



Retrieval and Clustering

Applied Machine Learning
Derek Hoiem



Previously

- K-nearest neighbor algorithm
- Application to classification and regression
- Basics of datasets and generalization

Today's lecture

- Retrieval
 - “Brute force”
 - Faiss library
 - Approximate: LSH
- Clustering
 - Kmeans
 - Hierarchical Kmeans
 - Agglomerative Clustering

Retrieval

- Given a new sample, find the closest sample(s) in a dataset
- Applications
 - Finding information (web search)
 - Prediction (e.g. nearest neighbor algorithm)
 - Clustering (kmeans)

“Brute force” search

- Compute distance between query and each dataset point and return closest point

Brute force search pseudo-code

getNearest(x_q, X)

```
dist_min = Inf
```

```
idx_min = -1
```

```
For each nth sample in X:
```

```
    dist = sum((X[n]-xq)**2) # sum square diff
```

```
    if dist < dist_min:
```

```
        dist_min = dist
```

```
        idx_min = n
```

```
return idx_min
```

FAISS library makes even brute force search very fast

- Multi-threading, BLAS libraries, SIMD vectorization, GPU implementations
- KNN for MNIST takes seconds
- Also contains many approximate search methods

```
import faiss                # make faiss available
index = faiss.IndexFlatL2(d) # build the index, d=size of vectors
# here we assume xb contains a n-by-d numpy matrix of type float32
index.add(xb)               # add vectors to the index
print index.ntotal
```

```
# xq is a n2-by-d matrix with query vectors
k = 4                        # we want 4 similar vectors
D, I = index.search(xq, k)  # actual search
print I
```

<https://engineering.fb.com/2017/03/29/data-infrastructure/faiss-a-library-for-efficient-similarity-search/>

Repos for fast search

- FAISS: <https://github.com/facebookresearch/faiss> (we use this for HW)
- ScaNN: <https://github.com/google-research/google-research/tree/master/scann>
- USearch: <https://github.com/unum-cloud/usearch>
- FAISS seems the gold standard, but ScaNN and USearch claim to be faster/lighter for some uses

Approximate search

- Many algorithms exist that are faster than brute force (especially for very large datasets) but don't guarantee returning exact nearest neighbors, e.g.
 - LSH – locality sensitive hashing
 - [HNSW](#) – hierarchical navigable small world graphs
 - IVF – inverted file (cluster-based)
- We'll discuss only LSH, but others are worth learning about if fast search is helpful for your application
- In practice, you should understand how different search methods work, but use an optimized library, not your own code

Locality Sensitive Hashing (LSH)

A fast approximate search method to return similar data points to query

Basic LSH process

1. Convert each data point into an array of bits, using the same conversion process/parameters for each
2. Map the arrays into buckets (e.g. with 10 bits, you have 2^{10} buckets)
 - Can use subsets of arrays to create multiple sets of buckets
3. On query, return points in the same bucket(s)
 - Can check additional buckets by flipping bits to find points within hash distances greater than 0

Random Projection LSH

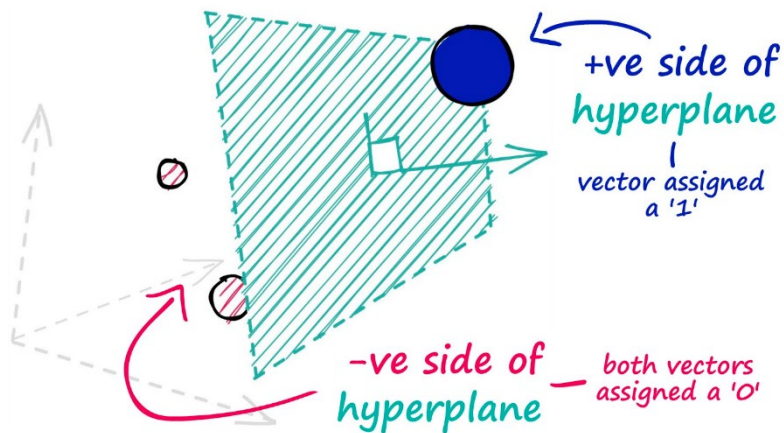
Data Preparation

Given dataset $X \in \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$ of $1 \times d$ datapoints:

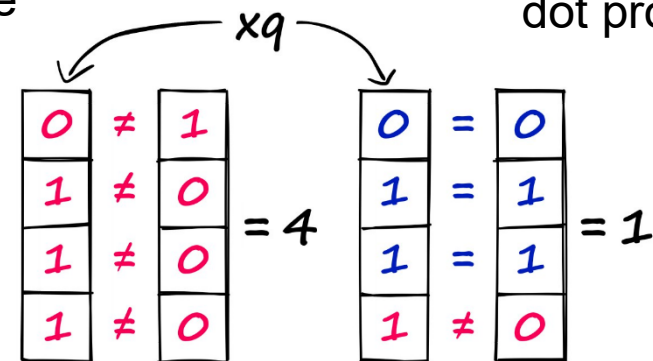
1. Center data on origin (subtract mean)
2. Create b random vectors $\mathbf{h}$ of length d `h = np.random.rand(d, nbits) - .5`
3. Convert each $\mathbf{x}_n$ to b bits: $\mathbf{x}_n \mathbf{h} > 0$

Query

1. Convert $\mathbf{x}_q$ to bits using $\mathbf{h}$
2. Check buckets based on bit vector and similar bit vectors to return most similar data points



Continuous vector is converted to bit vector, with each bit representing the datapoint's side of a hyperplane



Hamming distance == number of mismatches

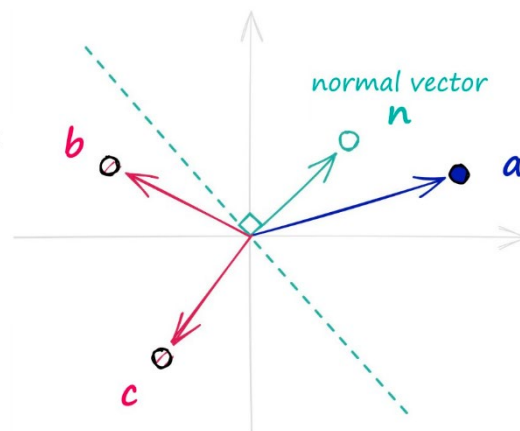
Point similarity is now a bit distance

dot-product

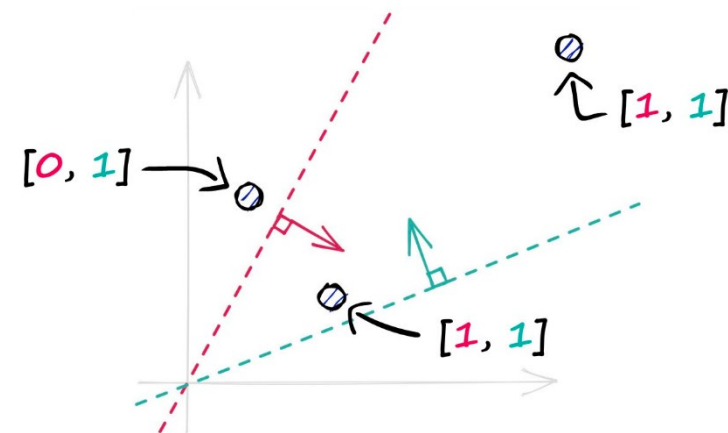
$$n \cdot a > 0$$

$$n \cdot b < 0$$

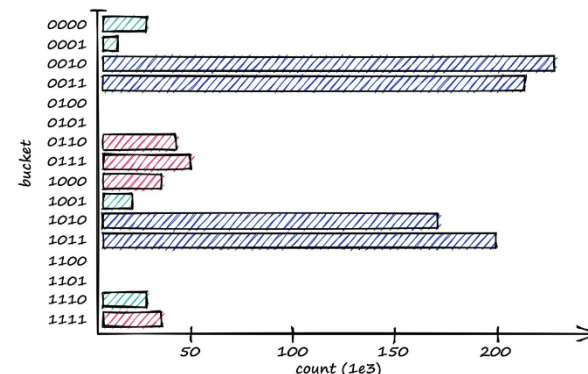
$$n \cdot c < 0$$



Representing the plane (that passes through origin) with a normal vector, the side is calculated by thresholding the dot product with datapoint



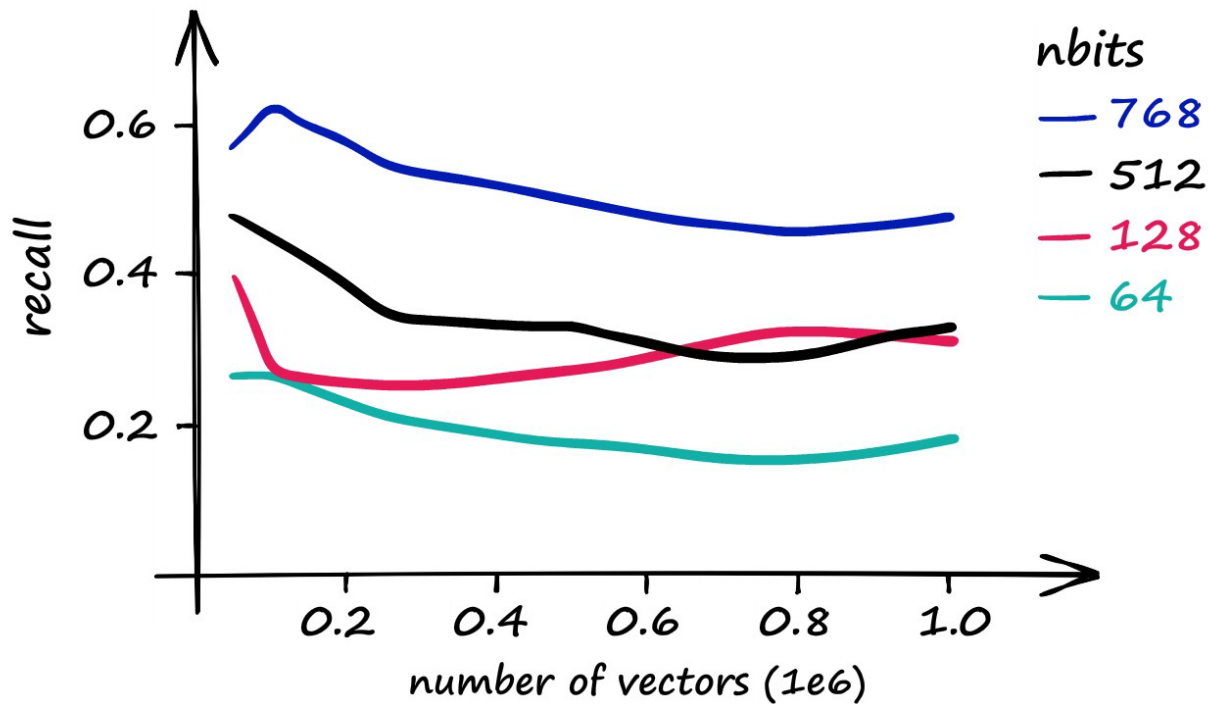
With multiple (randomly generated) planes, each datapoint is represented as a bit vector



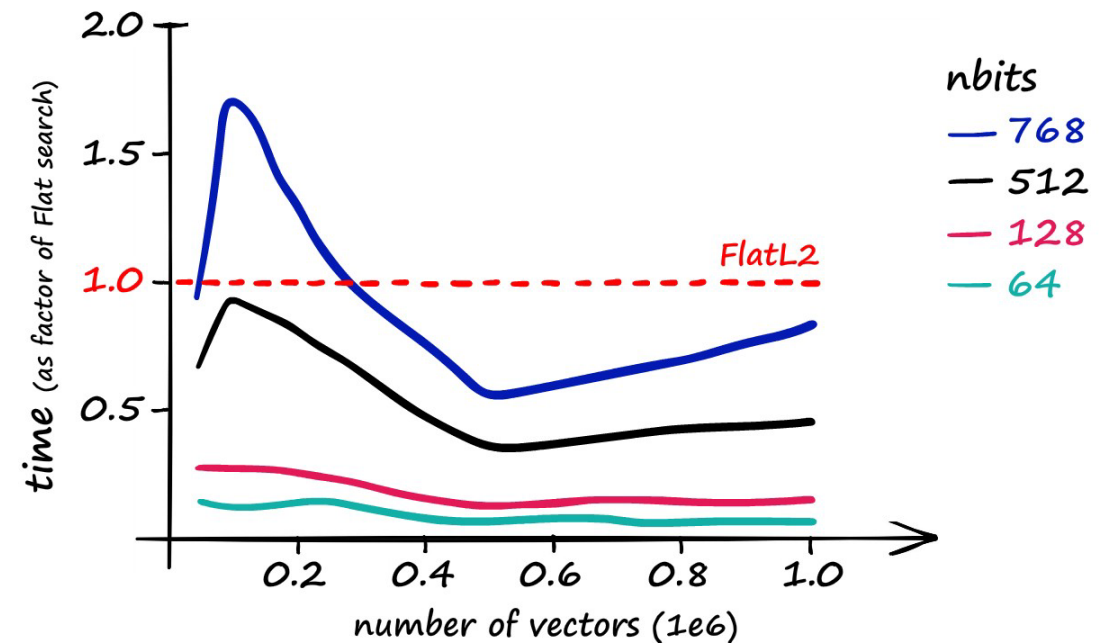
Can retrieve points in the same "bucket" and either randomly pick one or search within it

Key parameter: nbits

- Rule of thumb: nbits = dim is a decent choice (1 bit per feature dimension)



Recall vs. exact nearest neighbor



Time compared to brute force search

Nice video about LSH in faiss:

https://youtu.be/ZLfdQg_u7Eo

which is part of this very detailed and helpful post:

<https://www.pinecone.io/learn/locality-sensitive-hashing-random-projection/>

Comparison of search algorithms

Index	Idea	Time to Build	Memory (MB)	Query Time (ms)	Recall	Notes
Flat	Direct: Compare query to every data point	Fastest	~500	~18	1.0	Good for small datasets or where query time is irrelevant
LSH	Hash-based: Approximate search by approximating floating point vector with bit vector	Fast	20 - 600	1.7 - 30	0.4 - 0.85	Best for low dimensional data, or small datasets
HNSW	Graph-based: Search neighbors in a well-constructed graph	Slow	600 - 1600	0.6 - 2.1	0.5 - 0.95	Very good for quality, high speed, but large memory usage
IVF	Cluster-based: Assign query to cluster and also search nearby clusters	Medium	~520	1 - 9	0.7 - 0.95	Good scalable option. High-quality, at reasonable speed and memory usage

Time/memory/recall from [Nearest Neighbor Indexes for Similarity Search | Pinecone](#): The Sift1M dataset consists of a dimensionality of **128**, dataset size of **1M**. The above values result from querying for a single vector – returning 10 of the nearest matches.

Q1

<https://tinyurl.com/AML441-L4>

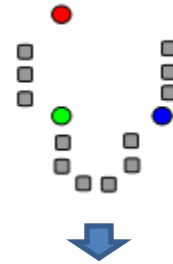


Clustering

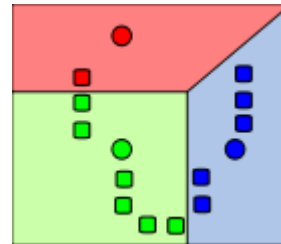
- Assign a label to each data point based on the similarities between points
- Why cluster
 - Represent data point with a single integer instead of a floating point vector
 - Saves space
 - Simple to count and estimate probability
 - Discover trends in the data
 - Make predictions based on groupings

K-means algorithm

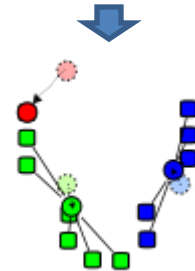
1. Randomly select K centers



2. Assign each point to nearest center (L2 dist)

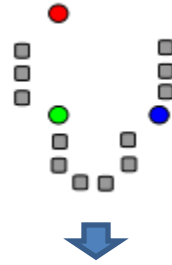


3. Compute new center (mean) for each cluster

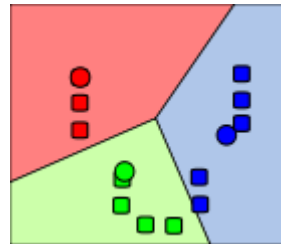


K-means algorithm

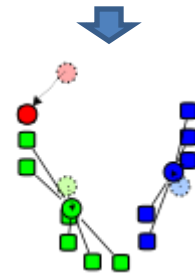
1. Randomly select K centers



2. Assign each point to nearest center (L2 dist)



3. Compute new center (mean) for each cluster



Back to 2

Pseudo-code

kmeans(X, K, maxiter)

```
# Create cluster centers  
center = X[:K].copy()
```

Until maxiter iterations or convergence:

```
    For each nth sample in X:  
        # get index of nearest center  
        idx[n] = get_nearest(X[n], centers)
```

```
    For each kth center:  
        # get mean of data points assigned to cluster k  
        center[k] = X[idx==k].mean(axis=0)
```

Convergence is if no idx changed in this iteration

```
return center, idx
```

What is the cost minimized by K means?

$$id^*, centers^* = \operatorname{argmin}_{id, centers} \sum_n ||centers_{id_n} - X_n||^2$$

1. Choose ids that minimizes square cost given centers
2. Choose centers that minimize square cost given ids

K-means Demo

<https://www.naftaliharris.com/blog/visualizing-k-means-clustering/>

What are some disadvantages of K-means in terms of clustering quality?

- All feature dimensions have equal weight
- Tends to break data into clusters of similar numbers of points (can be good or bad)
- Does not take into account any local structure
- Typically, no easy way to choose K
- Can be slow if the number of data points and clusters is large

Implementation issues

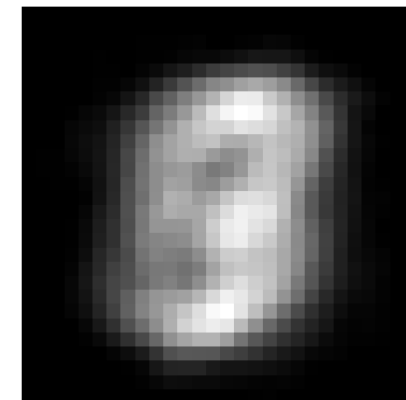
- How to choose K ?
 - Typically chosen by hand
 - Often based on the number of clusters you want considering time/space requirements
- How do you initialize
 - Randomly choose points
 - Iterative furthest point

Evaluating clustering with RMSE to cluster centers

- RMSE = root mean squared error
- Measures a kind of compression loss, i.e. how well is each data point represented by the closest cluster center

$$RMSE(X, C) = \sqrt{\frac{1}{N} \sum_{n \in \{0..N-1\}} \min_k \|C_k - X_n\|_2^2}$$

Example: K-means on MNIST



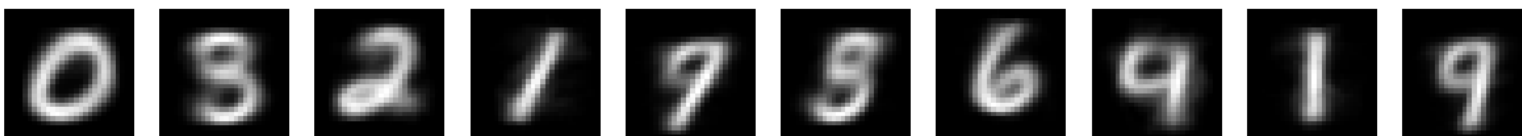
K=1, RMSE = 55.0



K=3, RMSE = 49.3



K=5, RMSE = 45.6



K=10, RMSE = 41.9

K=20, RMSE = 37.82



K=40, RMSE = 34.44



Evaluating clusters with purity

- We often cluster when there is no definitively correct answer, but a *purity* measure can be used to check the consistency with labels

$$purity = \sum_k \left(\max_y \sum_{n: id_n=k} \delta(label_n = y) \right) / N$$

- Purity is the count of data points with the most common label in each cluster, divided by the total number of data points (N)
- E.g., labels = {0, 0, 0, 0, 1, 1, 1, 1}, cluster ids = {0, 0, 0, 0, 0, 1, 1, 1},
purity = ?
purity = 7/8
- Purity can be used to select the number of clusters, or to compare approaches with a given number of clusters
 - A relatively small number of labels can be used to estimate purity, even if there are many data points

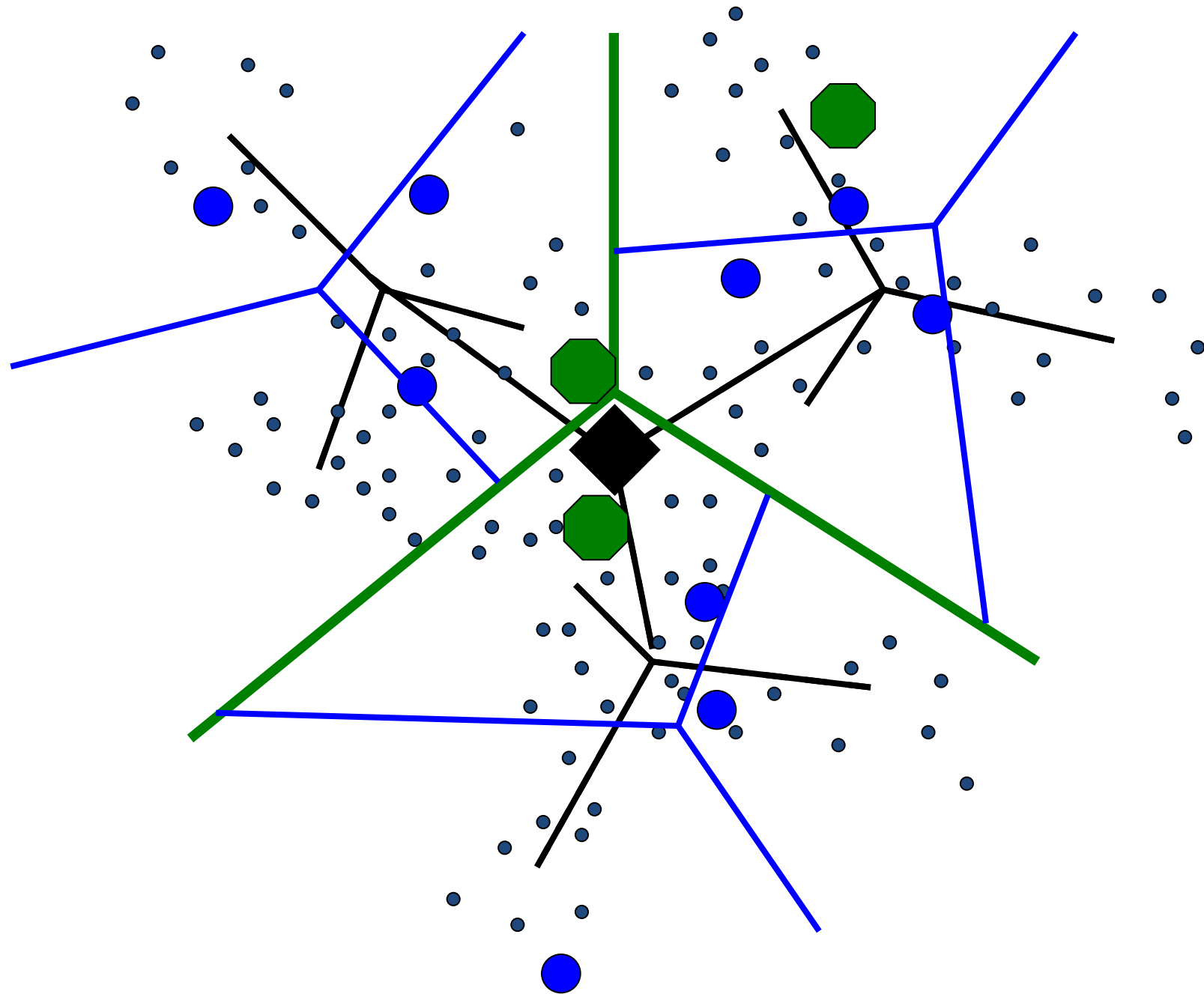
Q2

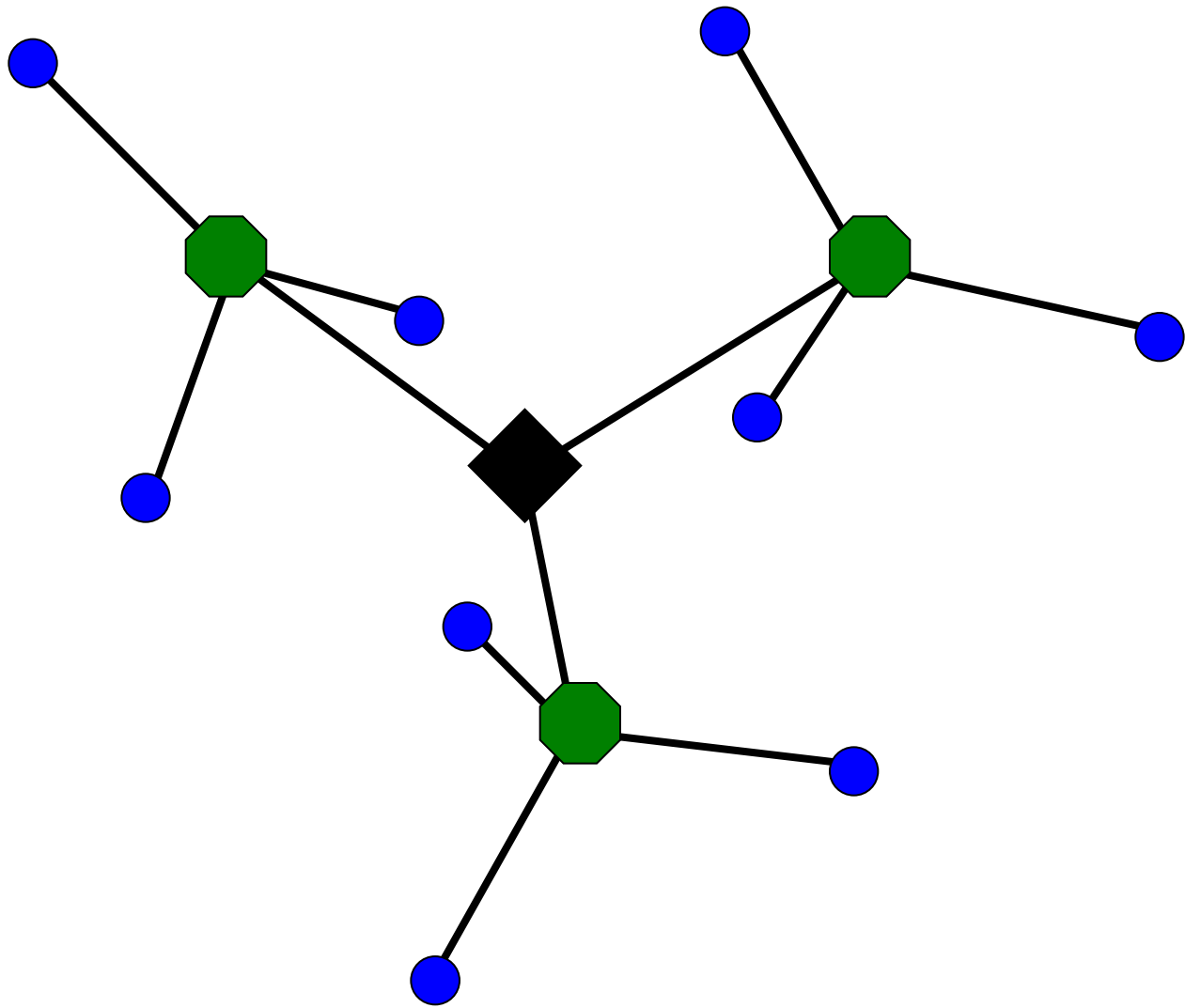
<https://tinyurl.com/AML441-L4>

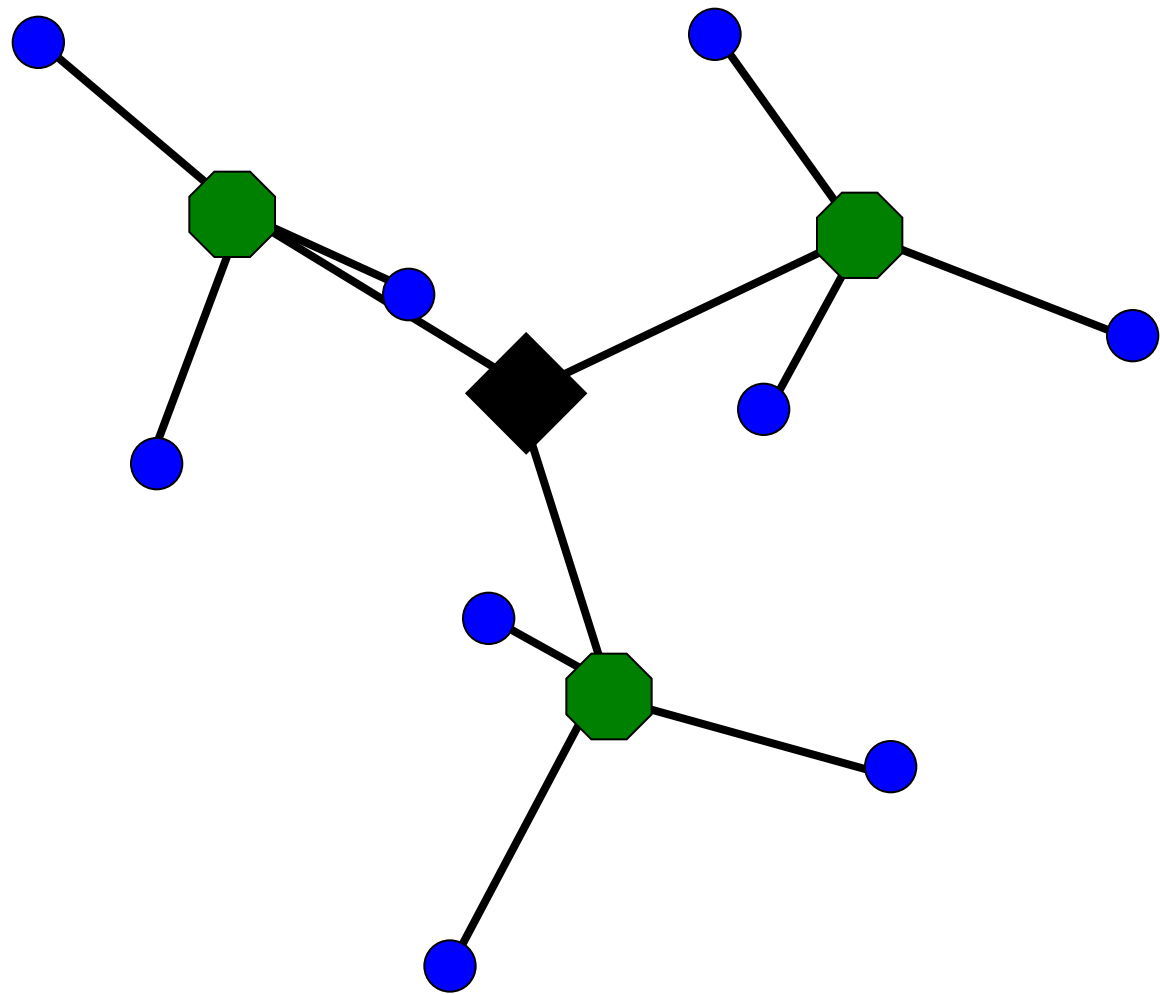


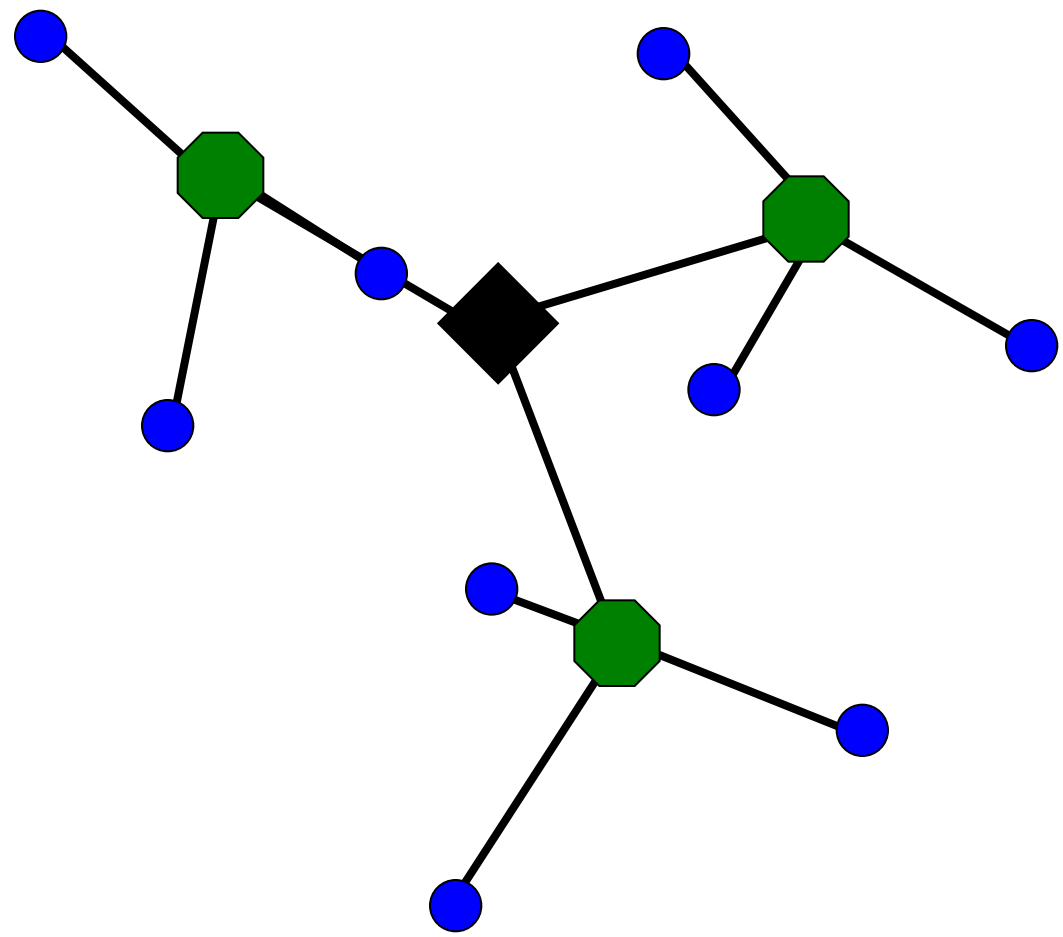
Hierarchical K-means

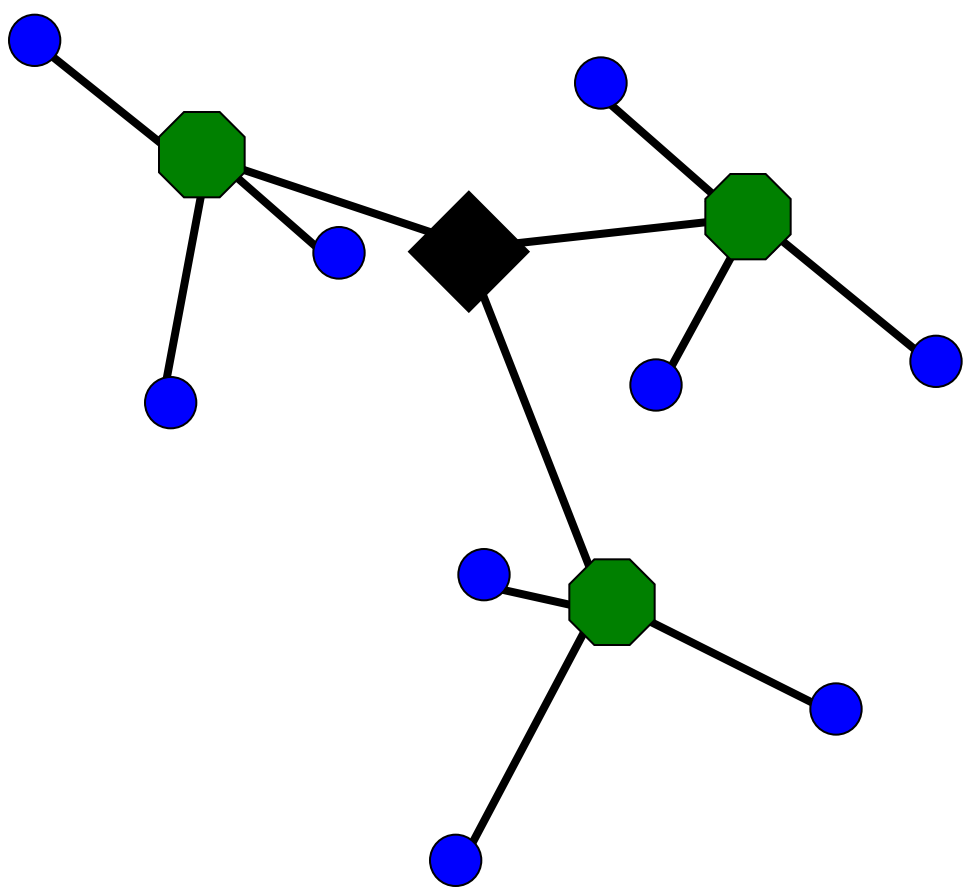
- Iteratively cluster points into K groups, then cluster each group into K groups

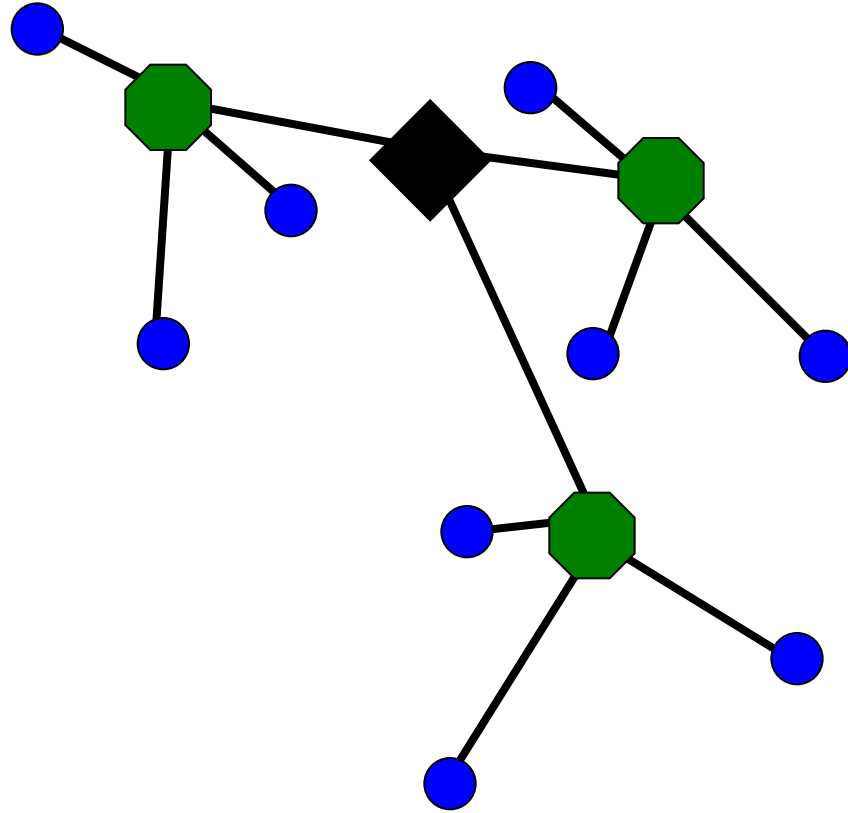


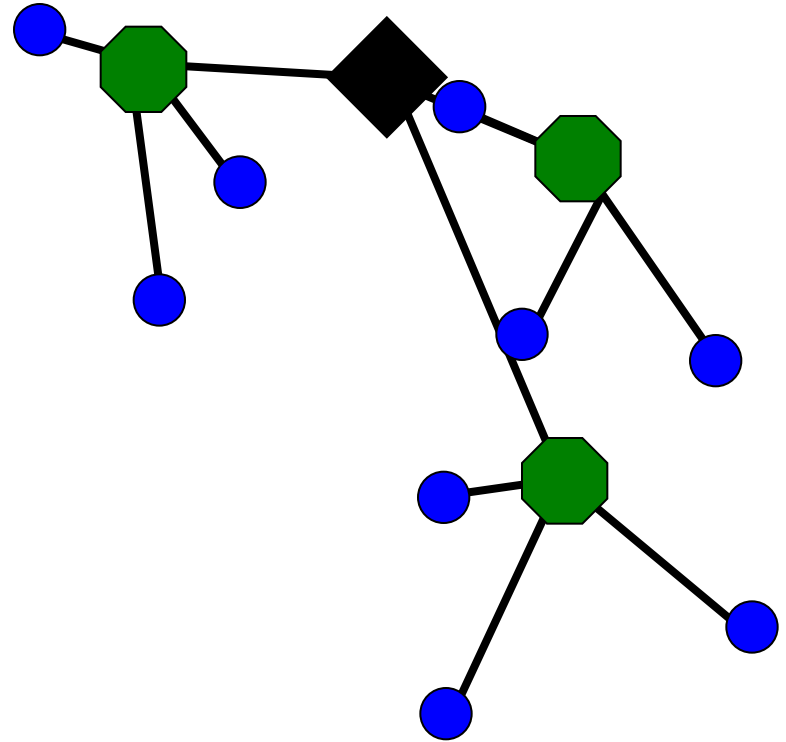


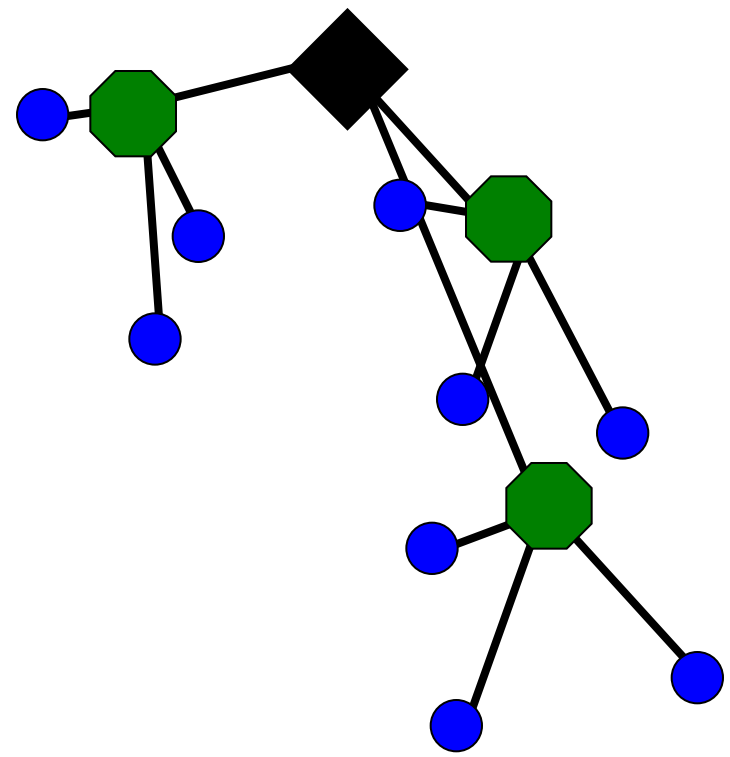


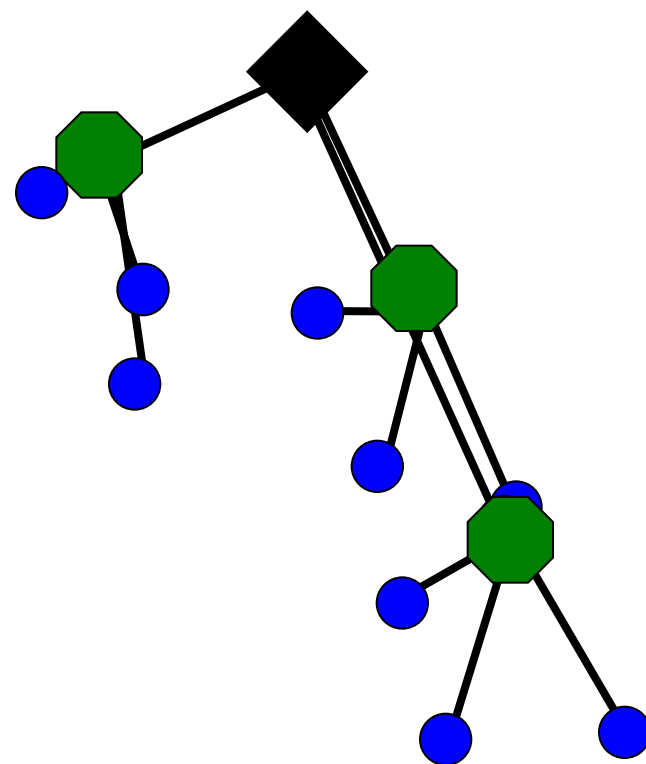


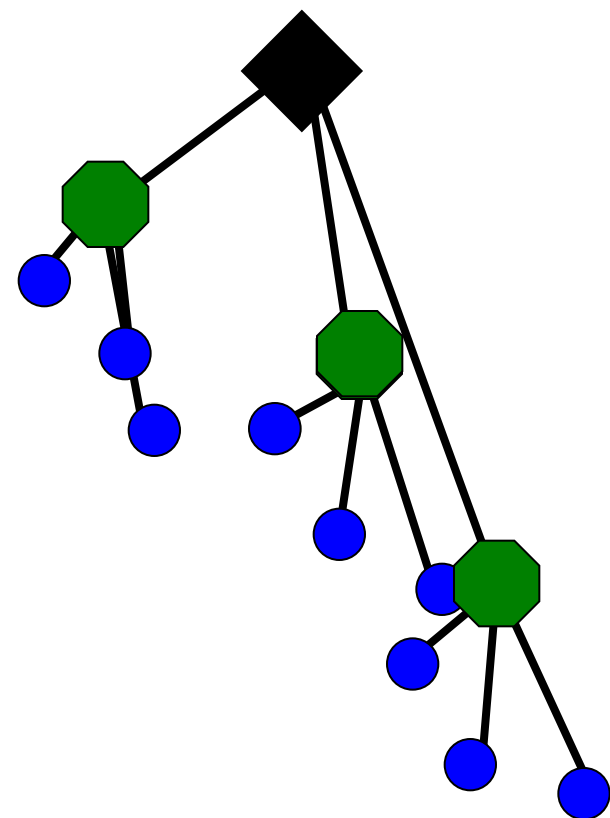


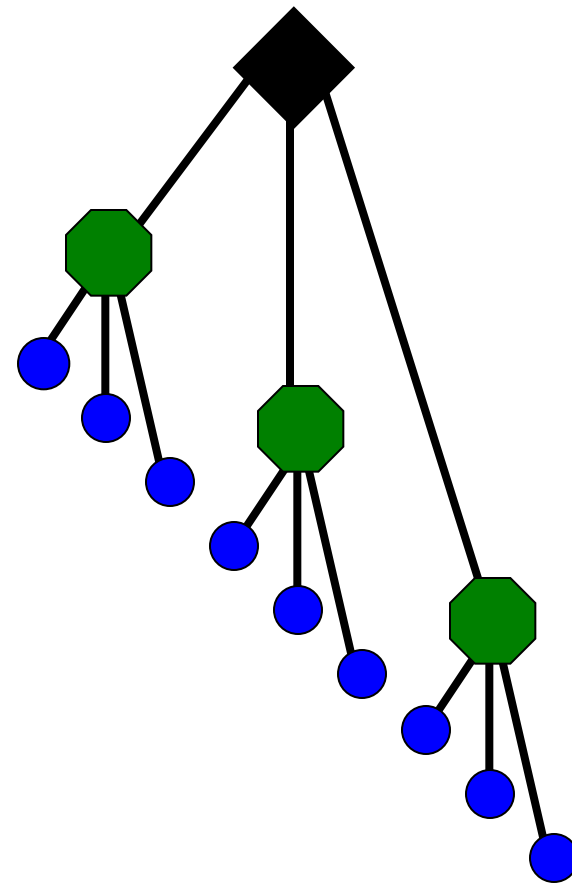












Advantages of Hierarchical K-Means

- Fast cluster training
 - With a branching factor of 10, can cluster into 1M clusters by clustering into 10 clusters $\sim 111,111$ times, each time using e.g. 10K data points
 - Vs. e.g. clustering 1B data points into 1M clusters
 - Kmeans is $O(K*N*D)$ per iteration so this is a 900,000x speedup!
- Fast lookup
 - Find cluster number in $O(\log(K)*D)$ vs. $O(K*D)$
 - 16,667x speedup in the example above



Are there any disadvantages of hierarchical Kmeans?

Yes, the assignment might not be quite as good, but often usually isn't a huge deal since K means is used to approximate data points with centroid anyway

Agglomerative clustering

- Iteratively merge the two most similar points or clusters
 - Can use various distance measures
 - Can use different “linkages”, e.g. distance of nearest points in two clusters or the cluster averages
 - Ideally the minimum distance between clusters should increase after each merge (e.g. if using the distance between cluster centers)
 - Number of clusters can be set based on when the cost to merge increases suddenly

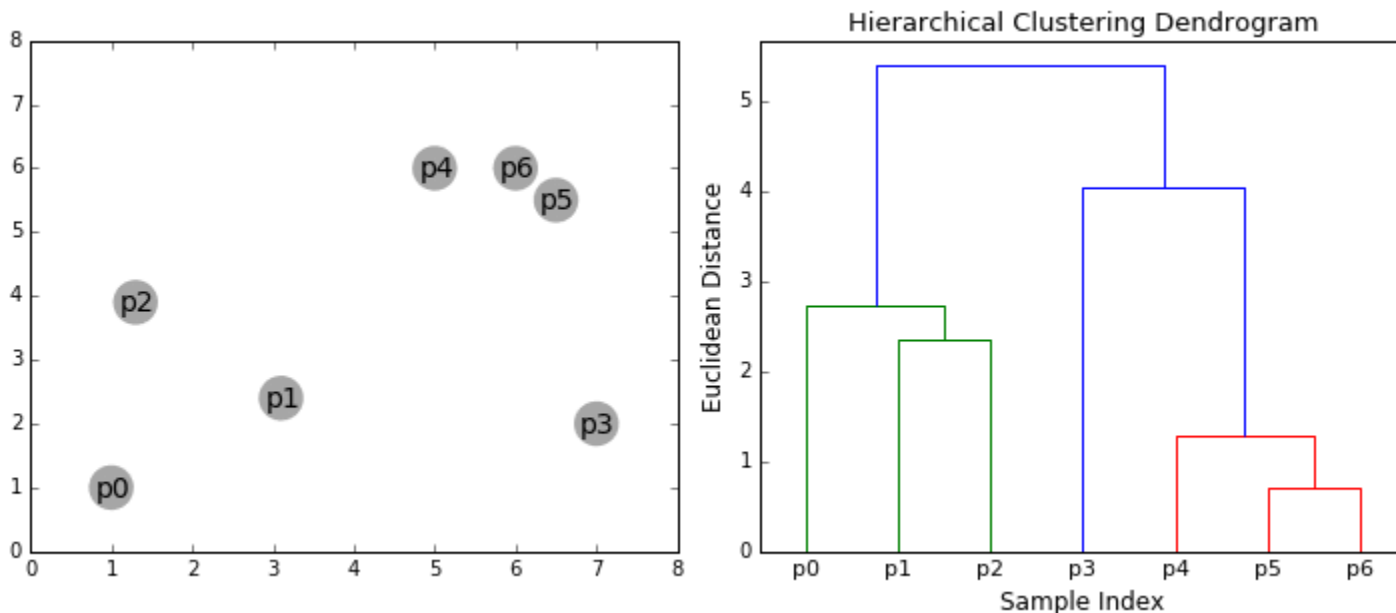
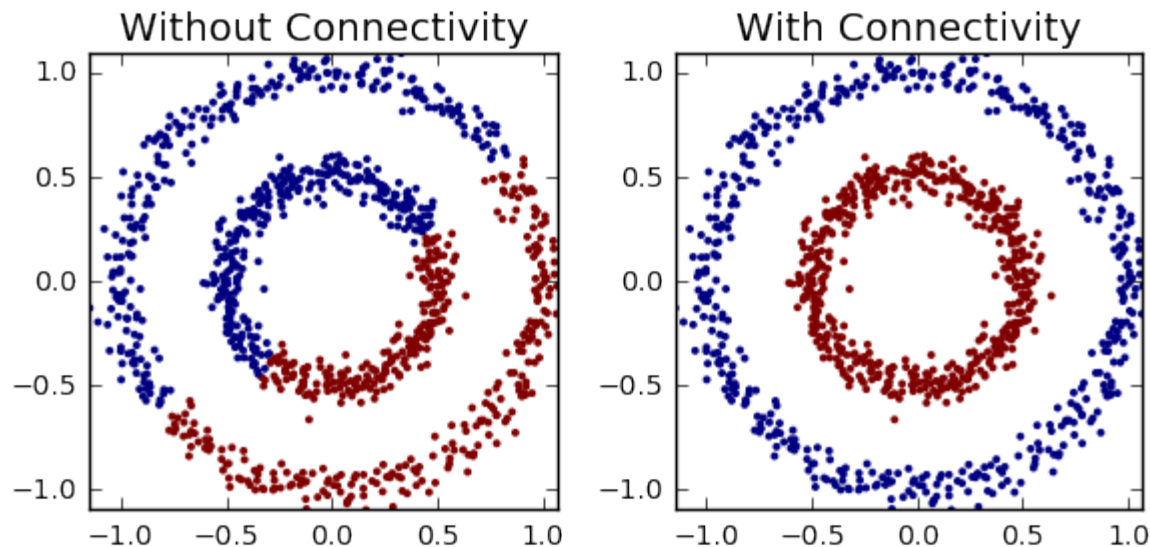


Illustration of Linkage Types

<https://courses.grainger.illinois.edu/CS441/fa2026/tutorials/agglomerative-clustering.html>

Agglomerative clustering

- With good choices of linkage, agglomerative clustering can reflect the data connectivity structure (“manifold”)



Clustering based on distance of 5
nearest neighbors between clusters

Applications of clustering

- K-means
 - Quantization (codebooks for image generation)
 - Search
 - Data visualization (show the average image of clusters of images)
- Hierarchical K-means
 - Fast search (document / image search)
- Agglomerative clustering
 - Finding structures in the data (image segmentation, grouping camera locations together)

Q3-Q4

<https://tinyurl.com/AML441-L4>



HW 1

- You should now be able to complete this HW

Things to remember

- Use highly optimized libraries like FAISS for search/retrieval
- Approximate search methods like LSH can be used to find similar points quickly
- Clustering groups similar data points
- K-means is the must-know method, but there are many others

