

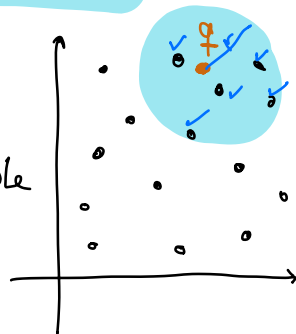
Similarity Search Queries

1. Rang queries

Given a query Point (vector) q

Find all the Points in the Vicinity of it

Table T:
Each black dot represents a tuple (vector) in Table T



Formally

$$\text{Find } \{t \in T \mid d(q, t) \leq r\}$$

SQL Extension

```
SELECT *
FROM T
WHERE d(q) ≤ r
```

e.g., pgvector standard for ℓ_2 -norm

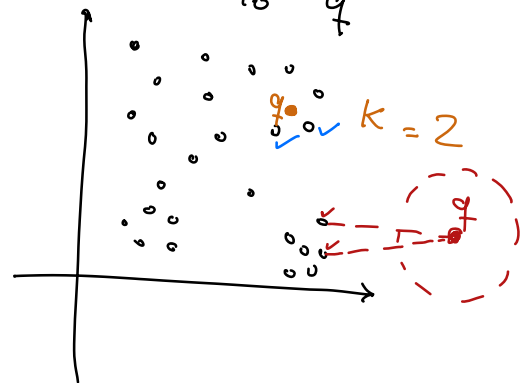
```
Select *
from T
where
```

Embedding $\langle - \rangle$ $[3, 5, 8, -1] \langle - \rangle 5$
Column Name d q r

2 - nearest Neighbor queries

Given a query point q , and a value K

Find the K nearest Points in the Table to q



Formally

Find

$$K\text{-argmin}_{t \in T} d(t, q)$$

SQL Extension

```
SELECT *
FROM T
ORDER BY d(q)
LIMIT k
```

In, pgvector, Euclidean (ℓ_2 -norm)

```
Select * from T
order by
```

Embedding $\langle - \rangle [-1, 3, -1, 2]$

Limit 10
 $\uparrow k$

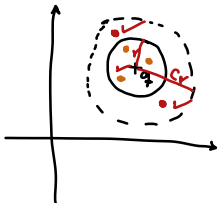
Efficient Processing of k-NN queries

Query Answering Schemes

- Exact: Given a value k , Find the k -nearest neighbor of Q
 $\rightarrow S = \underset{x \in T}{k\text{-argmin}} d(q, x)$
- Approximate: Returns a Set S' which may not be the exact kNN, but is "close enough"

measuring Approximation

① by distance
 $\forall x' \in S' \quad dist(q, x') \leq C \cdot dist(x, q)$



② by recall
 Recall@K: what Percentage of the returned Points belong to the kNN Set

$$Recall@K = \frac{|S \cap S'|}{K}$$

Question: How to Find kNN exactly?

Step 1:

Compute $d(q, x)$ for all Points in the Table

$$\Theta(nd)$$

→ number of columns
→ number of rows

Step 2:

apply the extended version of the median Finding Alg.

on the Scores to Find the top-k (k-NN)

$$\Theta(n)$$

⇒ Total time Complexity: $\Theta(nd)$
 X

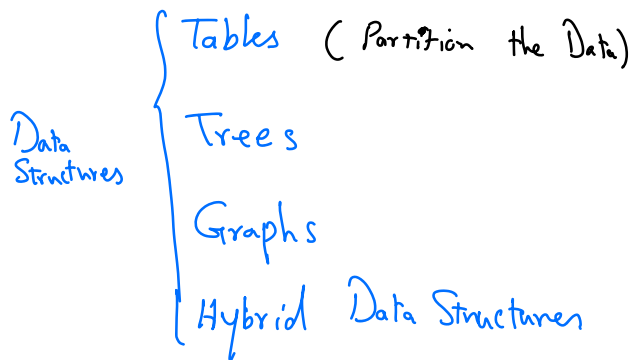
- This is very slow, particularly when n is large

⇒ Question: Can we preprocess the database & Create Indices that enable Sublinear query answering at least approximately?

Indexing the Database for fast ANN query answering

- To Partition The Data
 ↳ How?

- Build Some Data Structures on top of the Data

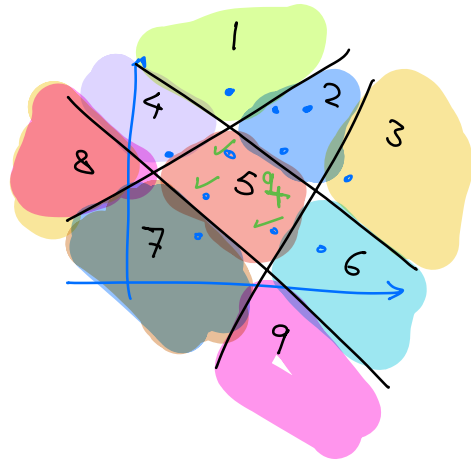
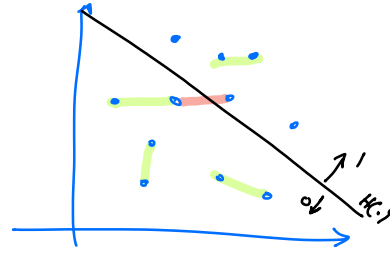


Local Sensitive Hashing (LSH)

A hash function H is in class of LSH, if

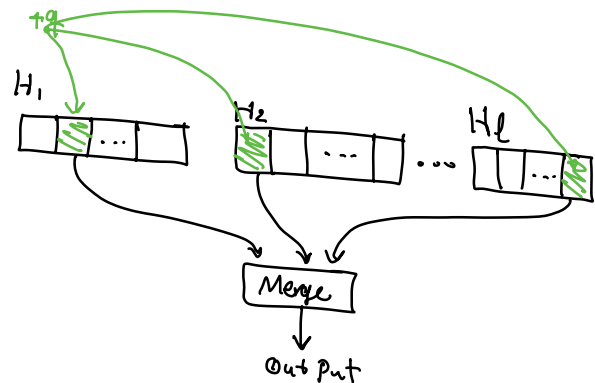
$$\begin{aligned} \checkmark a, b, \text{dist}(a, b) \leq r \\ P(H(a) = H(b)) \geq P_1 \\ \checkmark a, b, \text{dist}(a, b) \geq Cr \\ P(H(a) = H(b)) < P_2 \\ P_1 > P_2 \end{aligned}$$

e.g.,



Using One Hash function, it is likely that we miss Some of the NNs

* Resolution: Use multiple Hash Functions



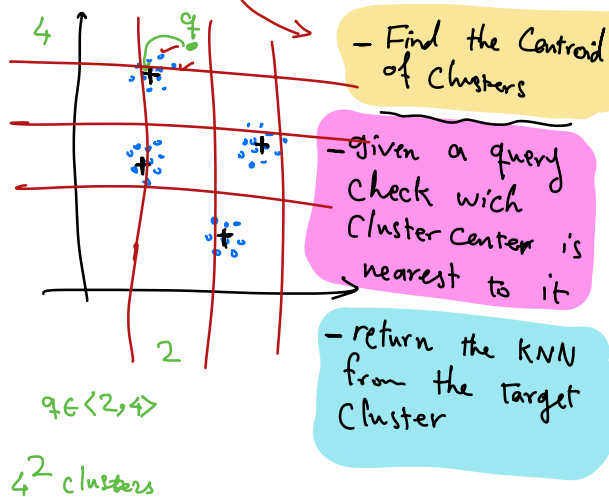
This is called LSH

* When Data Points are not evenly distributed in the Space (when many points are in close vicinities), LSH loses its power.

→ LSH does not Prune the Search Space in these Cases.

Resolutions

- Learned Hash Functions
- Quantization (K-means clustering)

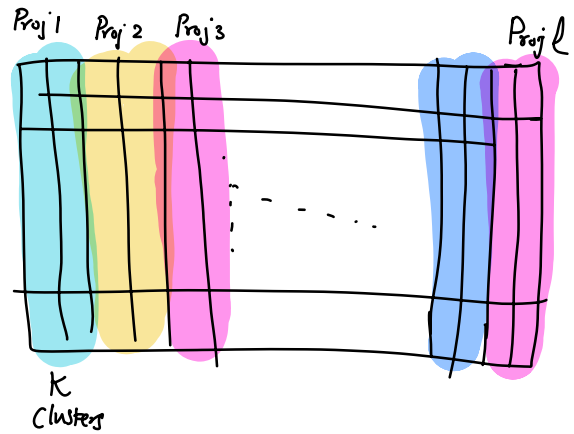


* Issue 1: K should be in $O(\frac{n}{\log n})$ for this approach to have $O(\log n)$ query time

* Issue 2: Curse of Dimensionality

Resolution:

Projection on lower Dimensions



$$\Rightarrow \text{No. Clusters} = K^l$$

$$K^l = \frac{n}{\log n}$$

$$\Rightarrow l = \log_K n - \log_K \log_K n$$

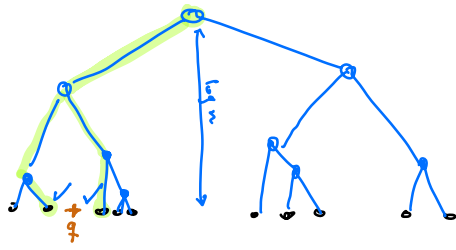
→ Given a query Point:

- Step 1: to find the Cluster Centers that q falls into

$$q \in \{c_1, c_2, \dots, c_l\}$$

c_i : is the closest Cluster Center to q , On projection i

- Step 2: check the Points within the same Cluster & return the top-k



To Construct the k-MV Tree

Construct Tree(S, i)

if $|S| \leq \tau$ ← Size Threshold
return S

Rule ← Select Rule(S, i)

← The rule for Partitioning S

Left ← $\{x \in S \mid \text{Rule}(x) = \text{false}\}$

Right ← $\{x \in S \mid \text{Rule}(x) = \text{true}\}$

return (rule, Construct Tree(Left, $i-1$),
Construct Tree(Right, $i-1$))

① Round Robin on Columns

Select Rule_{RR}(S, i)

j ← median of S based on
Column($i \bmod d$)

Rule ← $(x[i] \leq j)$

e.g. S

	1	2	3	4
t_1	3	...		
t_2	-1	...		
t_3	5	..		
t_4	8	..		
t_5	0	..		

$i=1$

S

$\{-1, 0, 3\}$

$\{t_2, t_5, t_1\}$

$\{5, 8\}$

$\{t_3, t_4\}$

Observation: This is not a Binary Search Tree!

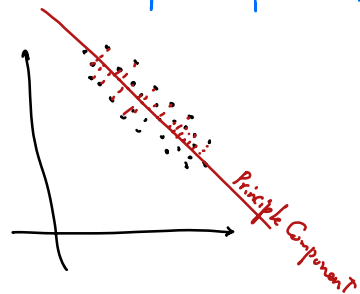
⇒ We Initially need to check both Left & Right Branches to a leaf of

→ This can become Linear $\Theta(nD)$

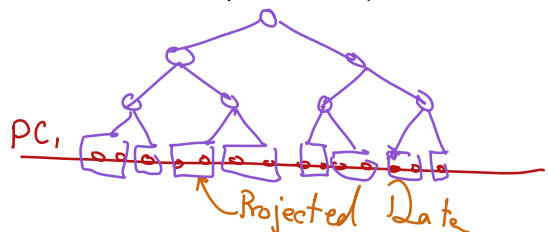
A Practical Resolution

- Use PCA (Principle Component Analysis)

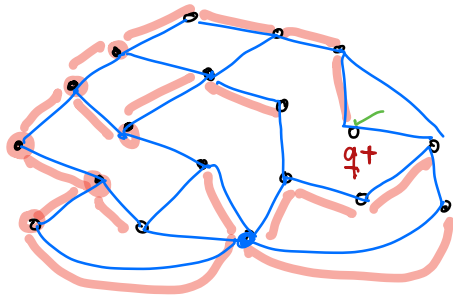
e.g.



- Instead of Creating the tree based indiv. Columns, Construct the Tree on Projection(s) of Data on Principle Components



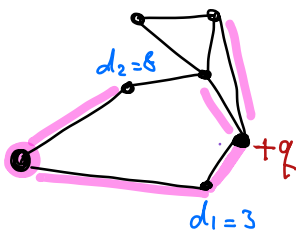
Graph Indices



Approach 1:

- Start from an arbitrary point in the Graph
- follow a BFS Traversal & find the NN of q .

$$\Theta(nd) \times$$



Approach 2: Greedy Search (DFS)

- at each iteration,
Next $\leftarrow$ the neighbor with
Smallest distance to q
- Continue the Search
on Next
- Stop when all neighbors
have larger distances

Time Complexity of Greedy Search

$$\Theta(L \deg(u))$$

$\rightarrow$ length of the DFS path
 $\rightarrow$ degree of the nodes of the

Properties of the Graph

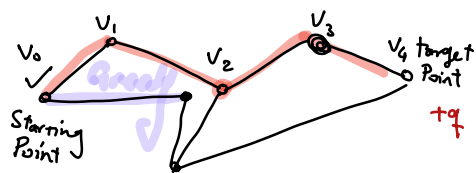
1- G should be Sparse
 $\deg(u) \sim \sim$ Small

2- L should be Small.
after a Small number
of hops, the alg. Stops

3- when the Greedy
Stops a close-to-NN
point is found

How about Connecting
each Point to
its K nearest
neighbors?

Monotonicity of the Search Path.



There should exist a monotonically decreasing
Path on $d(X, q)$ from
the Starting Point to $nn(q)$

$$\forall i \leq l \quad d(v_i, q) > d(v_{i+1}, q)$$

* Monotonic Search Network

A Graph, where the Greedy Path Satisfies the monotonicity Property.

→ if d is small ($d \leq 3$)
Such a graph with low
degrees exist

→ Use Voronoi Diagram

→ what if d is not small?

Navigable Small-world Graph

L should be small $\Theta(\log n)$

NSW

— In Practice

KNNG:

- Iteratively add the nodes to the Graph.
- for each node connect it to its k -nearest neighbors

HNSW (Hierarchical NSW)

