum_{x \in [l_i, u_i]} N_x \cdot \left( 1 - \frac{x}{u_i} \right) \right)$$

Dynamic programming to the rescue!

Complexity is  $O(|N|^2 \cdot n)$

- $N$  is the domain of all set sizes,  $n$  is the number of partitions

# OPTIMAL PARTITIONING FOR ANY DISTRIBUTION

Setting  $u_1, u_2, \dots, u_{n-1}$  to minimize the sum of upper bounds of  $n$  partitions:

$$\sum_{i=1}^n \left( \sum_{x \in [l_i, u_i]} N_x \cdot \left( 1 - \frac{x}{u_i} \right) \right)$$

Dynamic programming to the rescue!

Complexity is  $O(|N|^2 \cdot n)$

- $N$  is the domain of all set sizes,  $n$  is the number of partitions

The complexity is quite manageable in practice as:  $|N| \ll |u_n - l_1|$

- *Number of bank accounts*  $\ll$  *|Amount in Jeff Bezos' – Amount in Eric's|*
- All of Canada, U.S., and U.K. Open Data:  $|N| = 14\text{K}$ ,  $< 10$  min in Python on a laptop
- WDC Web Table:  $|N| = 4\text{K}$

# LSH ENSEMBLE: KEY POINTS

Given a query set  $Q$  and containment threshold  $t^*$

- Indexing: DP, MinHash LSH on partitions,  $s_u^* = \frac{t^*}{1 + \frac{u}{|Q|} - t^*}$
- Querying: probing hash tables; optionally remove false positives
- Also handles arbitrary  $|Q|$  and  $t^*$

# LSH ENSEMBLE: KEY POINTS

Given a query set  $Q$  and containment threshold  $t^*$

- Indexing: DP, MinHash LSH on partitions,  $s_u^* = \frac{t^*}{1 + \frac{u}{|Q|} - t^*}$
- Querying: probing hash tables; optionally remove false positives
- Also handles arbitrary  $|Q|$  and  $t^*$

Parallelize query over partitions

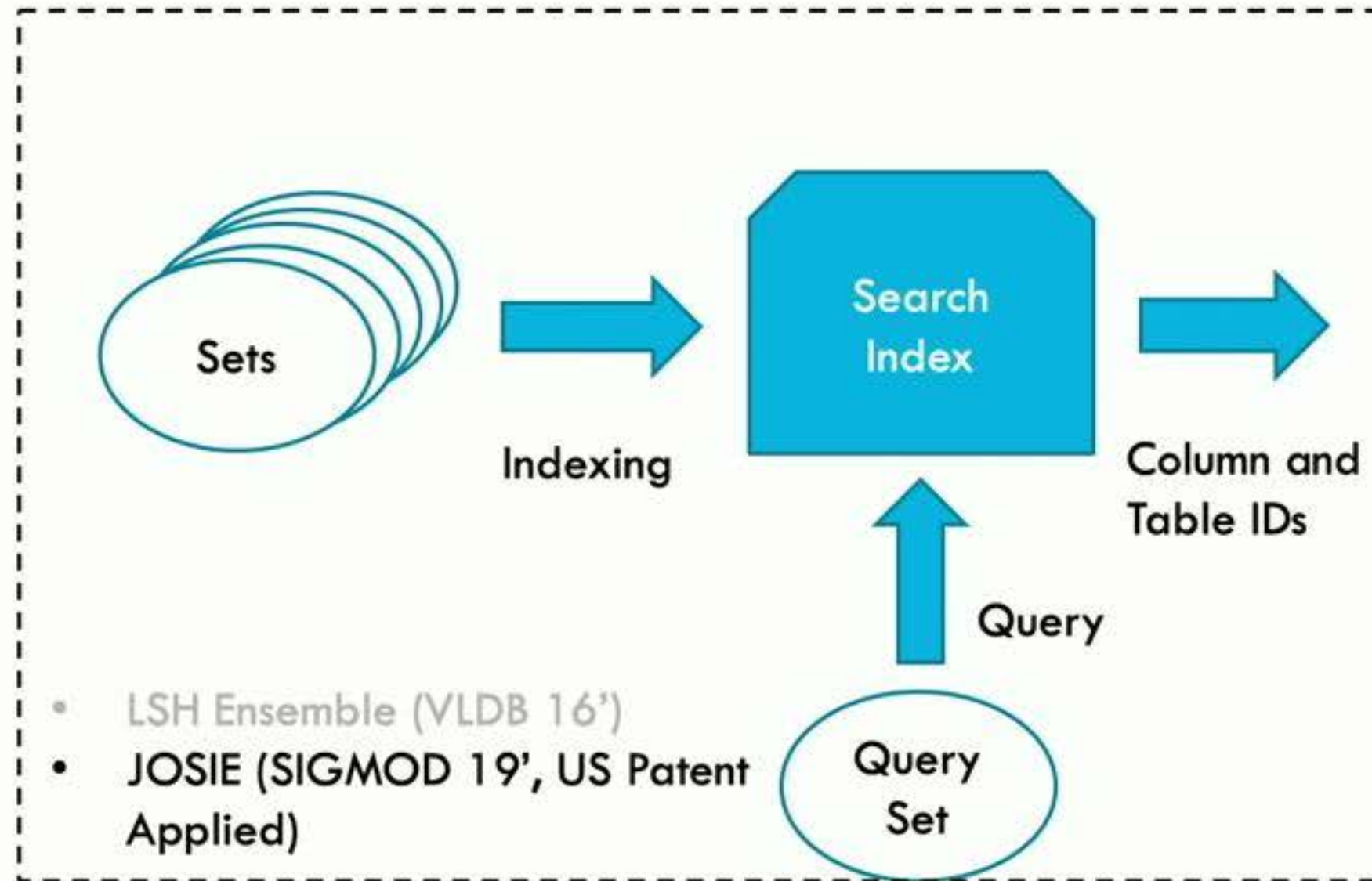
Small query memory footprints

- MinHash is only 1KB (128 hash functions)

Scale to massive number of sets

- Tested using 220 Million sets from WDC Web Table (parallel over 5 machines)

# A JOINABLE TABLE SEARCH SYSTEM



# JOSIE VS. LSH ENSEMBLE

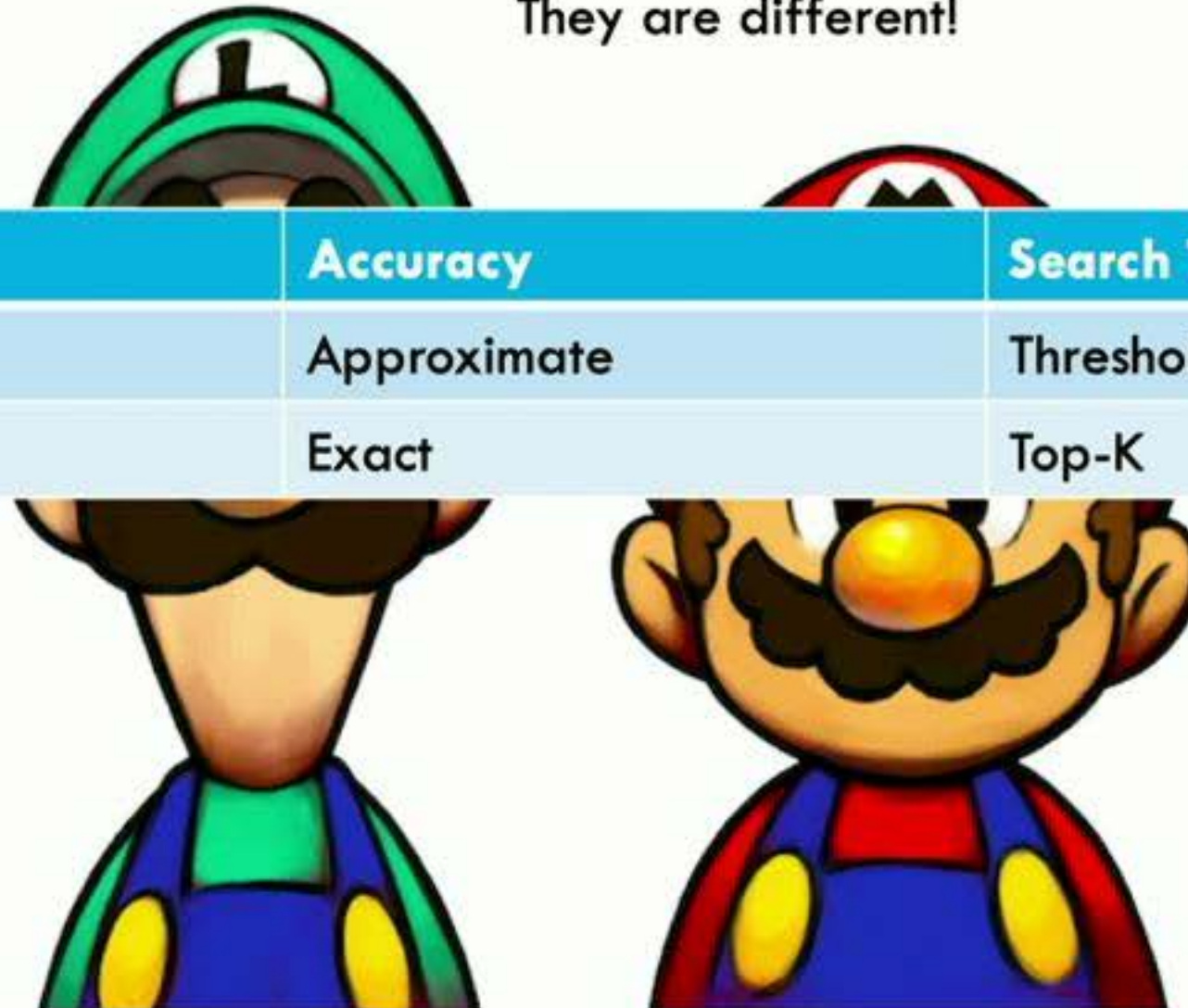
They are different!



# JOSIE VS. LSH ENSEMBLE

They are different!

| Algorithm    | Accuracy    | Search Type |
|--------------|-------------|-------------|
| LSH Ensemble | Approximate | Threshold   |
| JOSIE        | Exact       | Top-K       |



# WHEN TO USE TOP-K SEARCH?

Threshold search: user must specify a containment threshold

- User may not know a good threshold
- User may not understand containment

Top-K problem: just return the best  $k$  results

- No knowledge of relevance measure is required
- We showed that for small  $k$  ( $< 20$ ), our exact top- $k$  algorithm can be faster than LSH Ensemble with decreasing threshold hack!

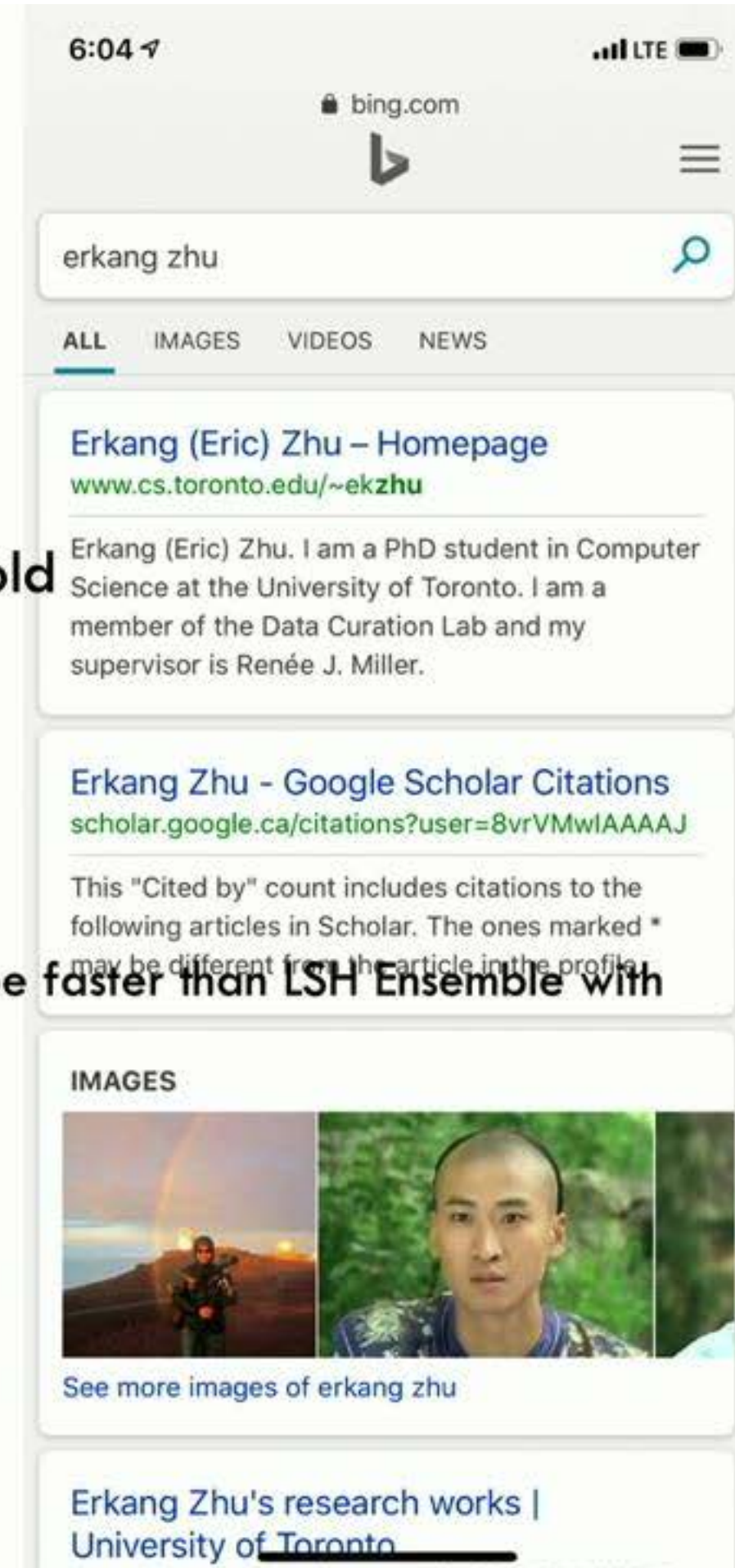
# WHEN TO USE TOP-K SEARCH?

**Threshold search: user must specify a containment threshold**

- User may not know a good threshold
- User may not understand containment

**Top-K problem: just return the best  $k$  results**

- No knowledge of relevance measure is required
- We showed that for small  $k$  ( $< 20$ ), our exact top-k algorithm can be faster than LSH Ensemble with decreasing threshold hack!



# WHEN TO USE TOP-K SEARCH?

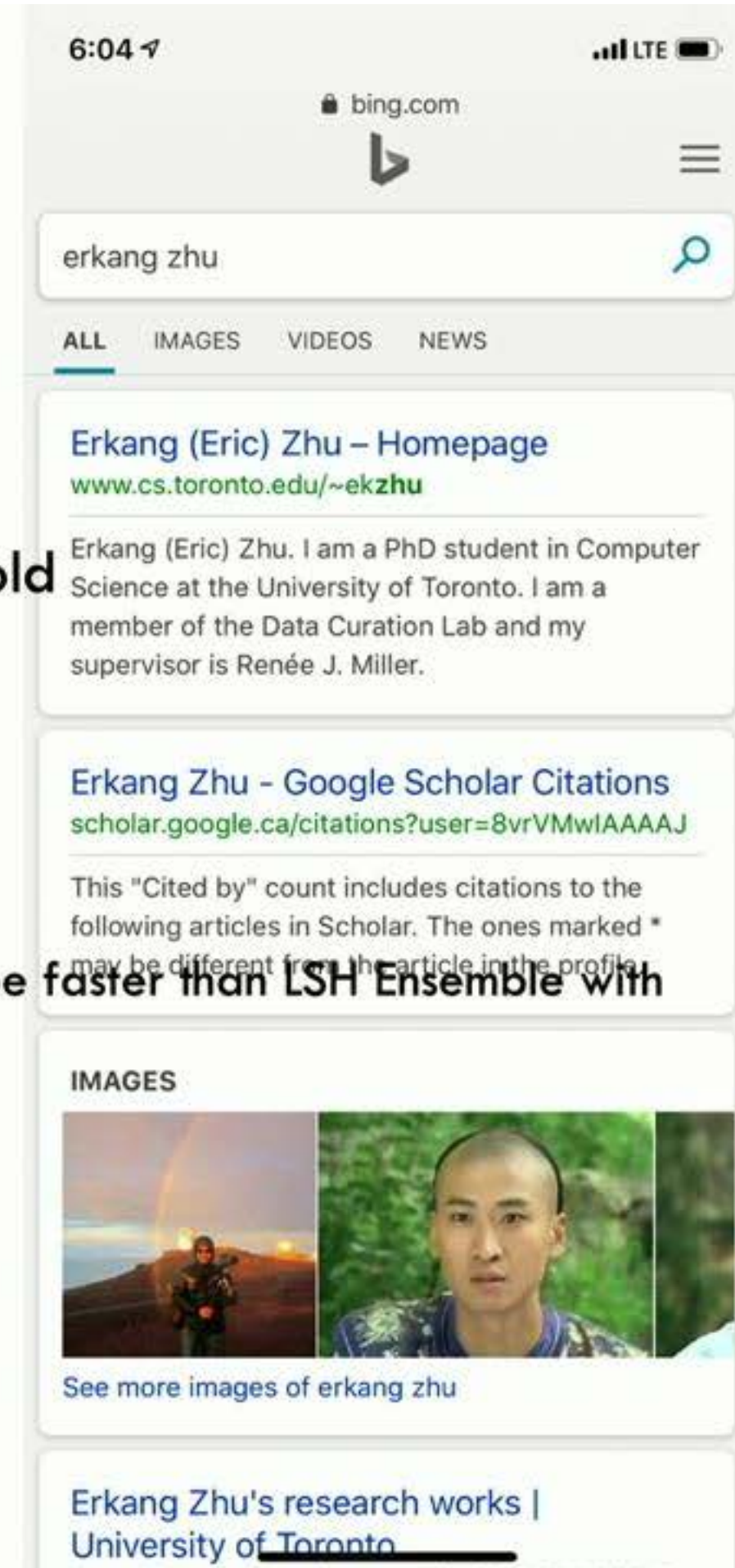
**Threshold search: user must specify a containment threshold**

- User may not know a good threshold
- User may not understand containment

**Top-K problem: just return the best  $k$  results**

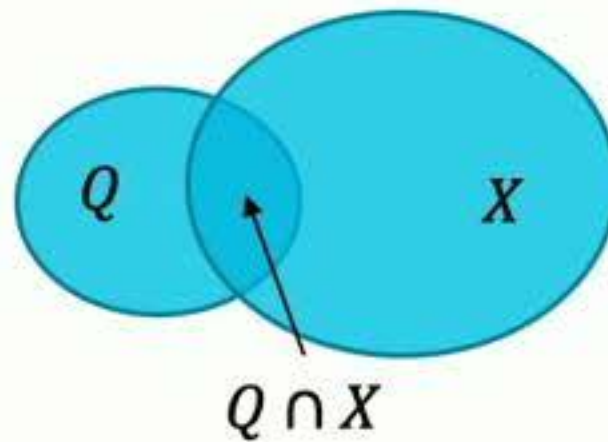
- No knowledge of relevance measure is required
- We showed that for small  $k$  ( $< 20$ ), our exact top-k algorithm can be faster than LSH Ensemble with decreasing threshold hack!

**Use top-k for less sophisticated users and small  $k$**

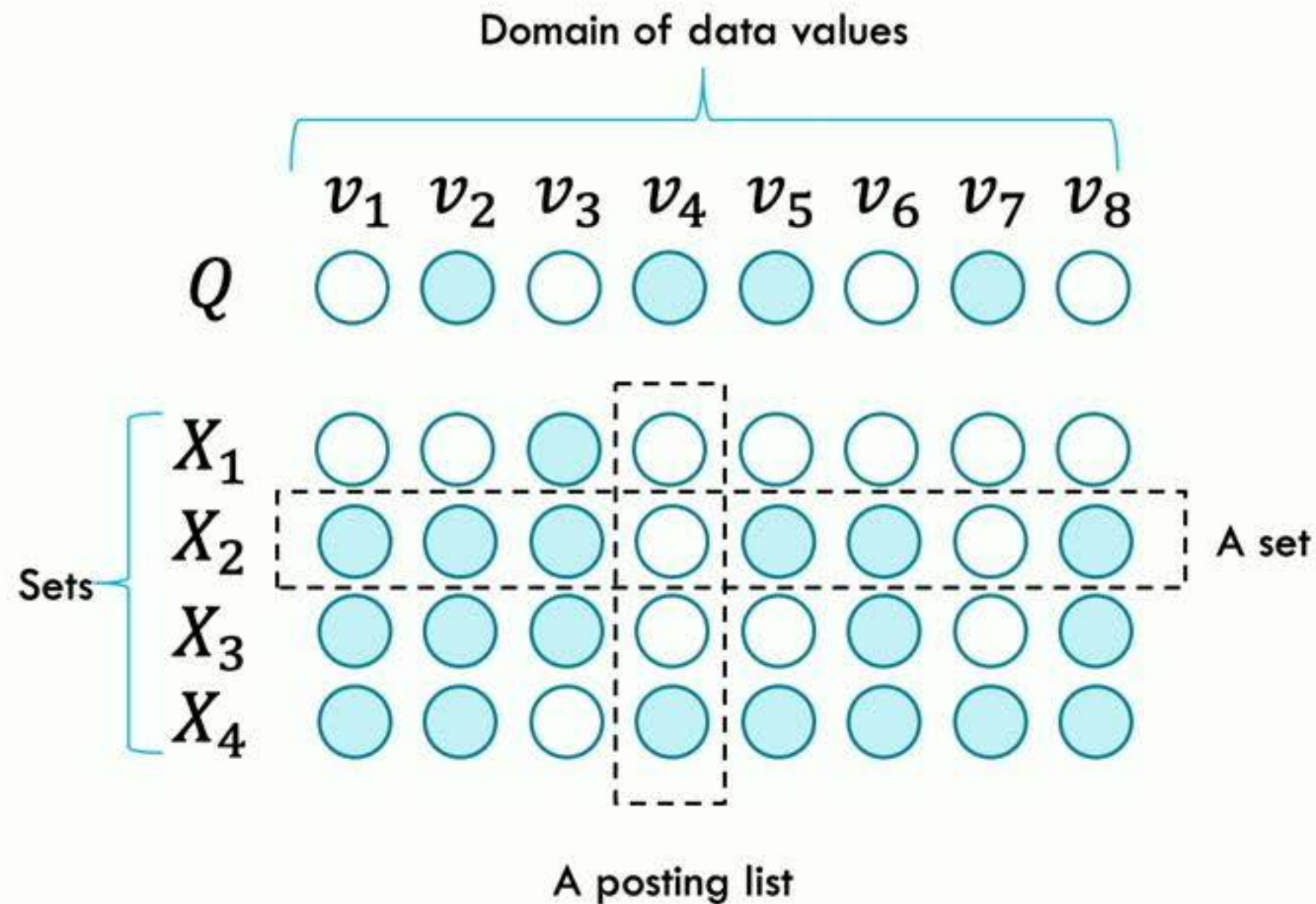


# TOP-K OVERLAP SET SIMILARITY SEARCH

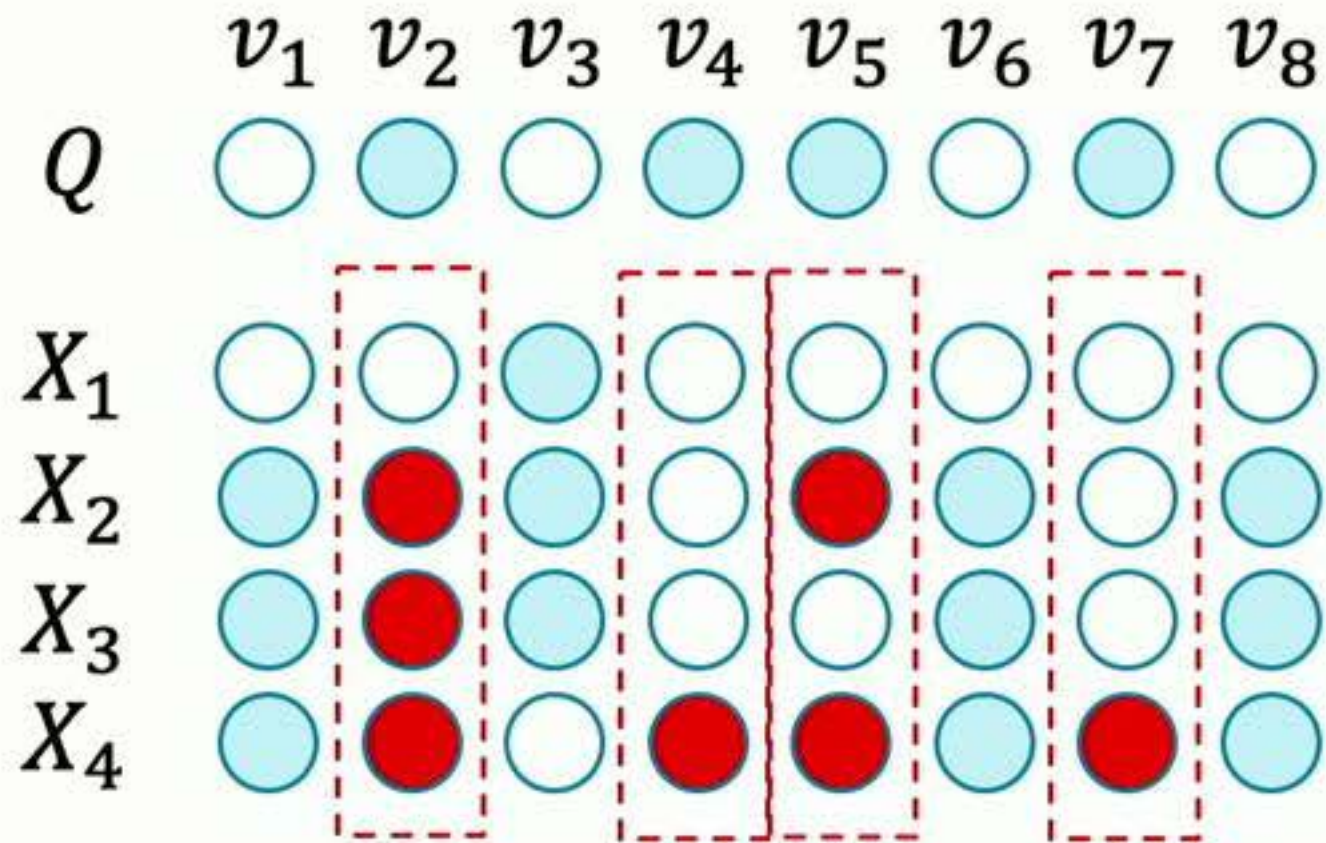
Overlap as the relevance measure:  $|Q \cap X|$ , equivalent to containment



# INVERTED INDEX — A MATRIX PERSPECTIVE



# BASELINES FOR FIND TOP-1

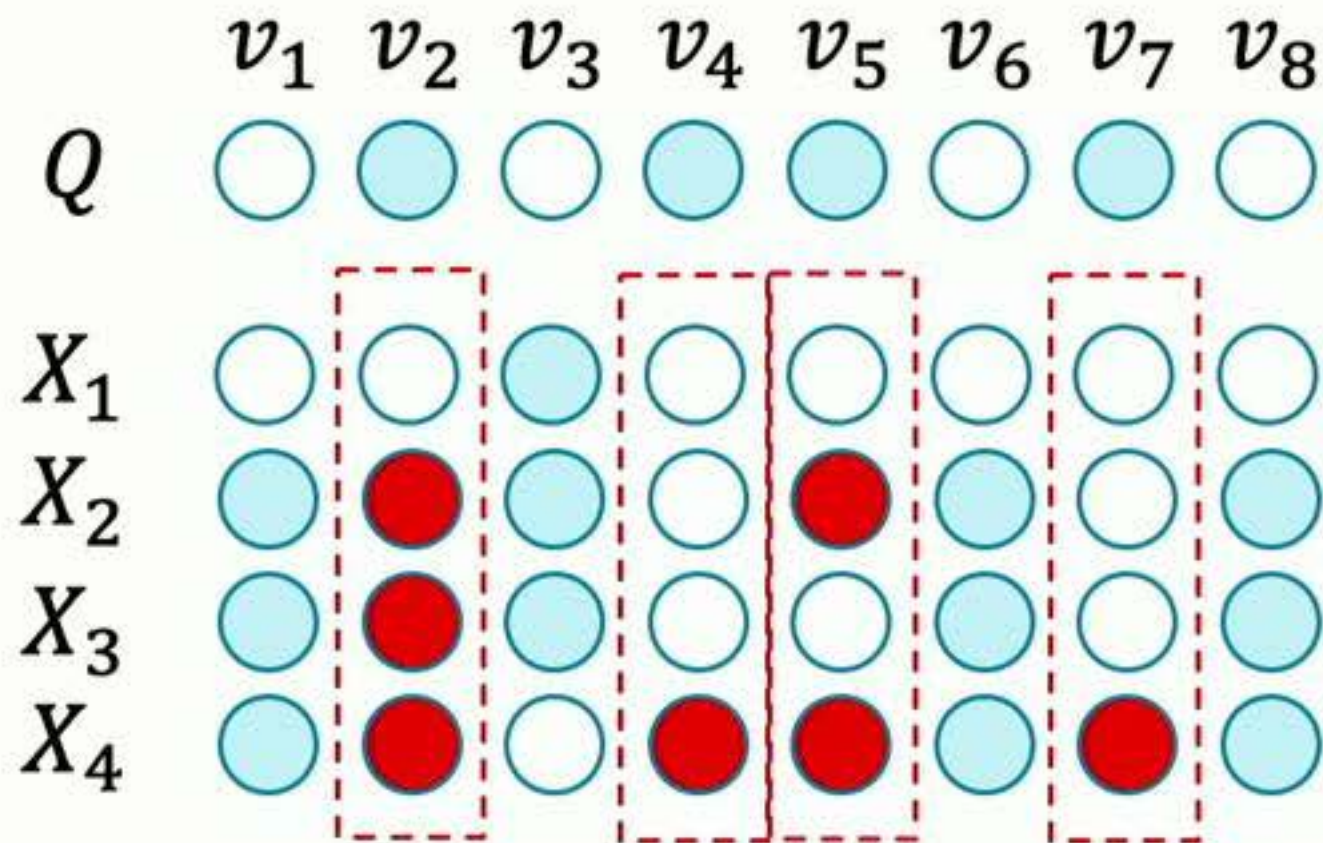


Posting list union<sup>1</sup>: reading all posting lists  
costs **7** values

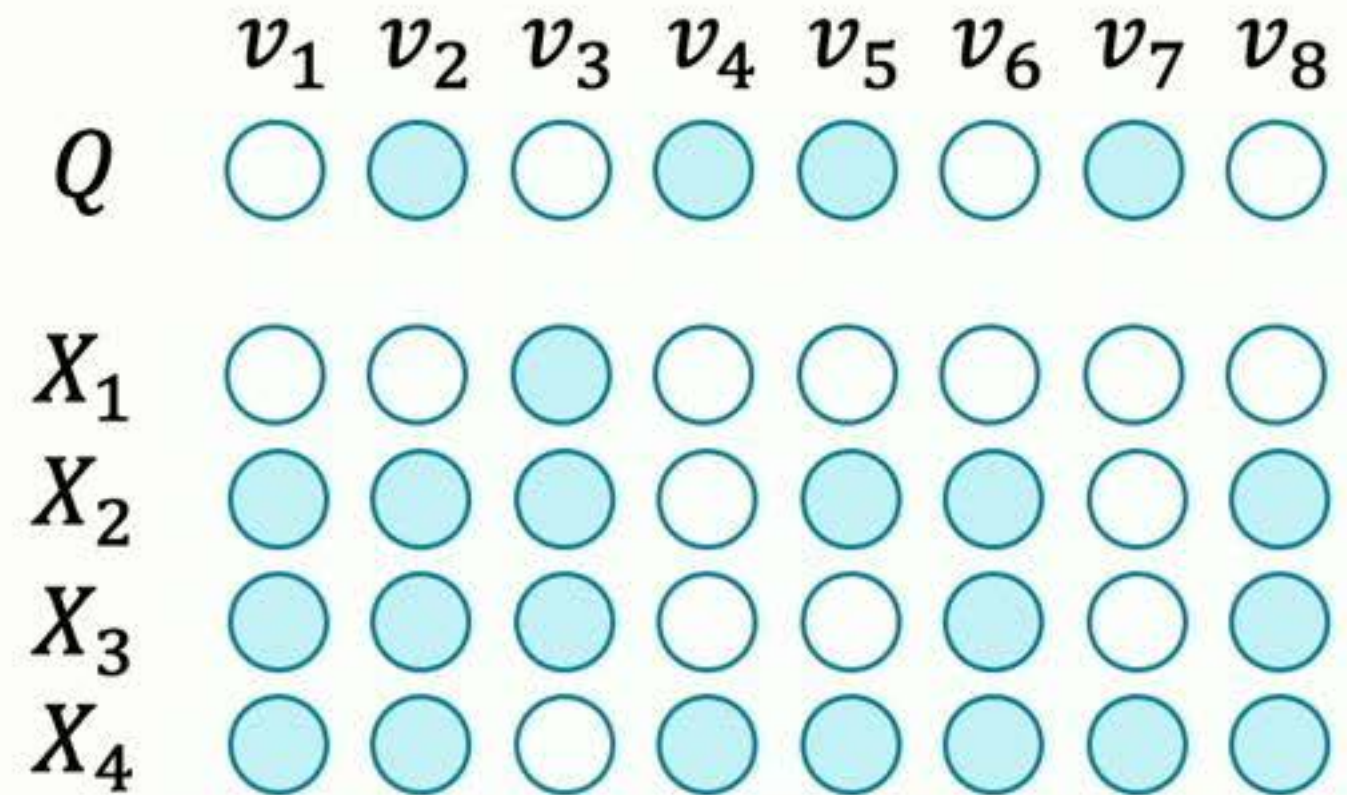
[1] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. Introduction to Information Retrieval. Cambridge University Press, 2008

[2] Chuan Xiao, Wei Wang, Xuemin Lin, and Haichuan Shang. Top-k set similarity joins. In ICDE, pages 916–927, 2009.

# BASELINES FOR FIND TOP-1



Posting list union<sup>1</sup>: reading all posting lists costs **7** values

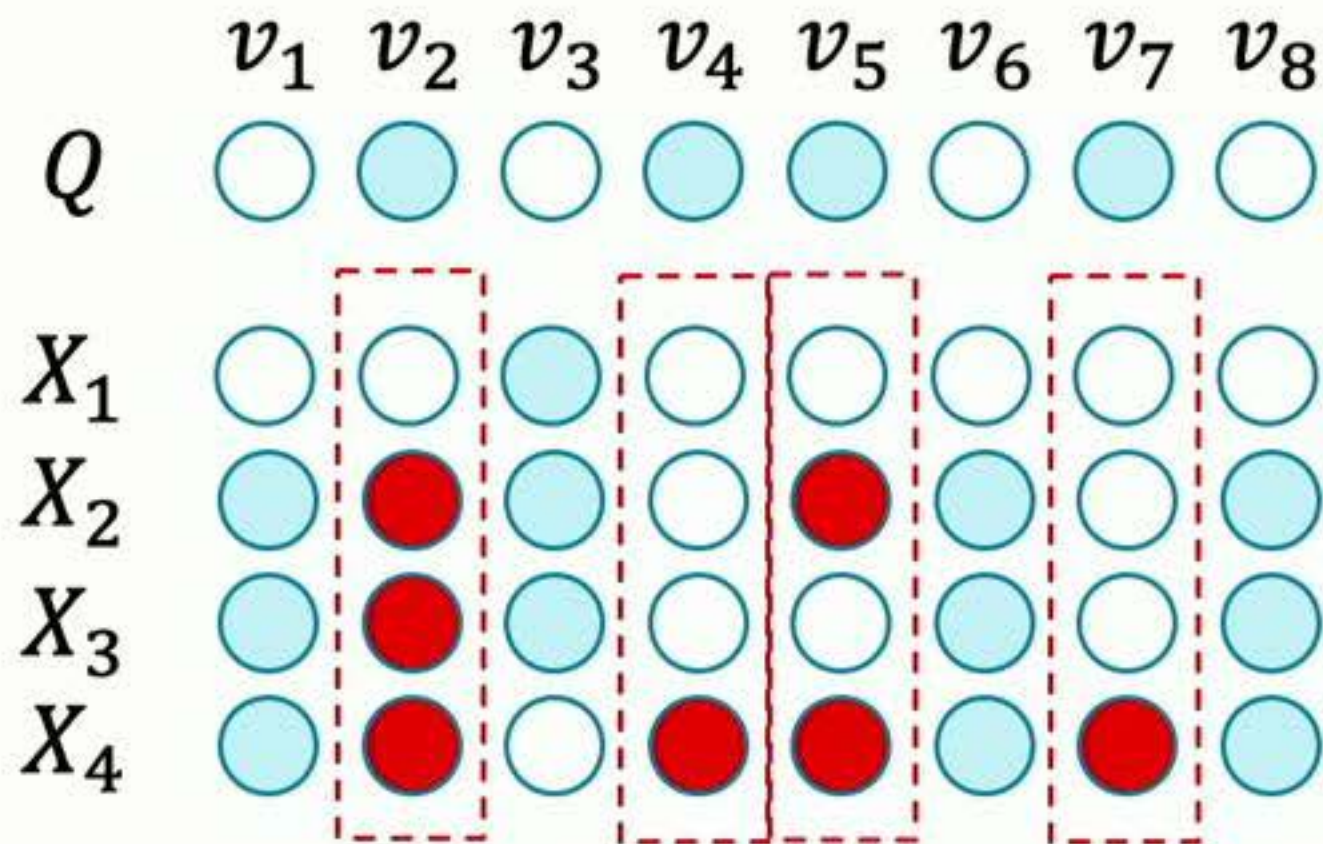


Prefix filtering<sup>2</sup>: reading candidate sets

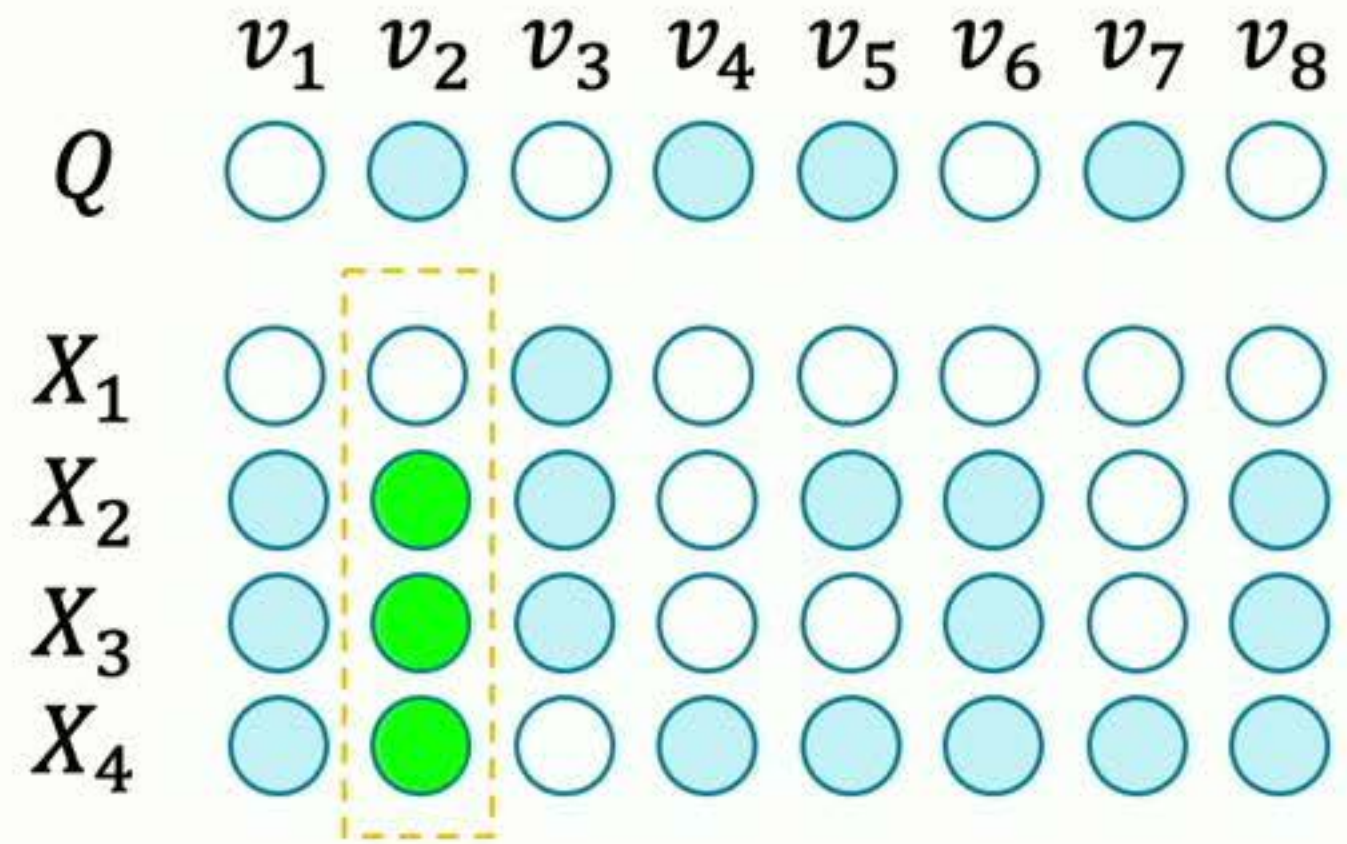
[1] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. Introduction to Information Retrieval. Cambridge University Press, 2008

[2] Chuan Xiao, Wei Wang, Xuemin Lin, and Haichuan Shang. Top-k set similarity joins. In ICDE, pages 916–927, 2009.

# BASELINES FOR FIND TOP-1



Posting list union<sup>1</sup>: reading all posting lists costs **7** values

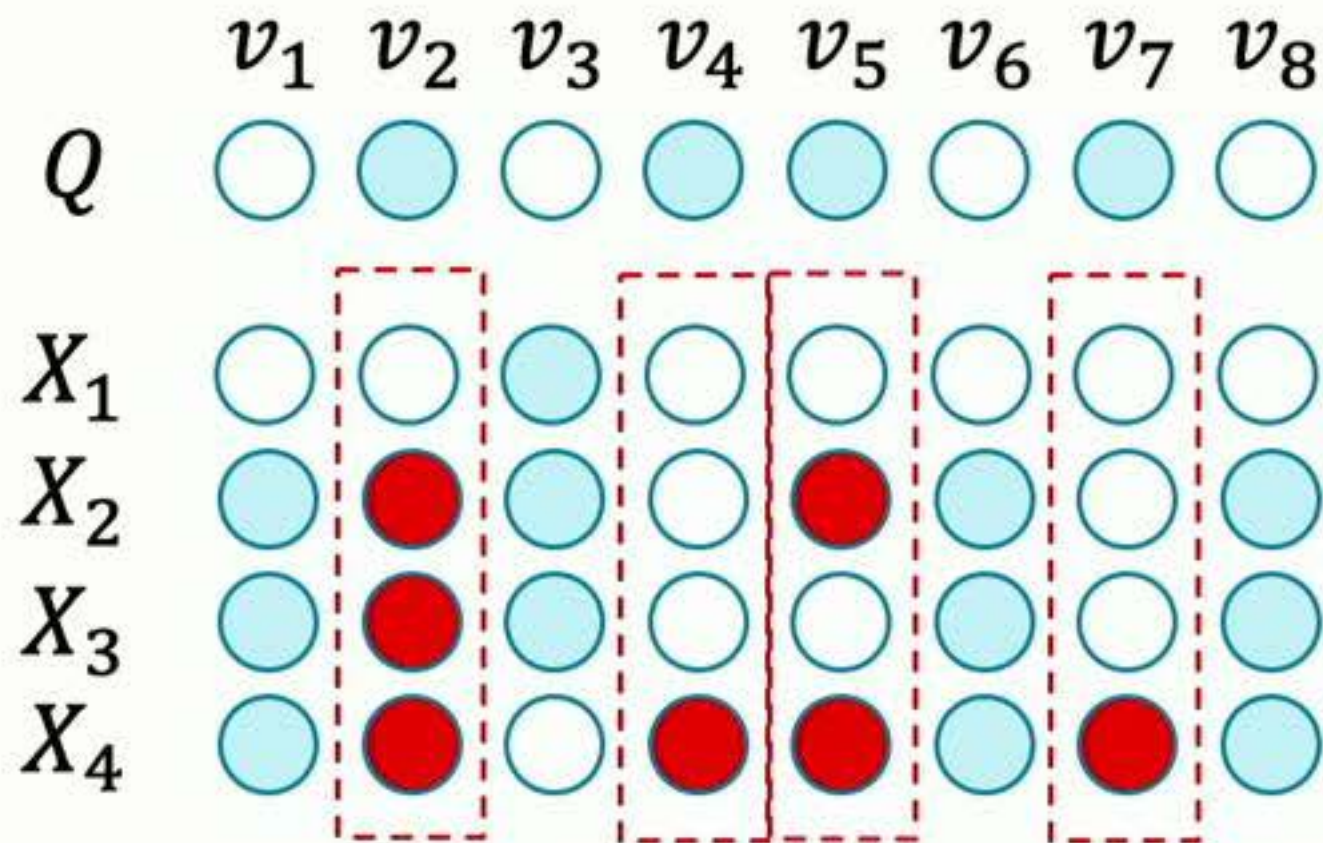


Prefix filtering<sup>2</sup>: reading candidate sets

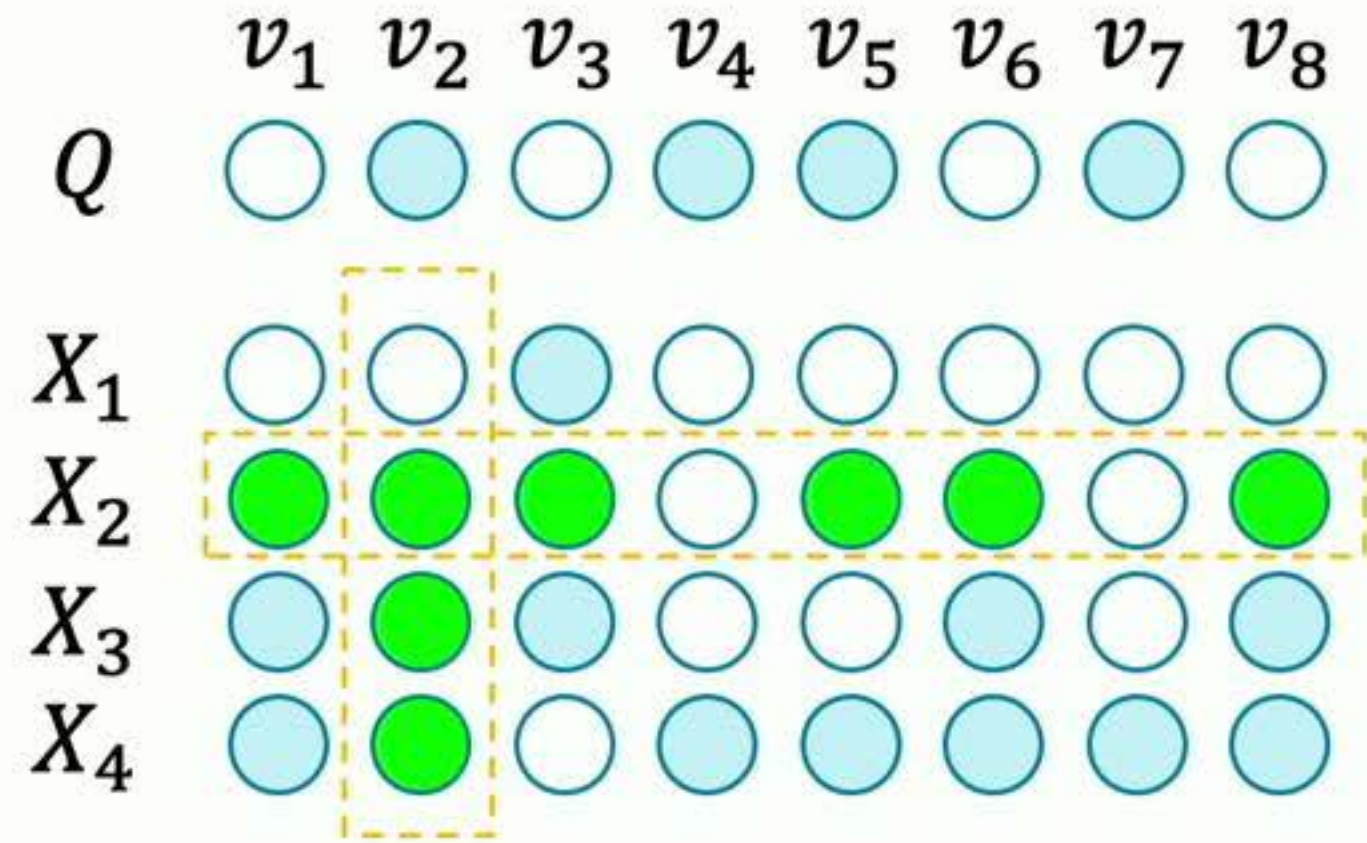
[1] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. Introduction to Information Retrieval. Cambridge University Press, 2008

[2] Chuan Xiao, Wei Wang, Xuemin Lin, and Haichuan Shang. Top-k set similarity joins. In ICDE, pages 916–927, 2009.

# BASELINES FOR FIND TOP-1



Posting list union<sup>1</sup>: reading all posting lists  
costs **7** values

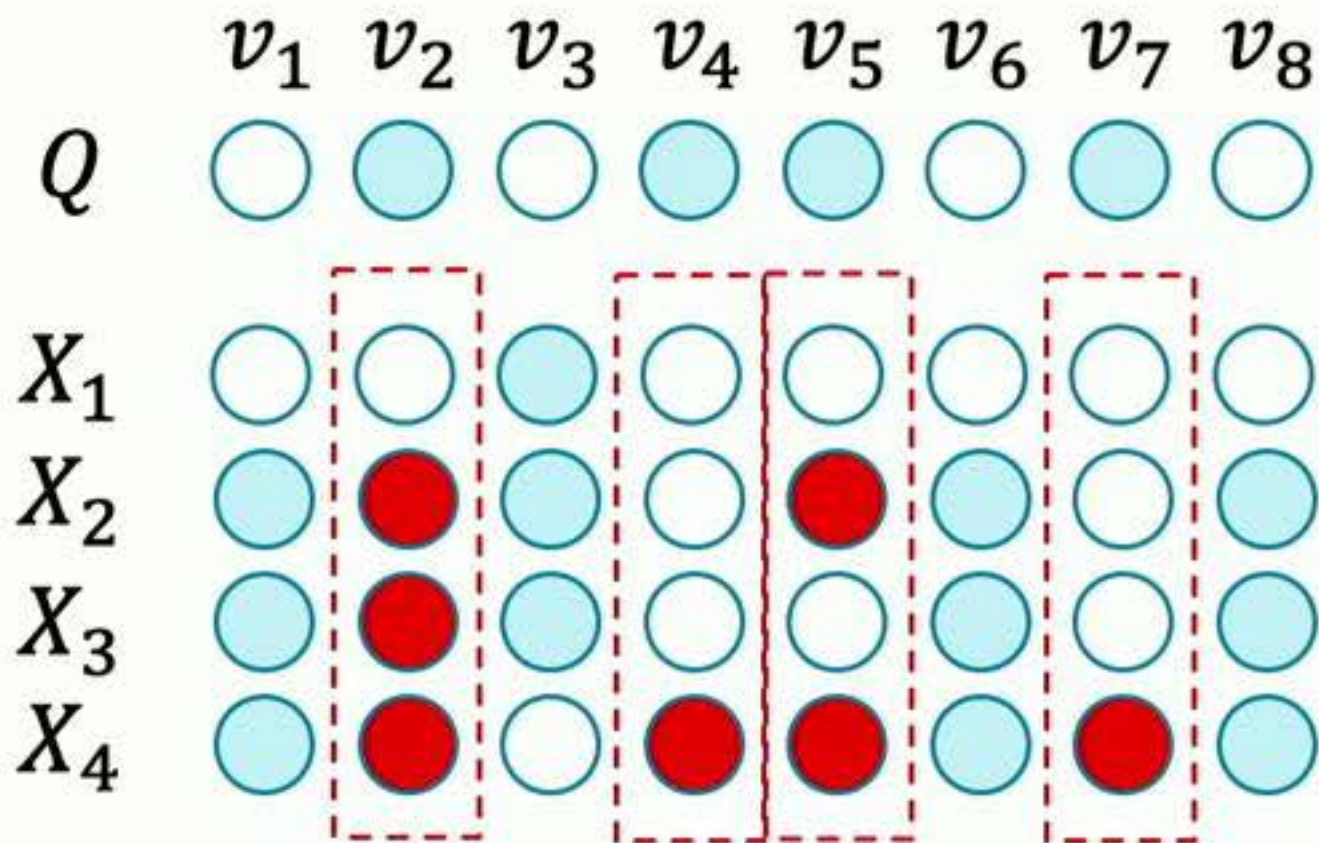


Prefix filtering<sup>2</sup>: reading candidate sets

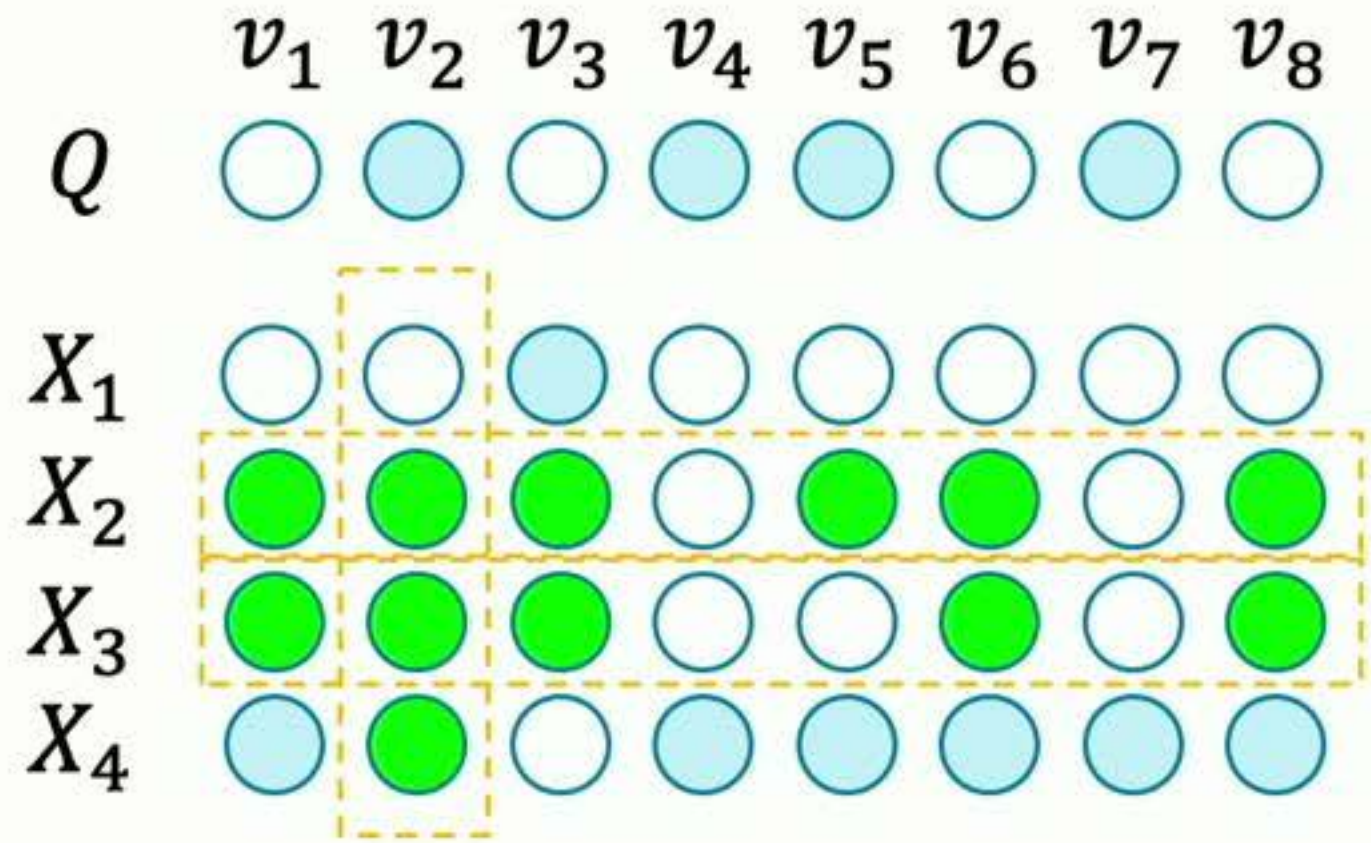
[1] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. Introduction to Information Retrieval. Cambridge University Press, 2008

[2] Chuan Xiao, Wei Wang, Xuemin Lin, and Haichuan Shang. Top-k set similarity joins. In ICDE, pages 916–927, 2009.

# BASELINES FOR FIND TOP-1



Posting list union<sup>1</sup>: reading all posting lists  
costs **7** values

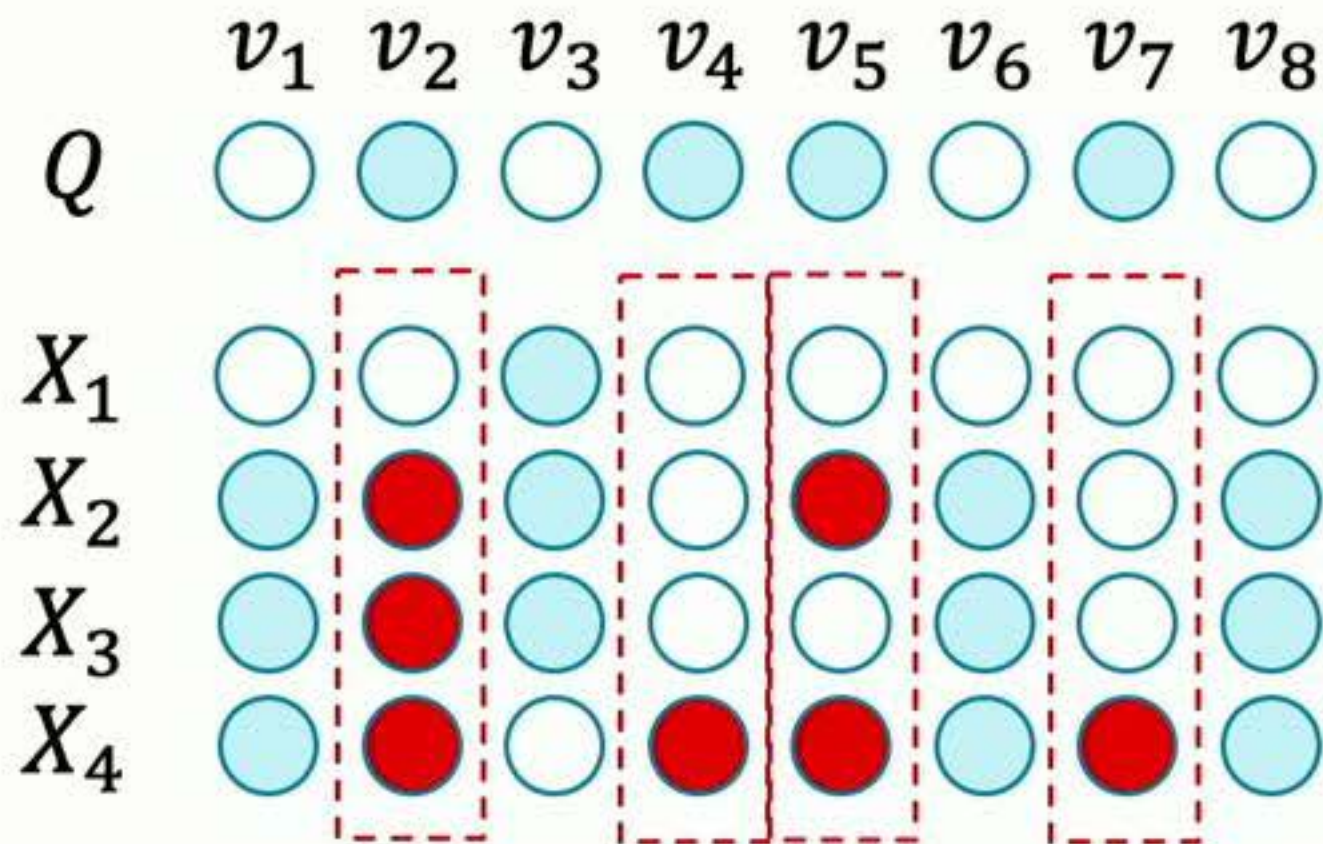


Prefix filtering<sup>2</sup>: reading candidate sets

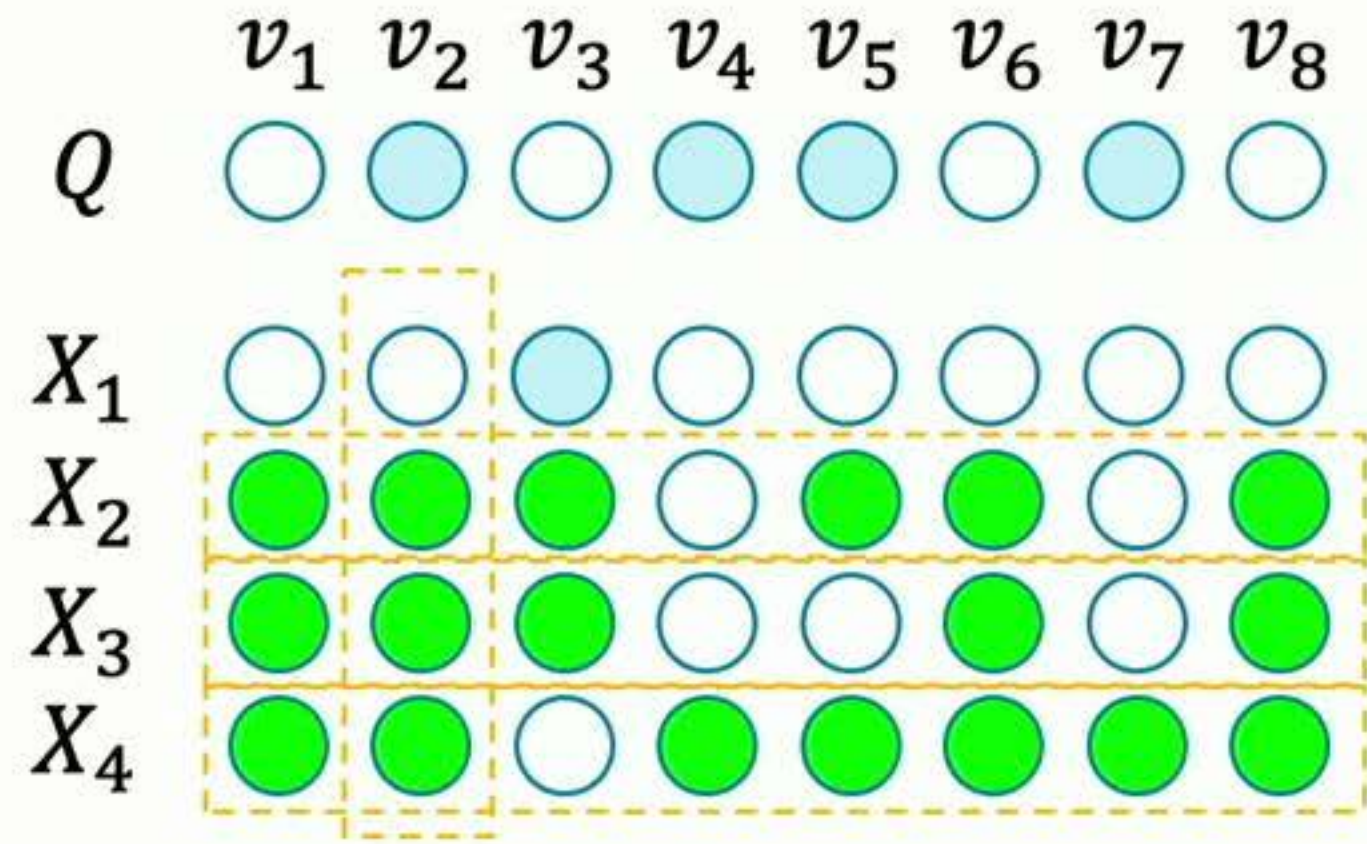
[1] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. Introduction to Information Retrieval. Cambridge University Press, 2008

[2] Chuan Xiao, Wei Wang, Xuemin Lin, and Haichuan Shang. Top-k set similarity joins. In ICDE, pages 916–927, 2009.

# BASELINES FOR FIND TOP-1



Posting list union<sup>1</sup>: reading all posting lists costs **7** values

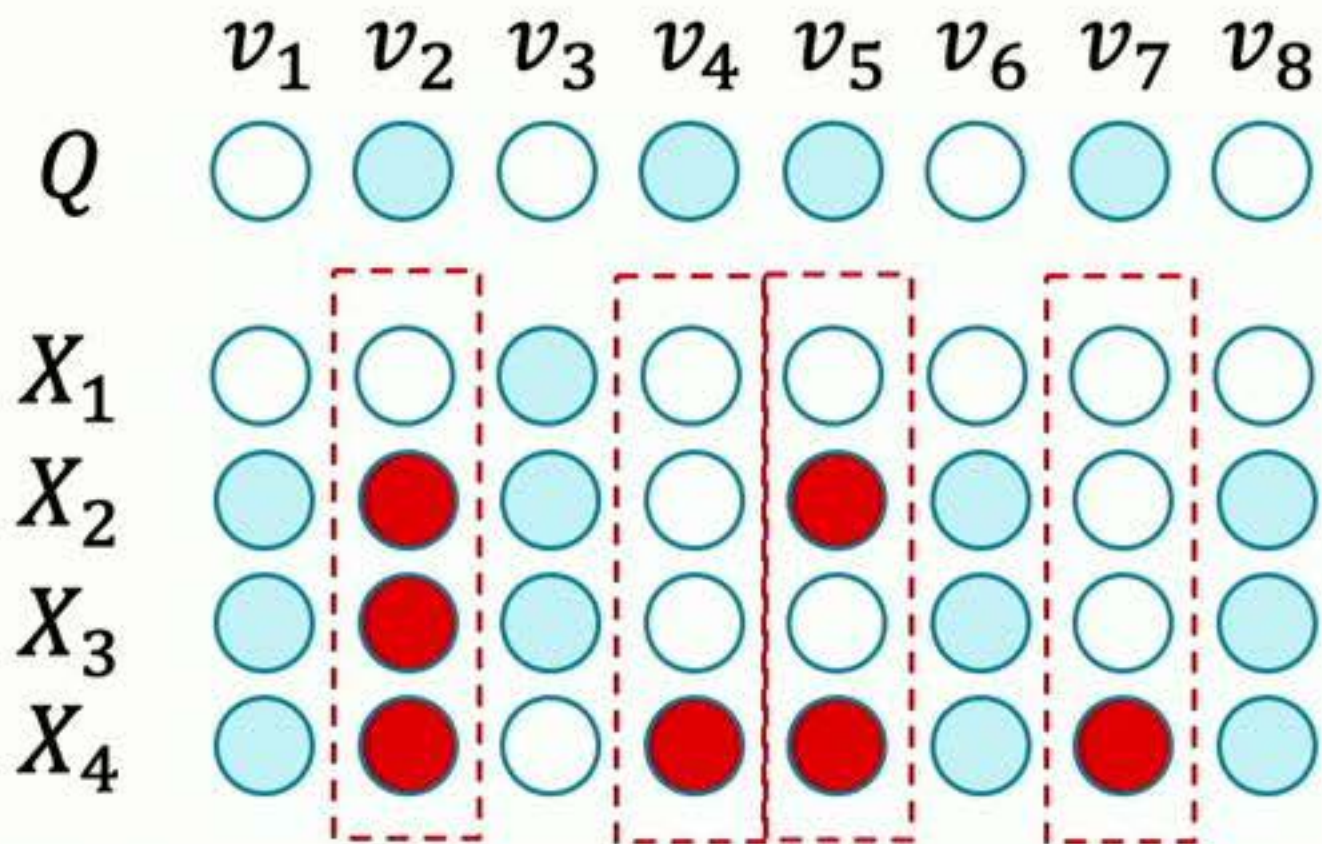


Prefix filtering<sup>2</sup>: reading candidate sets

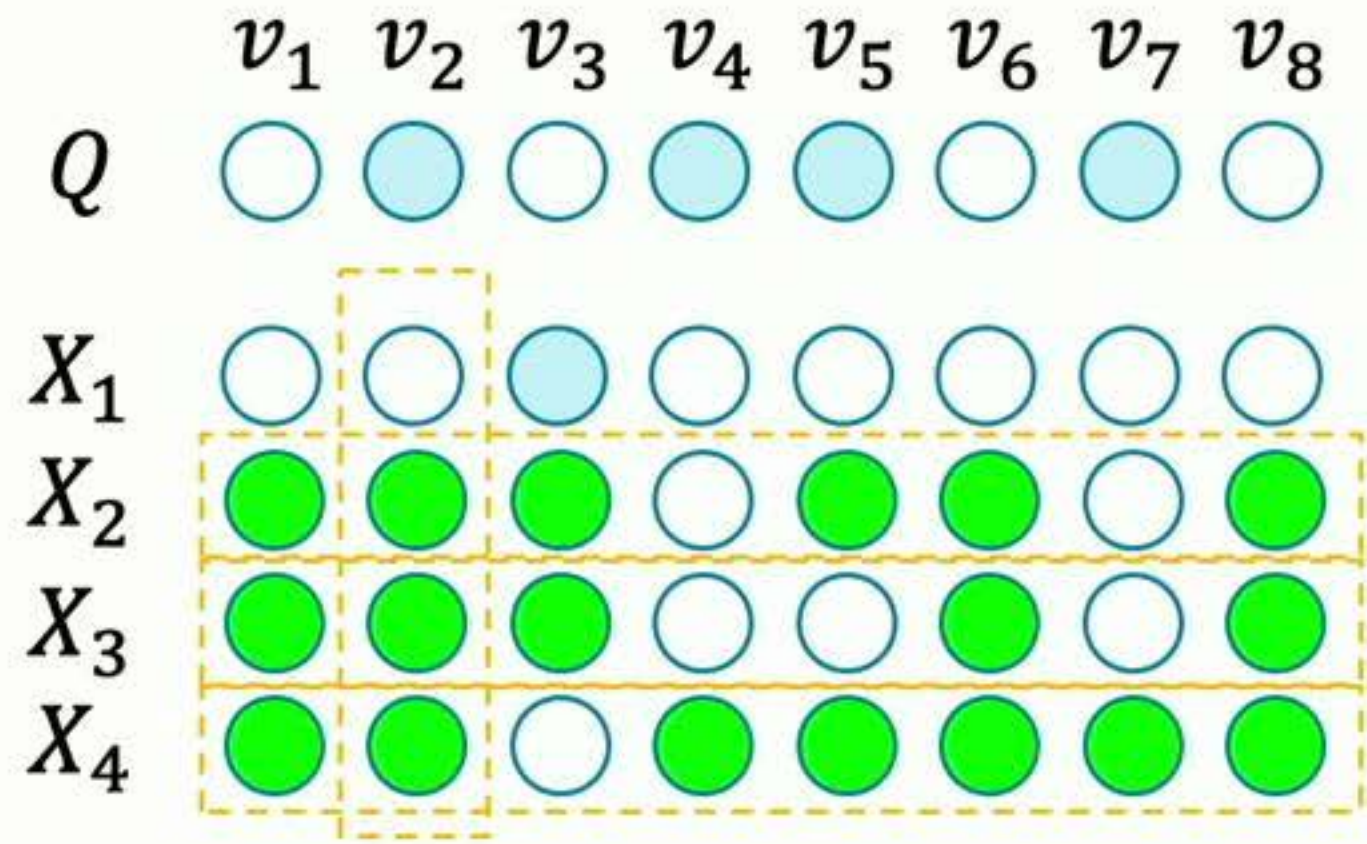
[1] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. Introduction to Information Retrieval. Cambridge University Press, 2008

[2] Chuan Xiao, Wei Wang, Xuemin Lin, and Haichuan Shang. Top-k set similarity joins. In ICDE, pages 916–927, 2009.

# BASELINES FOR FIND TOP-1



Posting list union<sup>1</sup>: reading all posting lists costs **7** values

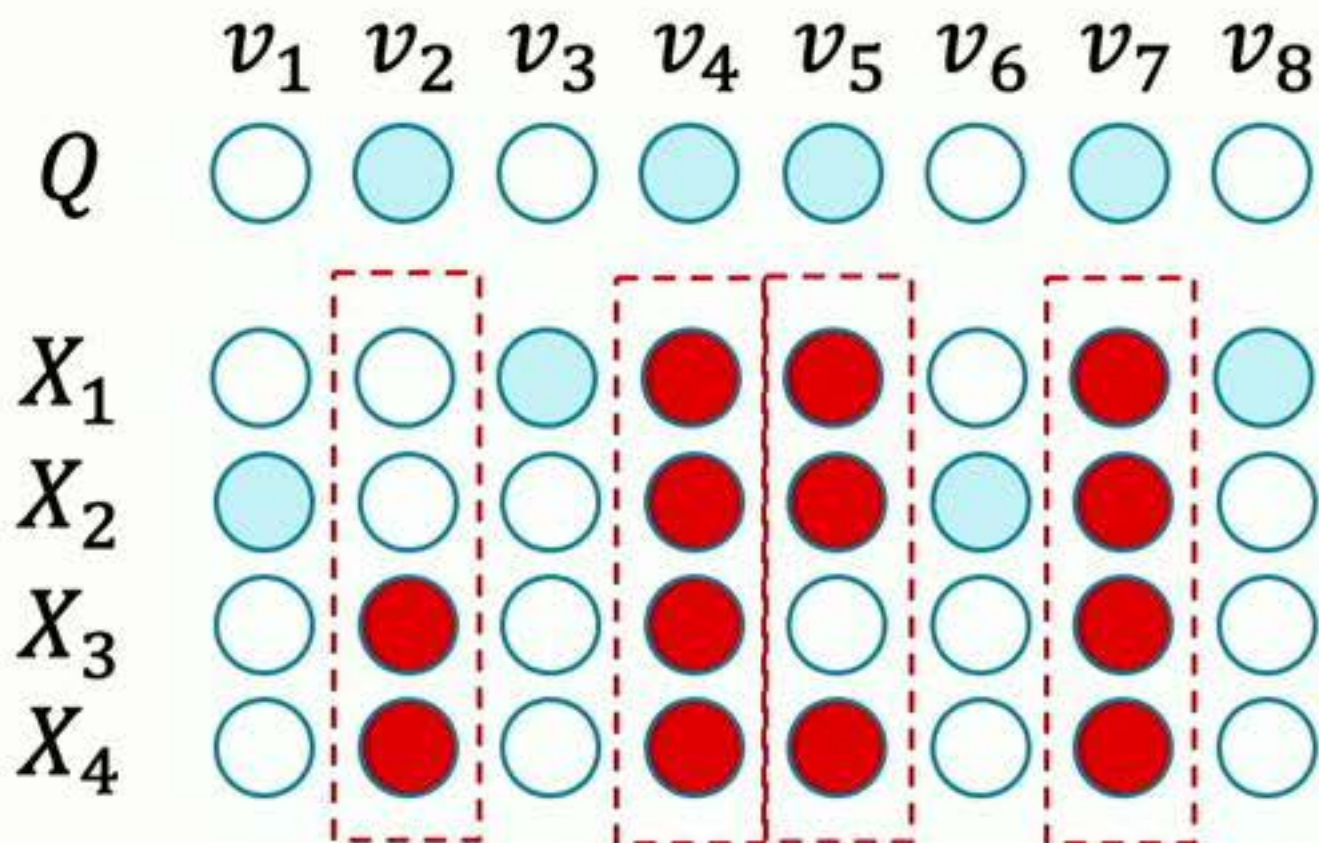


Prefix filtering<sup>2</sup>: reading candidate sets costs **18** values

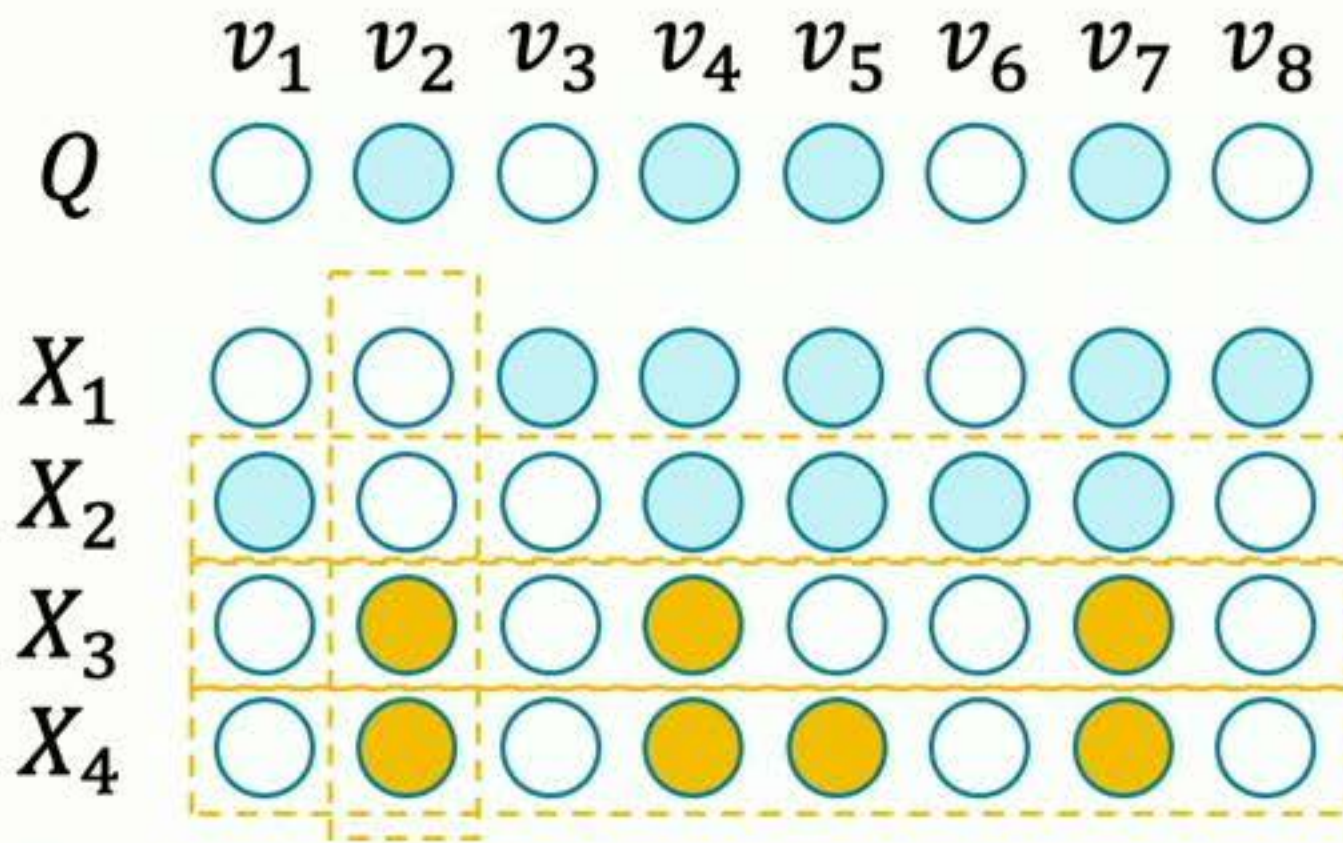
[1] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. Introduction to Information Retrieval. Cambridge University Press, 2008

[2] Chuan Xiao, Wei Wang, Xuemin Lin, and Haichuan Shang. Top-k set similarity joins. In ICDE, pages 916–927, 2009.

# BASELINES FOR FIND TOP-1



Posting list union<sup>1</sup>: reading all posting lists costs **13** values



Prefix filtering<sup>2</sup>: reading candidate sets costs **7** values

[1] Christopher D. Manning, Prabhakar Raghavan, and Hinrich Schütze. Introduction to Information Retrieval. Cambridge University Press, 2008

[2] Chuan Xiao, Wei Wang, Xuemin Lin, and Haichuan Shang. Top-k set similarity joins. In ICDE, pages 916–927, 2009.

# COST MATTERS IN DATA LAKES

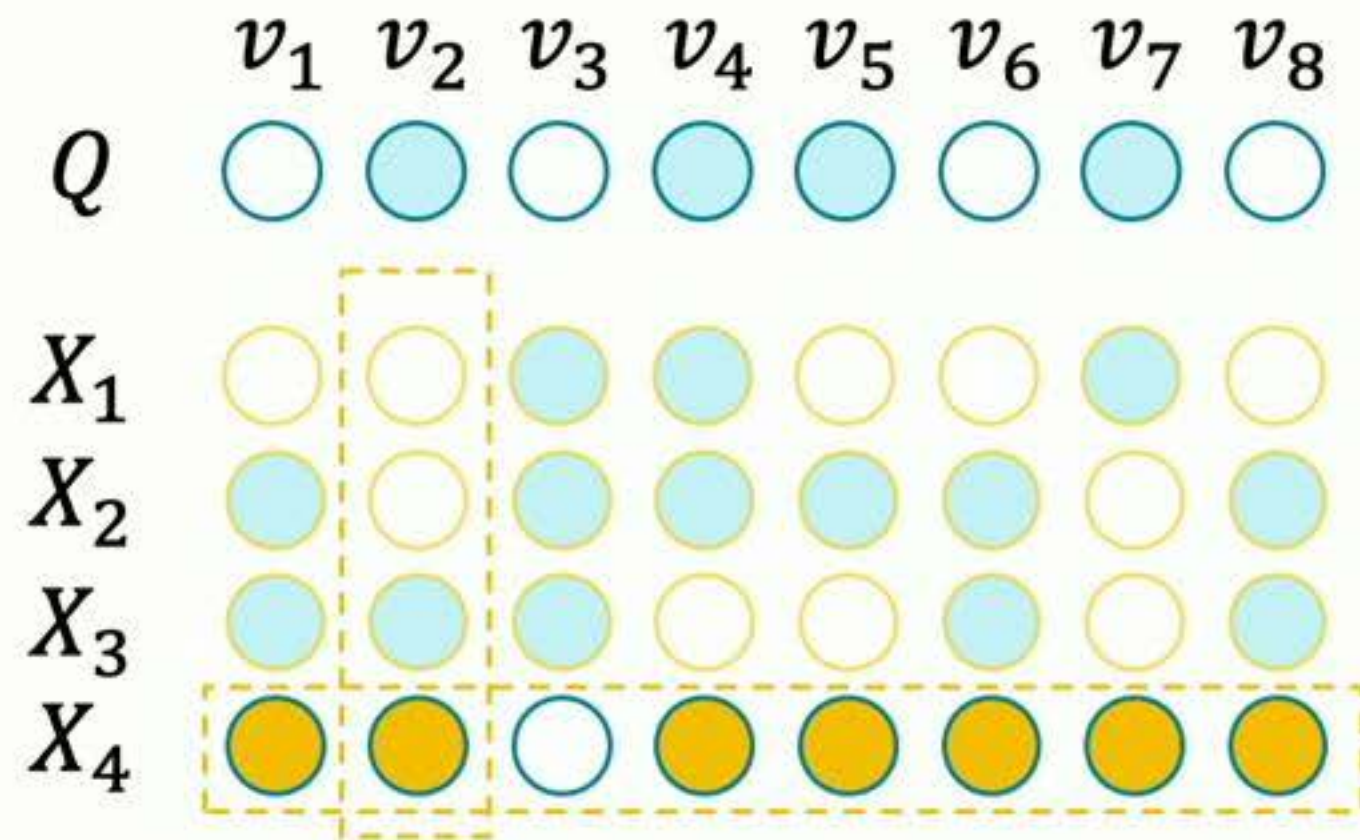
| Dataset               | # of Sets | Max Size | Avg. Size | # of Uniq. Values | Data Lakes |
|-----------------------|-----------|----------|-----------|-------------------|------------|
| Open Data*            | 745K      | 22M      | 1,540     | 562M              |            |
| WDC Web Table         | 163M      | 17K      | 10        | 184M              |            |
| AOL (Query Logs)      | 10M       | 245      | 3         | 3.9M              |            |
| ERON (Emails)         | 517K      | 3,162    | 135       | 1.1M              |            |
| DBLP (Bibliographies) | 100K      | 1,625    | 86        | 6,864             |            |

\*215,393 Open Data tables from Canadian, US, and UK Open Data Portal

*Read Bottleneck* – Large number of sets and data values makes index access expensive; large posting lists and sets are expensive to read

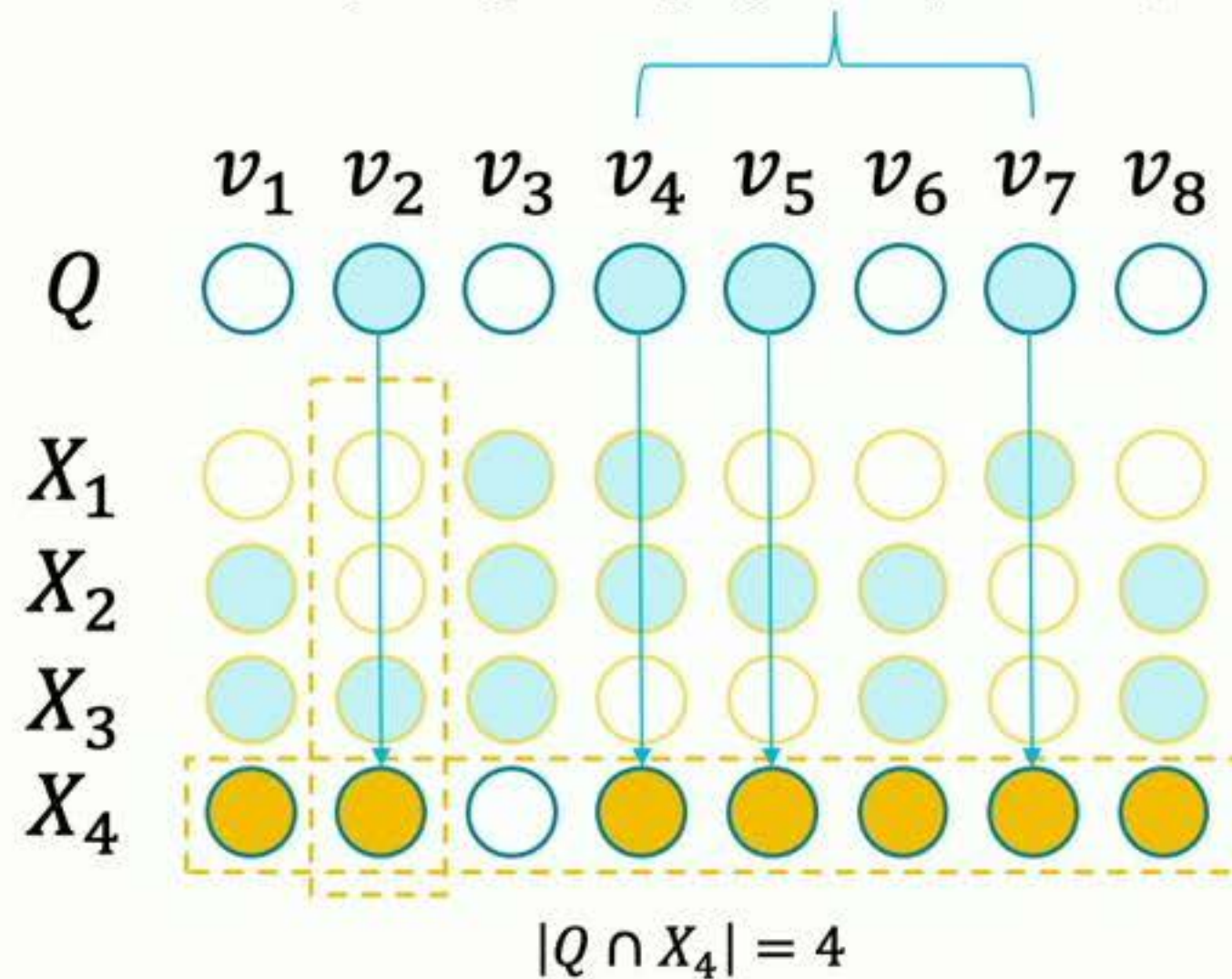
An efficient algorithm must reduce read cost

# READING SET ELIMINATES POSTING LISTS



# READING SET ELIMINATES POSTING LISTS

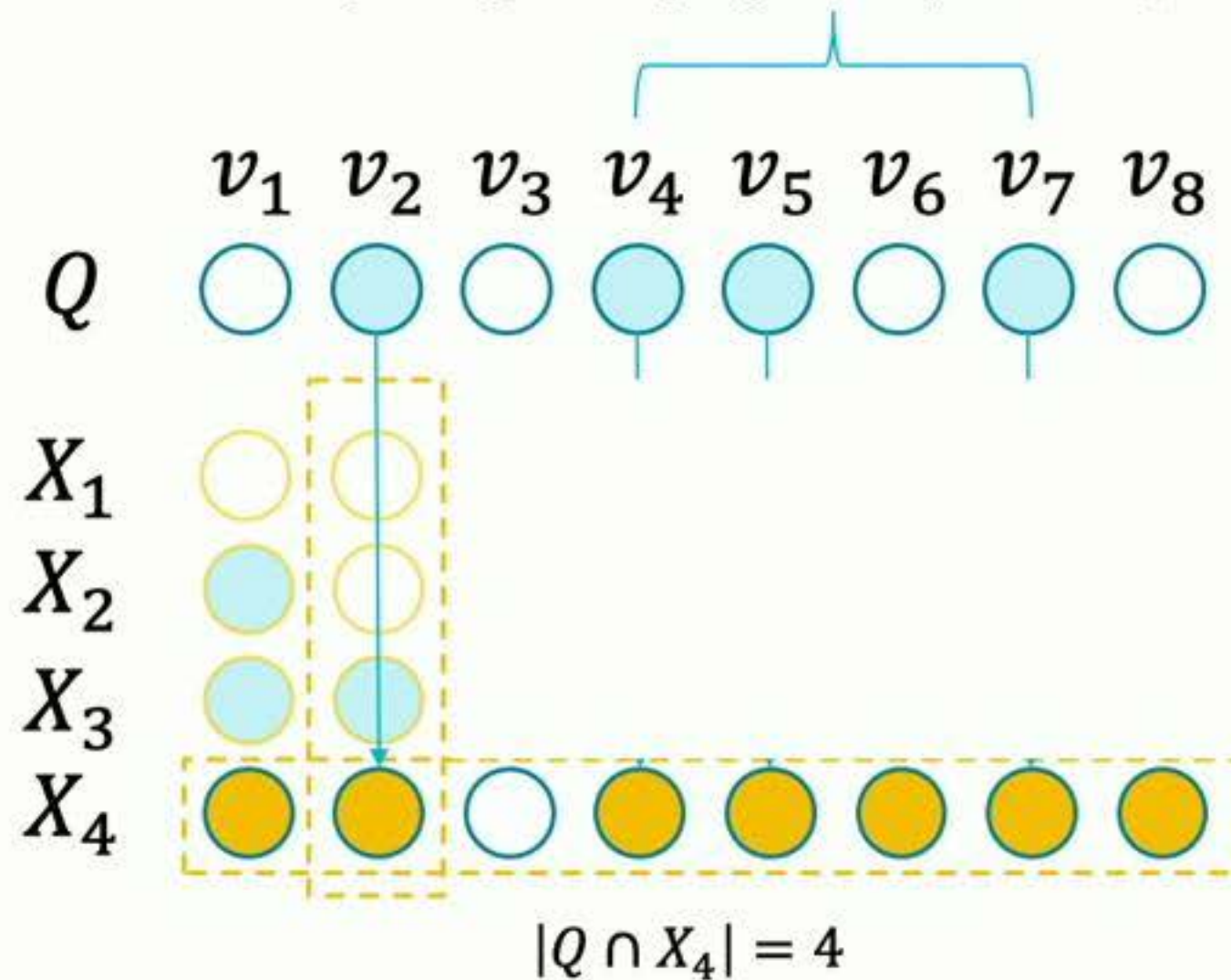
Maximum overlap from any unseen candidates  
in posting lists  $v_4, v_5$  and  $v_7$  is  $3 < 4$ ;



# READING SET ELIMINATES POSTING LISTS

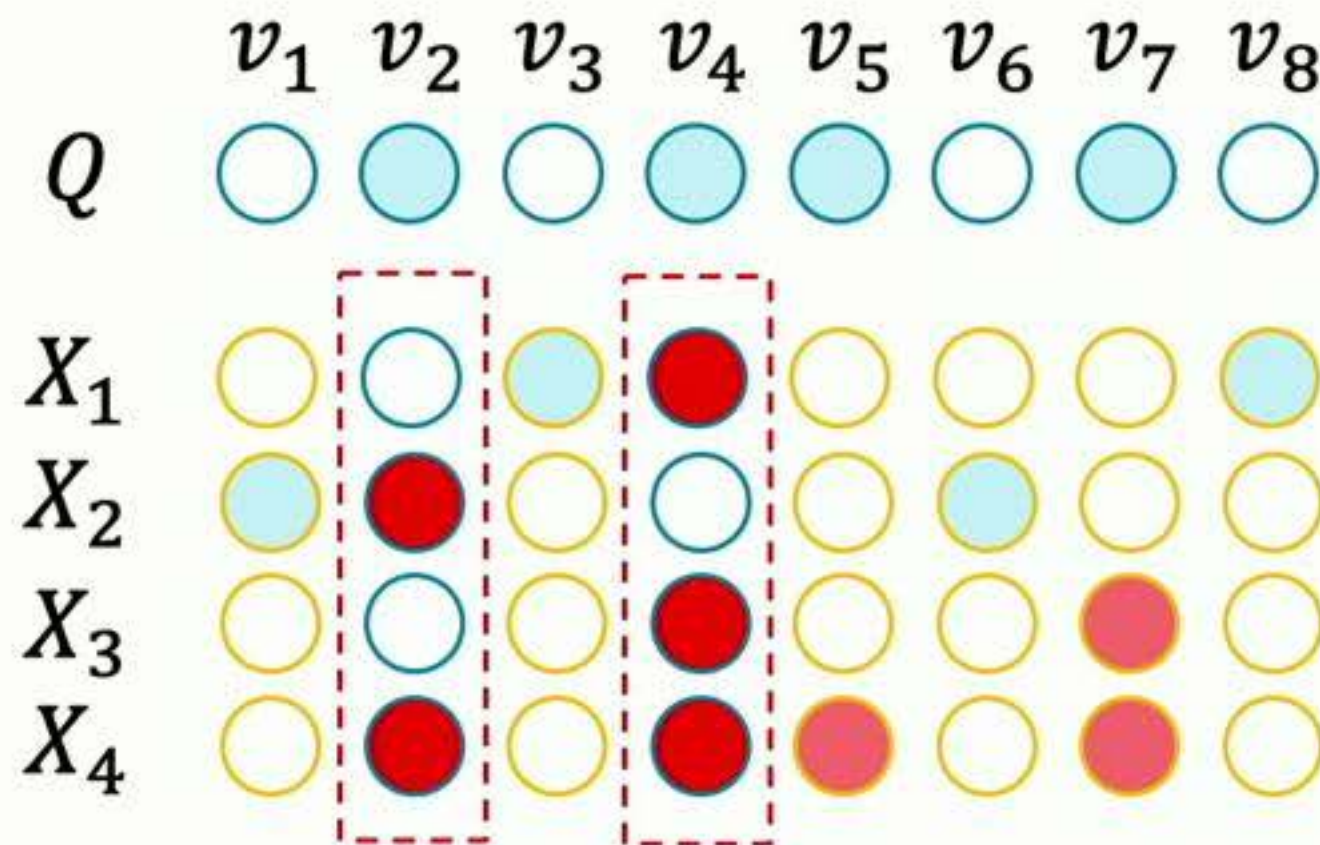
Maximum overlap from any unseen candidates  
in posting lists  $v_4, v_5$  and  $v_7$  is  $3 < 4$ ;

$X_4$  is the top-1; **eliminate**  
posting lists  $v_4, v_5$  and  $v_7$



# READING POSTING LISTS ELIMINATE SETS

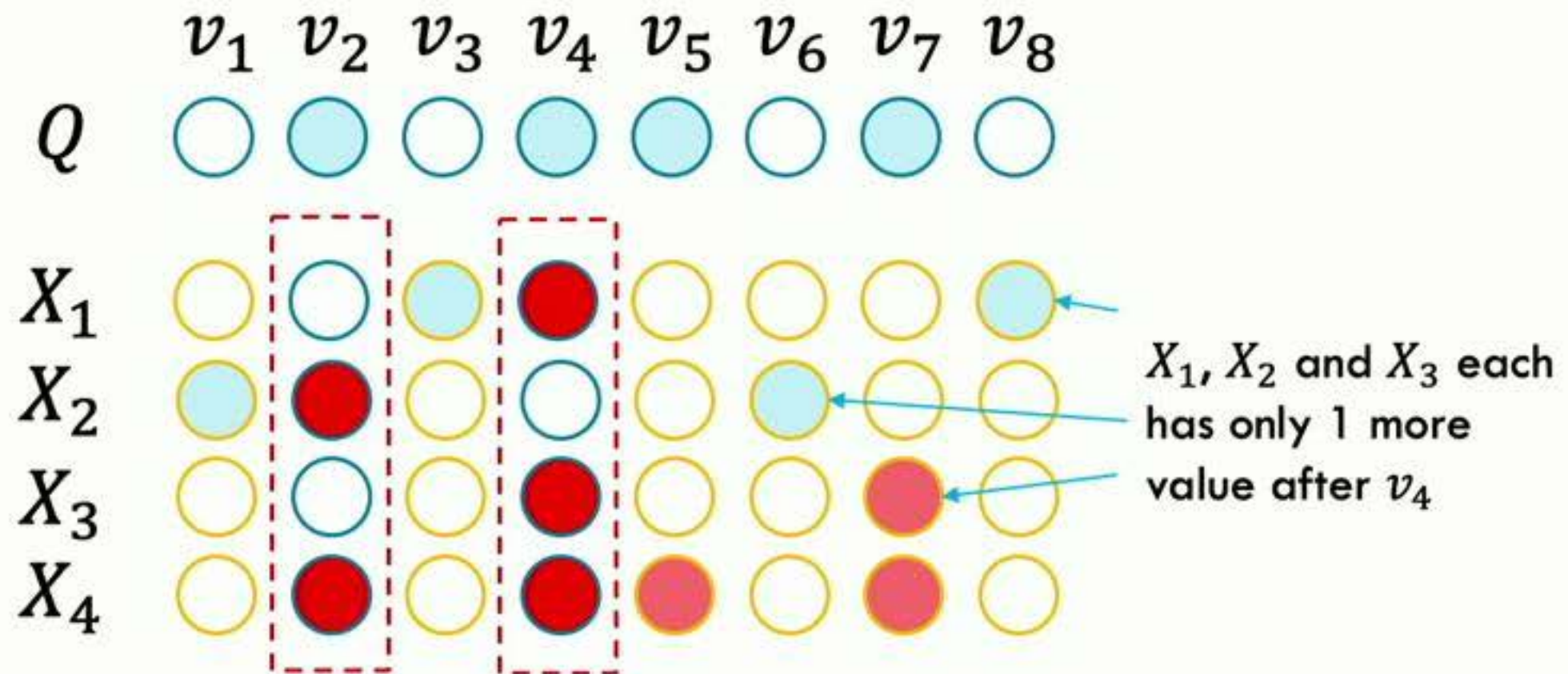
Current candidate sets are:  
 $\{X_1, X_2, X_3, X_4\}$



Assume the current top-1  
has 4 overlaps

# READING POSTING LISTS ELIMINATE SETS

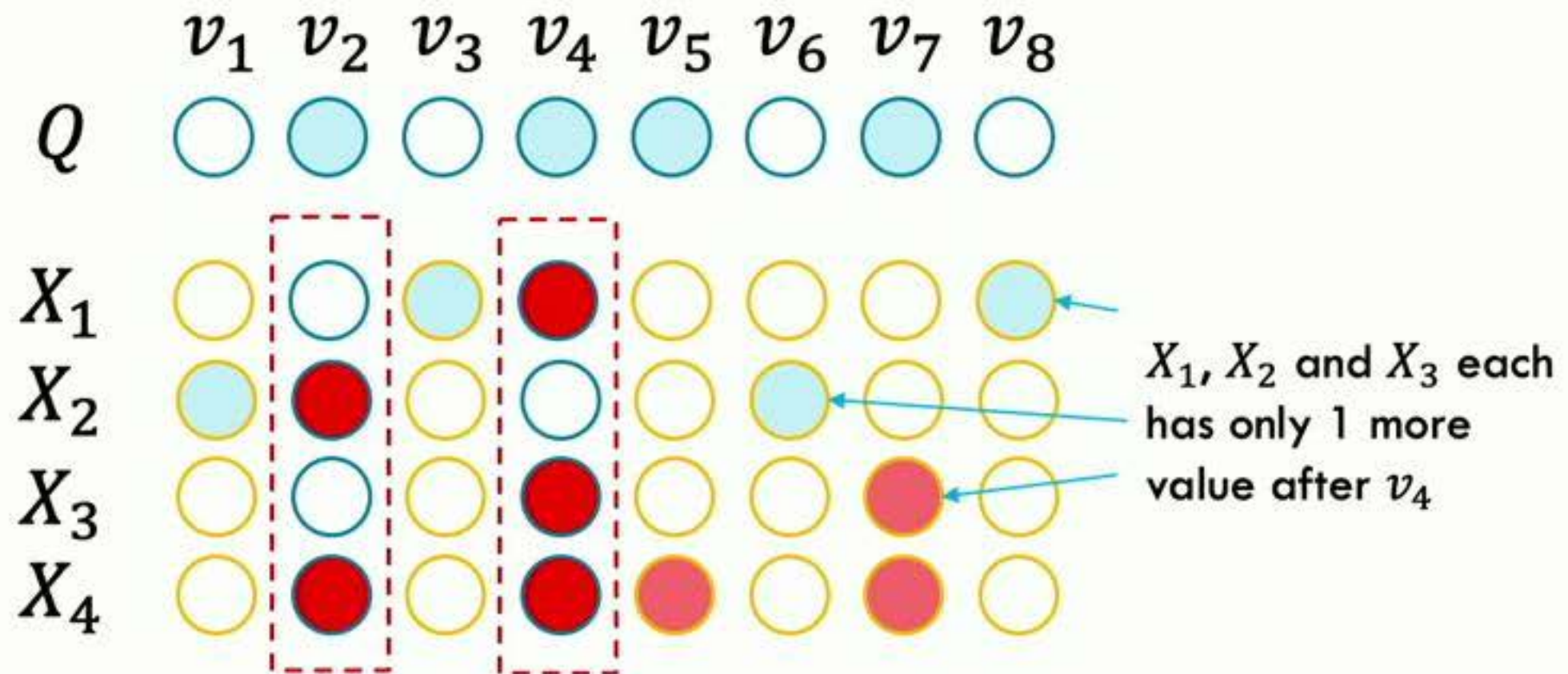
Current candidate sets are:  
 $\{X_1, X_2, X_3, X_4\}$



# READING POSTING LISTS ELIMINATE SETS

Current candidate sets are:  
 $\{X_1, X_2, X_3, X_4\}$

$X_1$ ,  $X_2$  and  $X_3$  can have at most 2 overlaps, less than the current top-1; **eliminate**

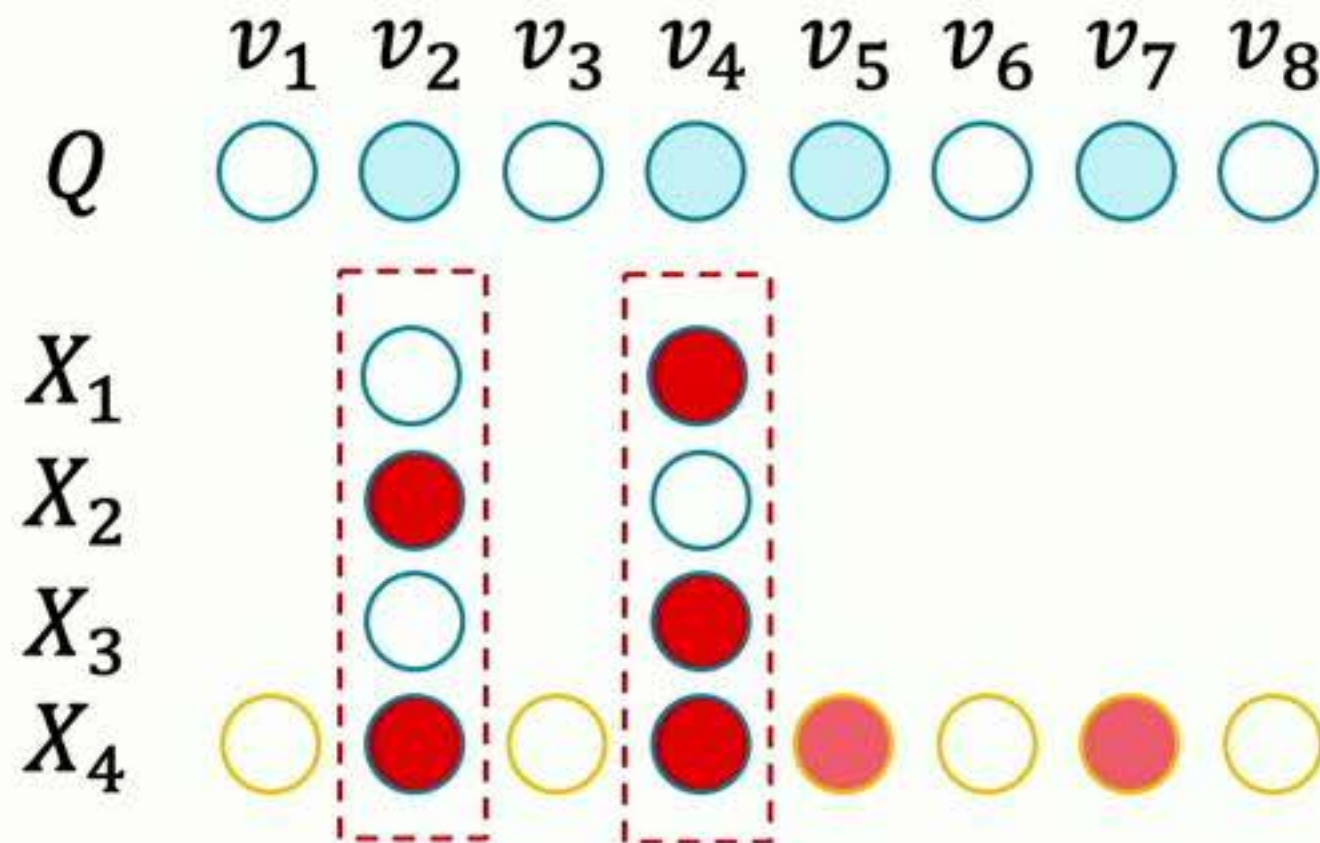


Assume the current top-1  
has 4 overlaps

# READING POSTING LISTS ELIMINATE SETS

Current candidate sets are:  
 $\{ \quad, X_4 \}$

$X_1, X_2$  and  $X_3$  can have at  
most 2 overlaps, less than  
the current top-1; **eliminate**



$X_1, X_2$  and  $X_3$  each  
has only 1 more  
value after  $v_4$

Assume the current top-1  
has 4 overlaps

# JOSIE - JOIN SEARCH VIA INTERSECTION EST.<sup>1</sup>

Measure “work” as total read cost of remaining posting lists and candidates

Measure “progress” as reduction in work

- Reading posting lists reveals intersecting values – reducing work in reading candidates
- Reading candidates improves prefix filter – reducing work in reading posting lists

Measure “price” as the read cost associated with progress

# JOSIE - JOIN SEARCH VIA INTERSECTION EST.<sup>1</sup>

Measure “work” as total read cost of remaining posting lists and candidates

Measure “progress” as reduction in work

- Reading posting lists reveals intersecting values – reducing work in reading candidates
- Reading candidates improves prefix filter – reducing work in reading posting lists

Measure “price” as the read cost associated with progress

```
while work > 0:  
    estimate progress  
    choose reading posting lists or candidates  
        that maximize (progress - price)
```

# JOSIE - JOIN SEARCH VIA INTERSECTION EST.<sup>1</sup>

Measure “work” as total read cost of remaining posting lists and candidates

Measure “progress” as reduction in work

- Reading posting lists reveals intersecting values – reducing work in reading candidates
- Reading candidates improves prefix filter – reducing work in reading posting lists

Measure “price” as the read cost associated with progress

```
while work > 0:  
    estimate progress  
    choose reading posting lists or candidates  
        that maximize (progress - price)
```

Cost model

# JOSIE - JOIN SEARCH VIA INTERSECTION EST.<sup>1</sup>

Measure “work” as total read cost of remaining posting lists and candidates

Measure “progress” as reduction in work

- Reading posting lists reveals intersecting values – reducing work in reading candidates
- Reading candidates improves prefix filter – reducing work in reading posting lists

Measure “price” as the read cost associated with progress

```
while work > 0:  
    estimate progress  
    choose reading posting lists or candidates  
        that maximize (progress - price)
```

Cost model

# JOSIE - PERFORMANCE

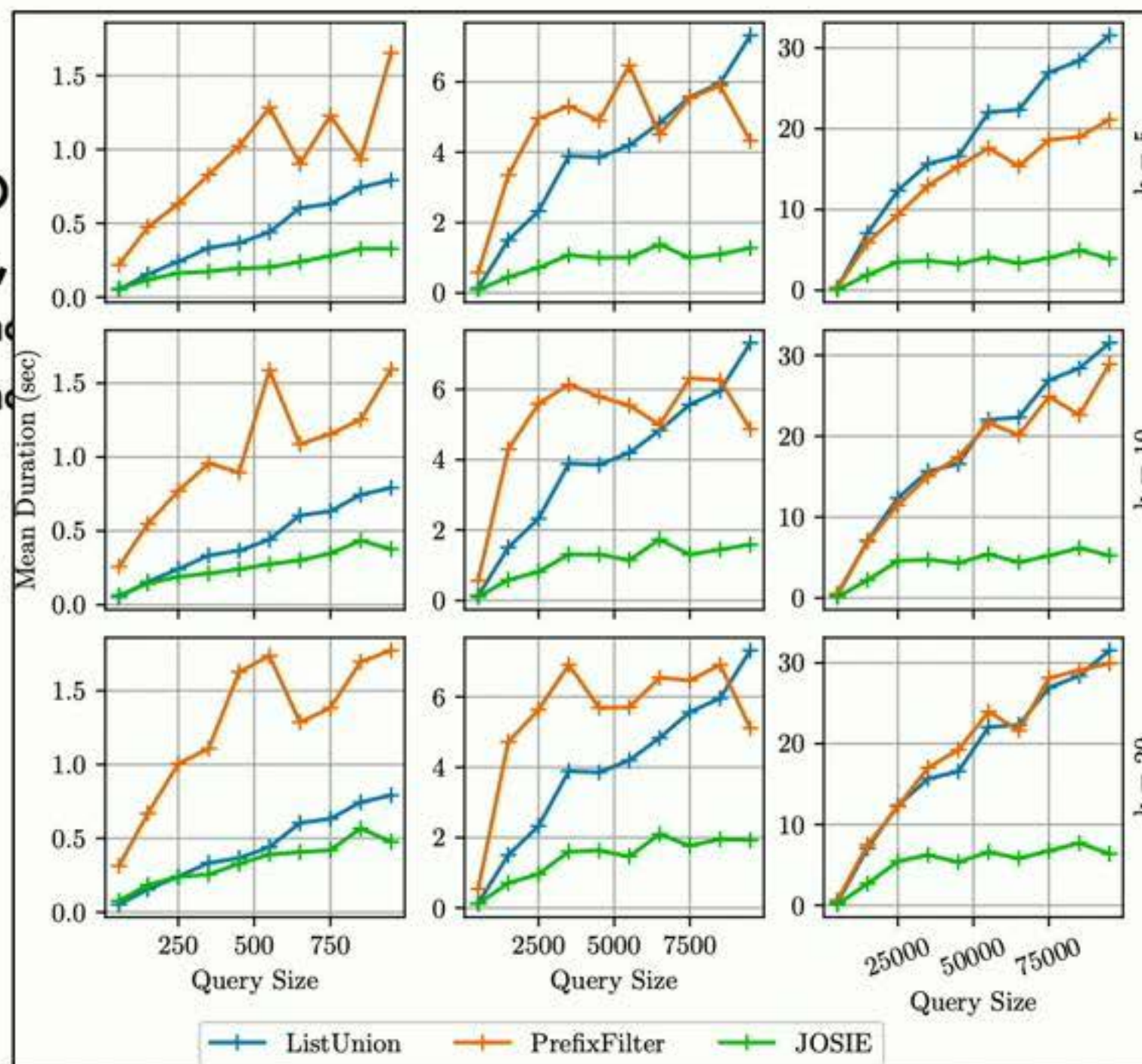
Evaluated on an Open Data lake (745K sets, 562M posting lists) and WDC Web Tables (163M sets, 184M posting lists)

- Up to 5× performance over baseline algorithms
- Up to 2× performance over LSH Ensemble (approximate), small queries ( $<5K$ ) and  $k < 20$

# JOSIE - PERFORMANCE

Evaluated on an Open Data Benchmark  
Tables (163M sets,

- Up to 5× performance
- Up to 2× performance



and WDC Web

and  $k < 20$

On Open Data  
Benchmark

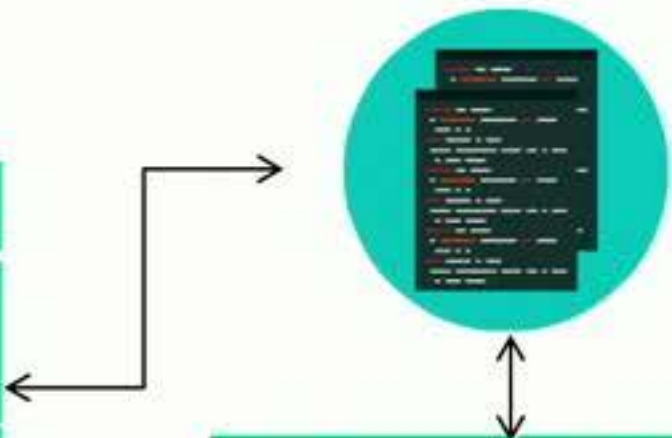
# OTHER WORK

## Auto-Join: Join Table by Leveraging Transformation (VLDB 17')

- Collaboration with Yeye He and Surajit Chaudhuri @ MSR

Learning transformation **without** user input

| Company       | Headquarter   | Ticker         |
|---------------|---------------|----------------|
| Alphabet Inc. | Mountain View | GOOGL (NASDAQ) |
| Apple Inc.    | Palo Alto     | AAPL (NASDAQ)  |
| Alibaba       | Hangzhou      | BABA (NYSE)    |



| Exchange | Ticker | Open     | Close    | Market Cap |
|----------|--------|----------|----------|------------|
| NASDAQ   | GOOGL  | 1,032.50 | 1,042.32 | 766 B      |
| NASDAQ   | AAPL   | 171.10   | 175.12   | 892 B      |
| NYSE     | BABA   | 179.37   | 174.13   | 427 B      |

# OTHER WORK

## Table Union Search on Open Data<sup>1</sup> (VLDB 17')

- Given a query table, find tables that can be “stacked” with the query table and aligning columns
- Top-K and Approximate

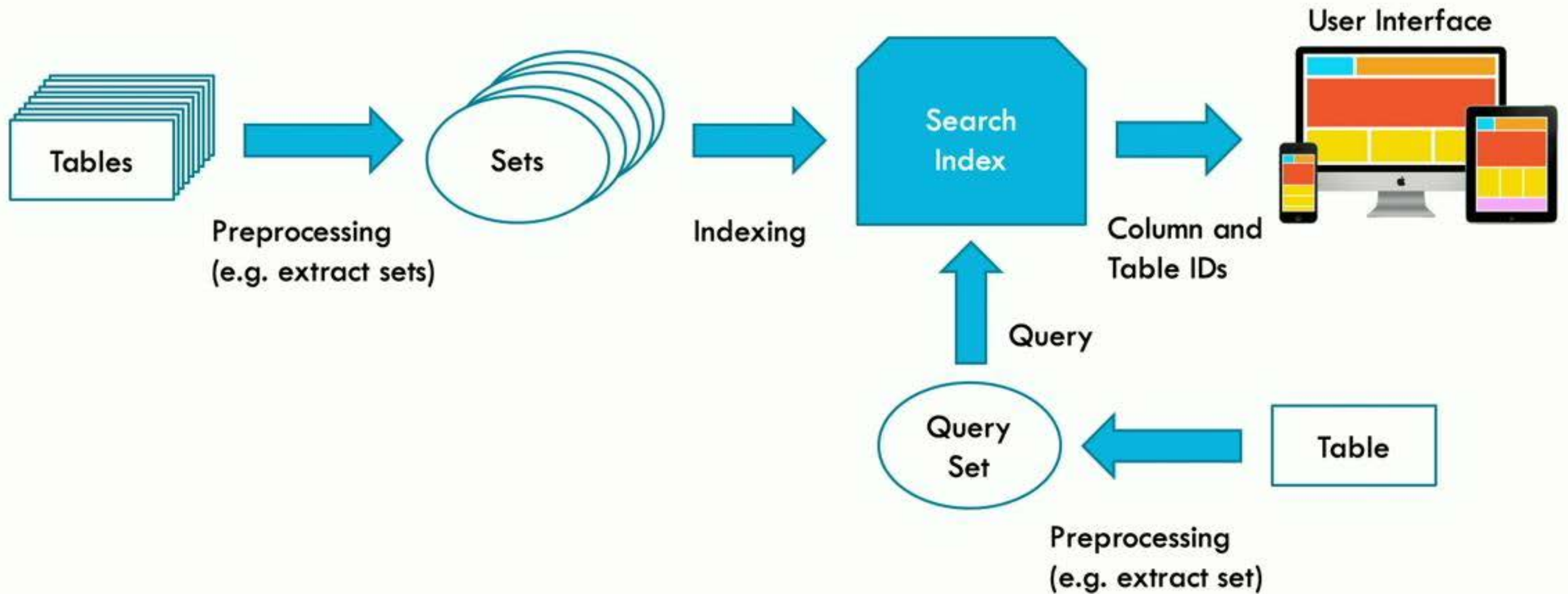
## Optimizing Organizations for Navigating Data Lakes<sup>2</sup> (VLDB 19' Revision)

- Generate a navigational DAG of topics for browsing the data lakes effectively
- Complementary to search

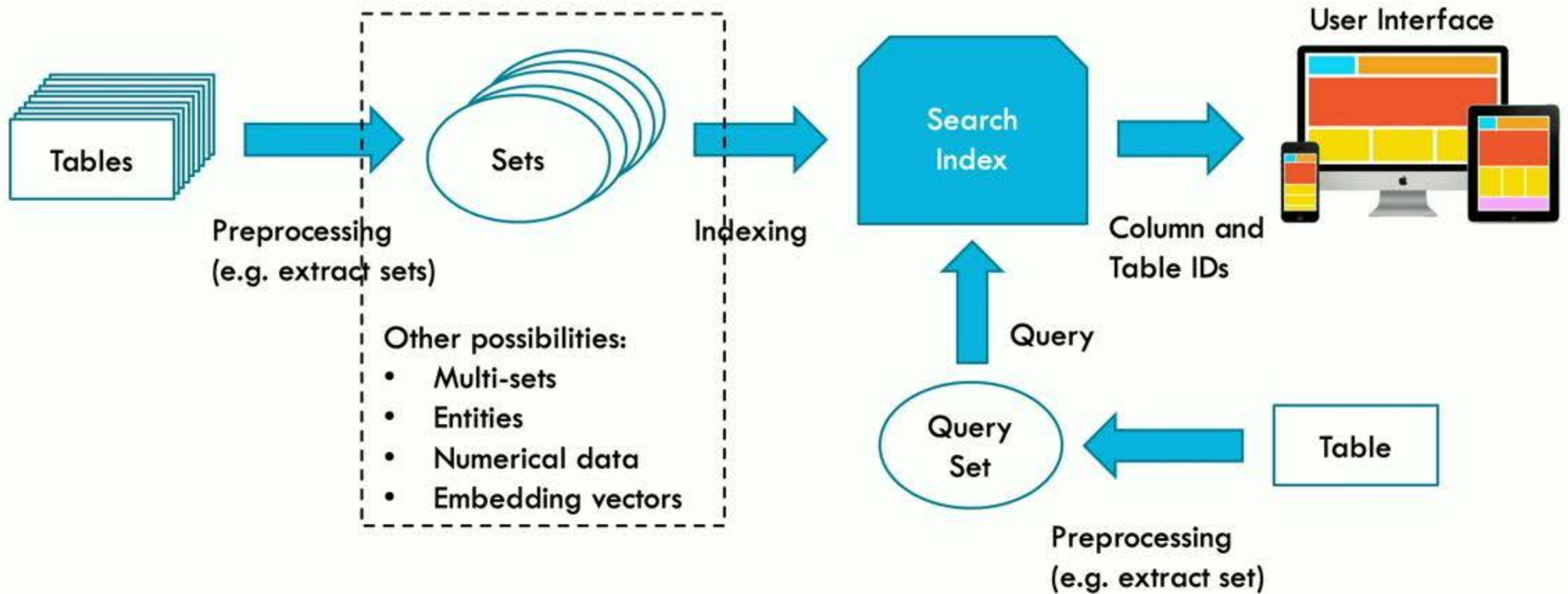
[1] Fatemeh Nargesian, **Erkang Zhu**, Ken Pu, Renée J. Miller, “Table Union Search on Open Data”, VLDB 2017

[2] Fatemeh Nargesian, Ken Pu, **Erkang Zhu**, Bahar Ghadiri Bashardoost, Renée J. Miller, “Optimizing Organizations for Navigating Data Lakes” VLDB 2019 (Revision), arXiv:1812.07024

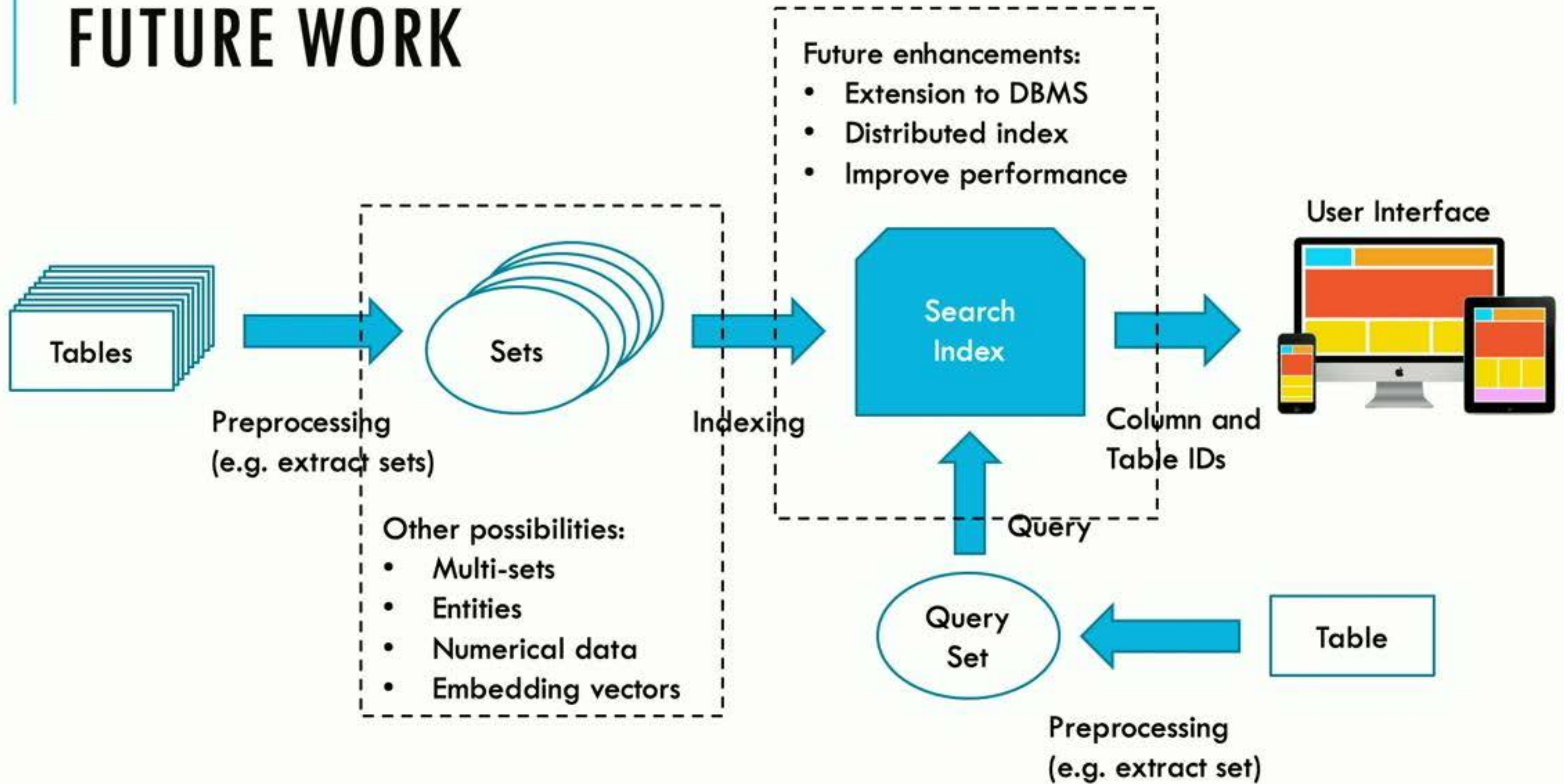
# FUTURE WORK



# FUTURE WORK



# FUTURE WORK



# FUTURE WORK

