



10-601 Introduction to Machine Learning

Machine Learning Department
School of Computer Science
Carnegie Mellon University

Clustering (K-Means)

Clustering Readings:

Murphy 25.5
Bishop 12.1, 12.3
HTF 14.3.0
Mitchell --

Matt Gormley
Lecture 15
March 8, 2017

Reminders

- **Homework 5: Readings / Application of ML**
 - **Release: Wed, Mar. 08**
 - **Due: Wed, Mar. 22 at 11:59pm**

Outline

- **Clustering: Motivation / Applications**
- **Optimization Background**
 - Coordinate Descent
 - Block Coordinate Descent
- **Clustering**
 - Inputs and Outputs
 - Objective-based Clustering
- **K-Means**
 - K-Means Objective
 - Computational Complexity
 - K-Means Algorithm / Lloyd's Method
- **K-Means Initialization**
 - Random
 - Farthest Point
 - K-Means++

Clustering, Informal Goals

Goal: Automatically partition **unlabeled** data into groups of similar datapoints.

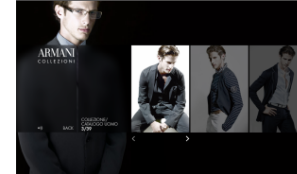
Question: When and why would we want to do this?

Useful for:

- Automatically organizing data.
- Understanding hidden structure in data.
- Preprocessing for further analysis.
 - Representing high-dimensional data in a low-dimensional space (e.g., for visualization purposes).

Applications (Clustering comes up everywhere...)

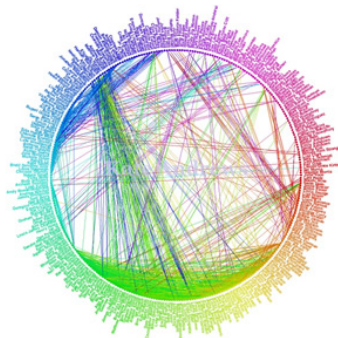
- Cluster news articles or web pages or search results by topic.



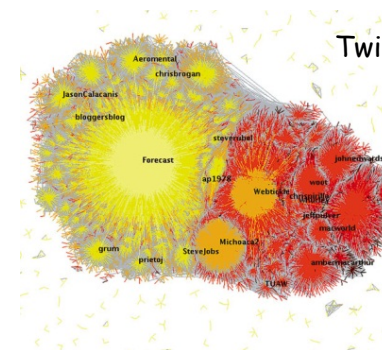
- Cluster protein sequences by function or genes according to expression profile.



- Cluster users of social networks by interest (community detection).



Facebook network



Twitter Network

Applications (Clustering comes up everywhere...)

- Cluster customers according to purchase history.



- Cluster galaxies or nearby stars (e.g. Sloan Digital Sky Survey)



- And many many more applications....

Optimization Background

Whiteboard:

- Coordinate Descent
- Block Coordinate Descent

Clustering

Whiteboard:

- Inputs and Outputs
- Objective-based Clustering

K-Means

Whiteboard:

- K-Means Objective
- Computational Complexity
- K-Means Algorithm / Lloyd's Method

K-Means Initialization

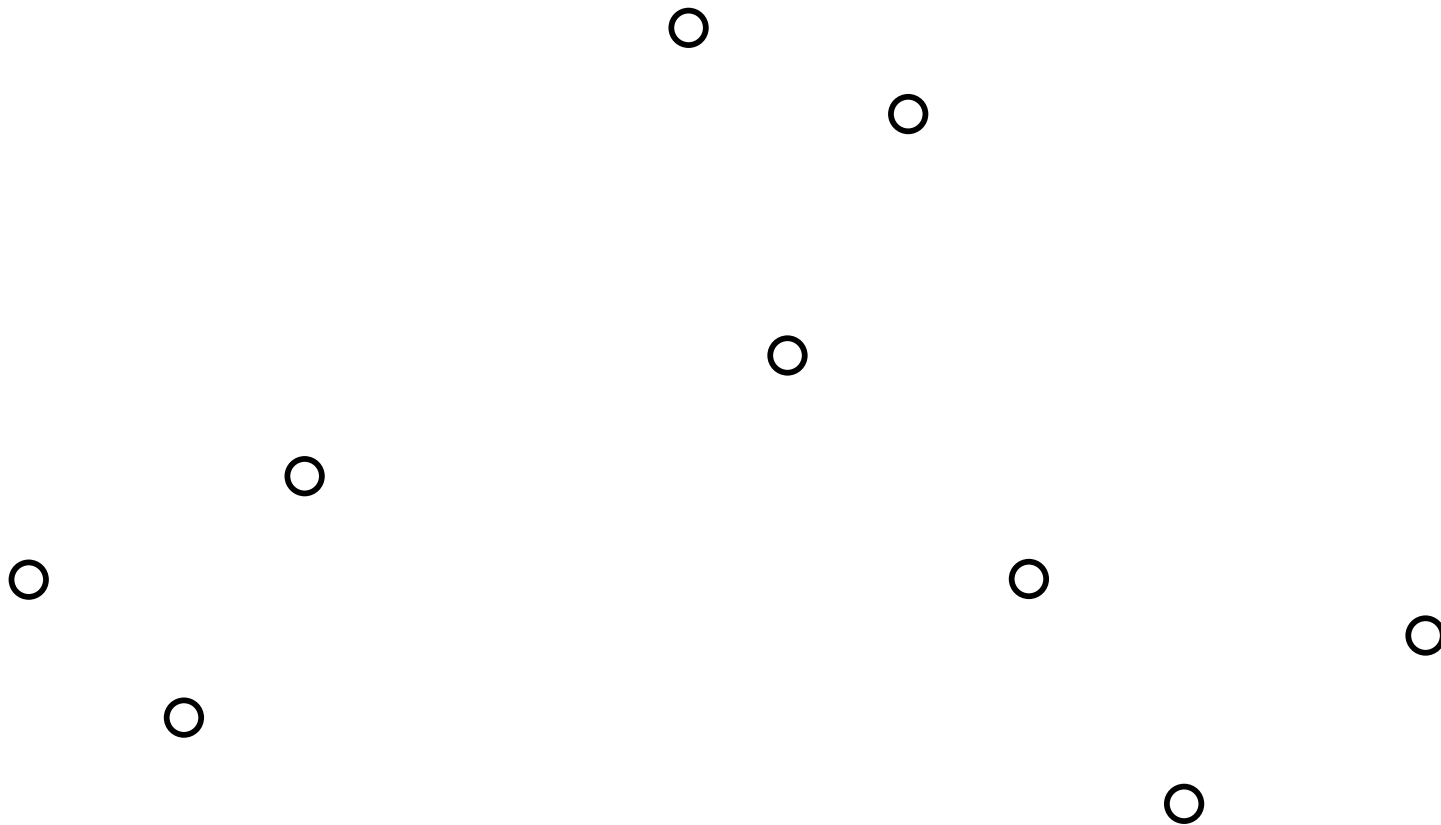
Whiteboard:

- Random
- Furthest Traversal
- K-Means++

Lloyd's method: Random Initialization

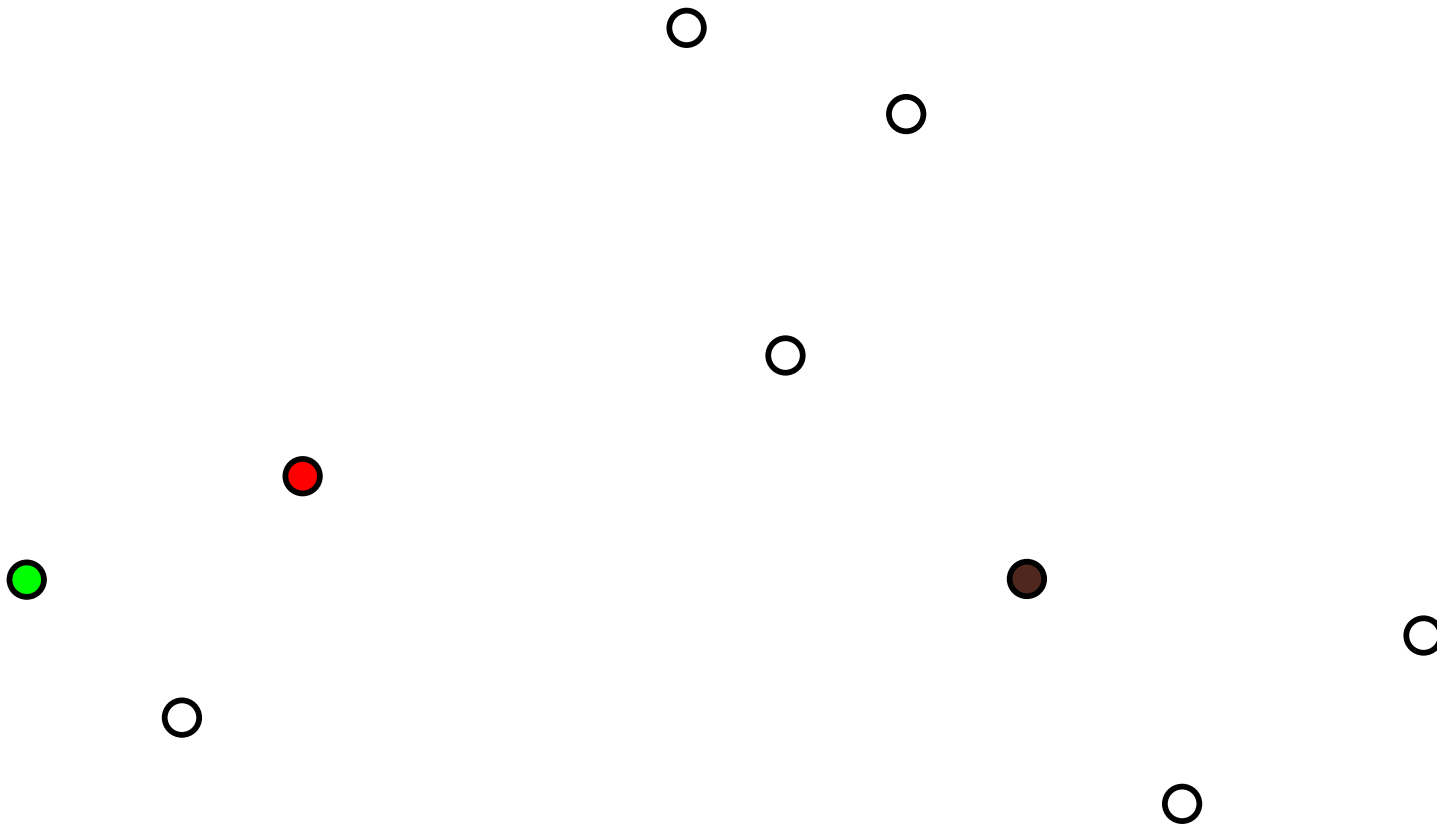
Lloyd's method: Random Initialization

Example: Given a set of datapoints



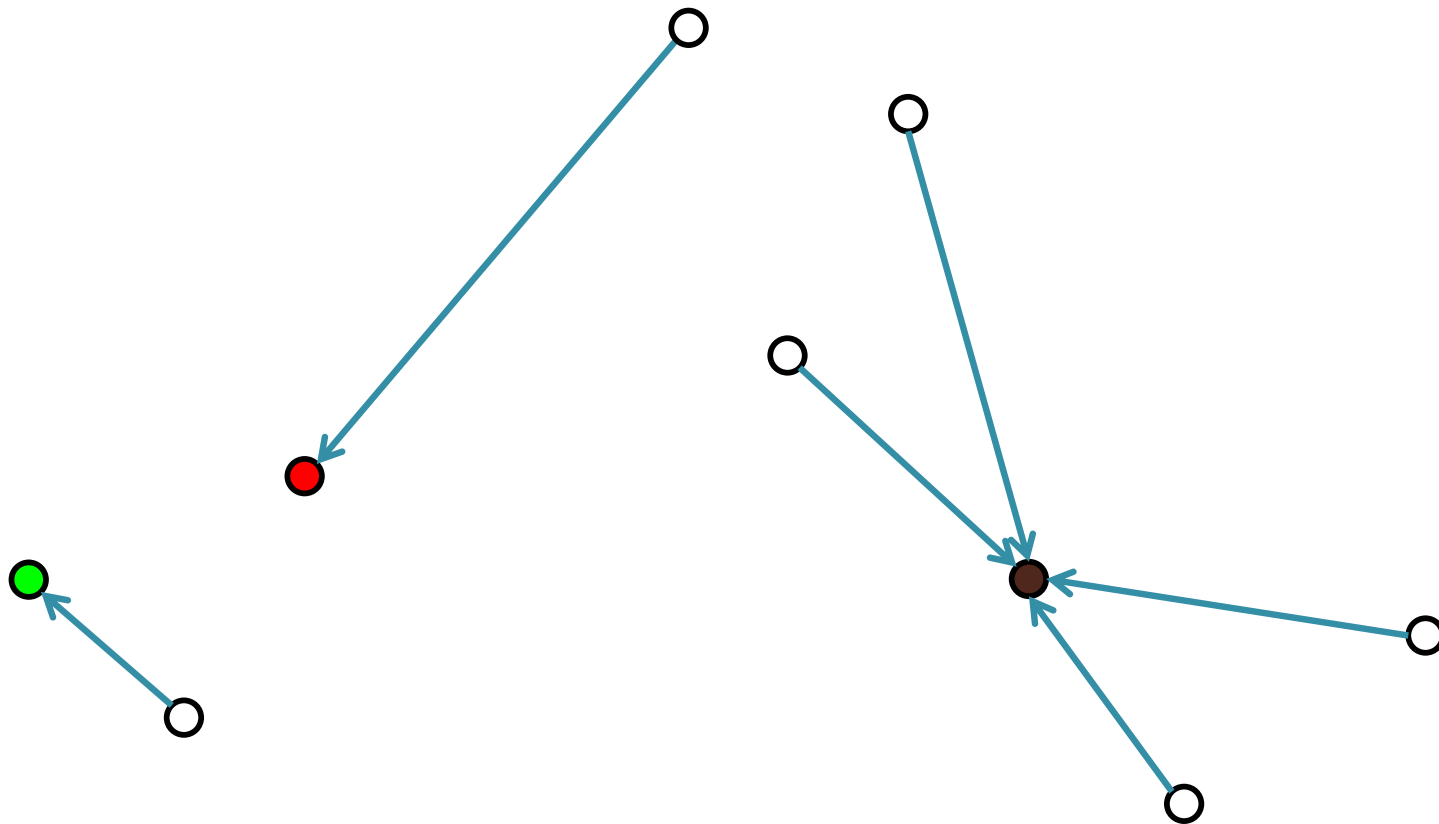
Lloyd's method: Random Initialization

Select initial centers at random



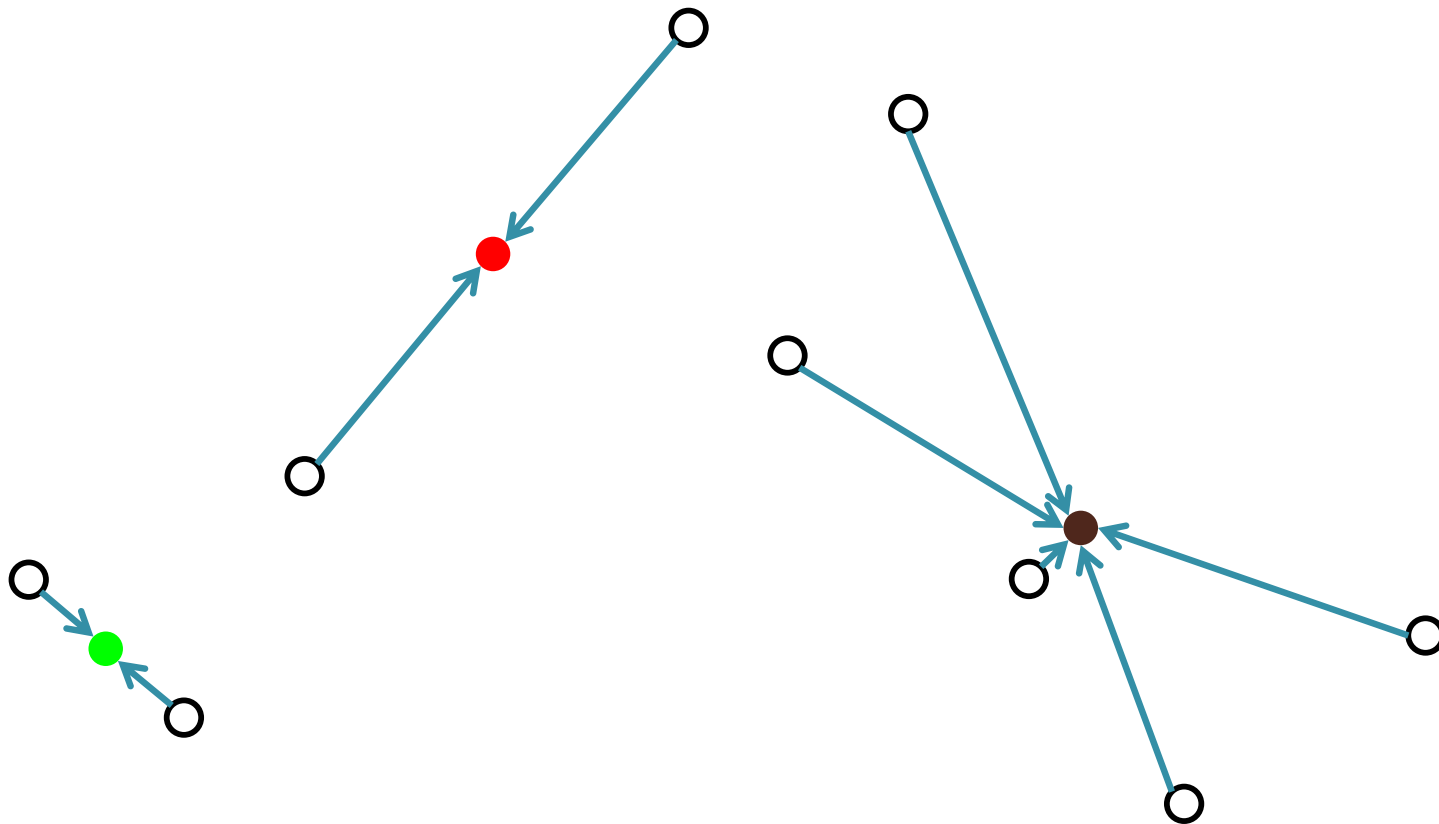
Lloyd's method: Random Initialization

Assign each point to its nearest center



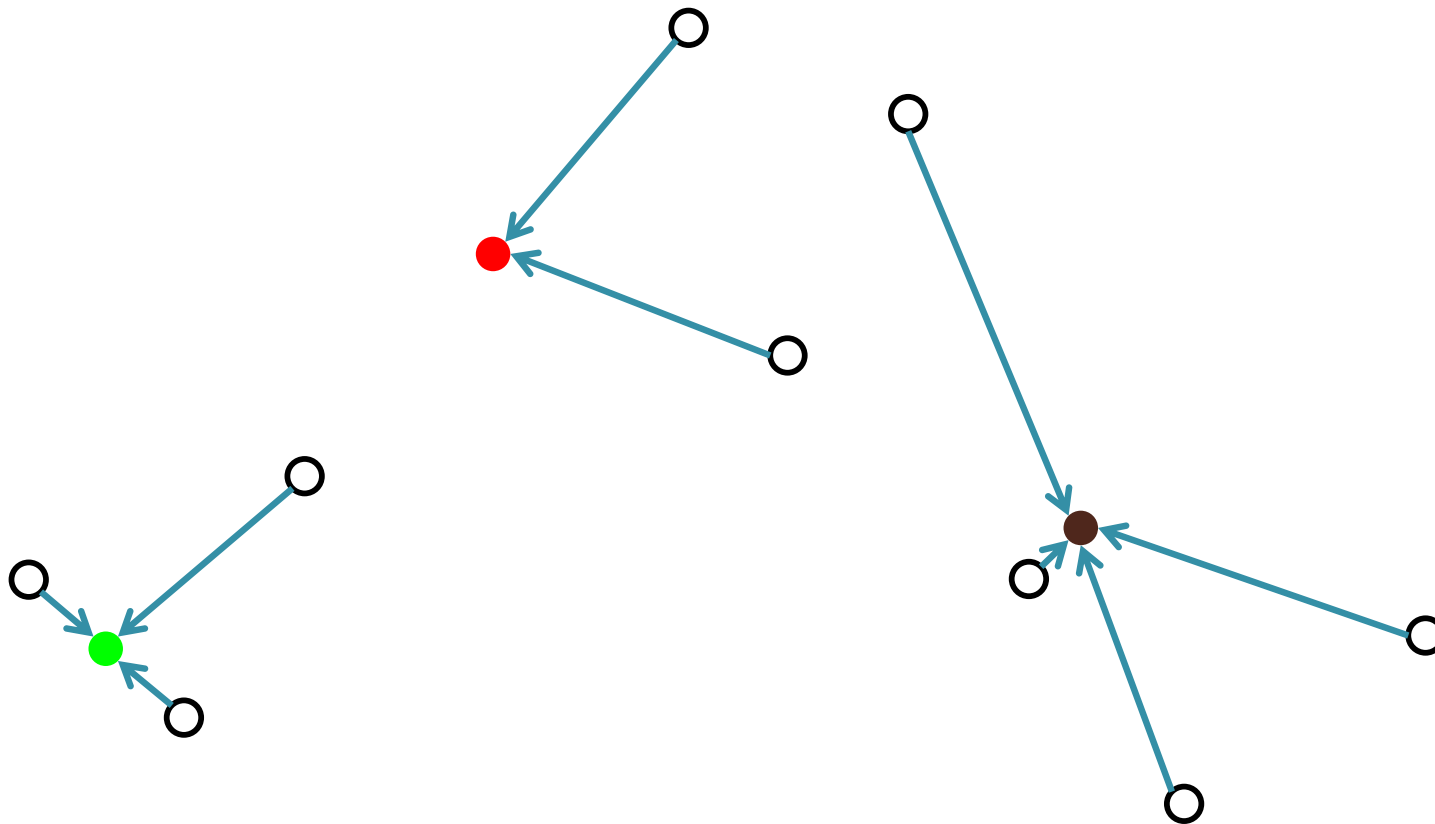
Lloyd's method: Random Initialization

Recompute optimal centers given a fixed clustering



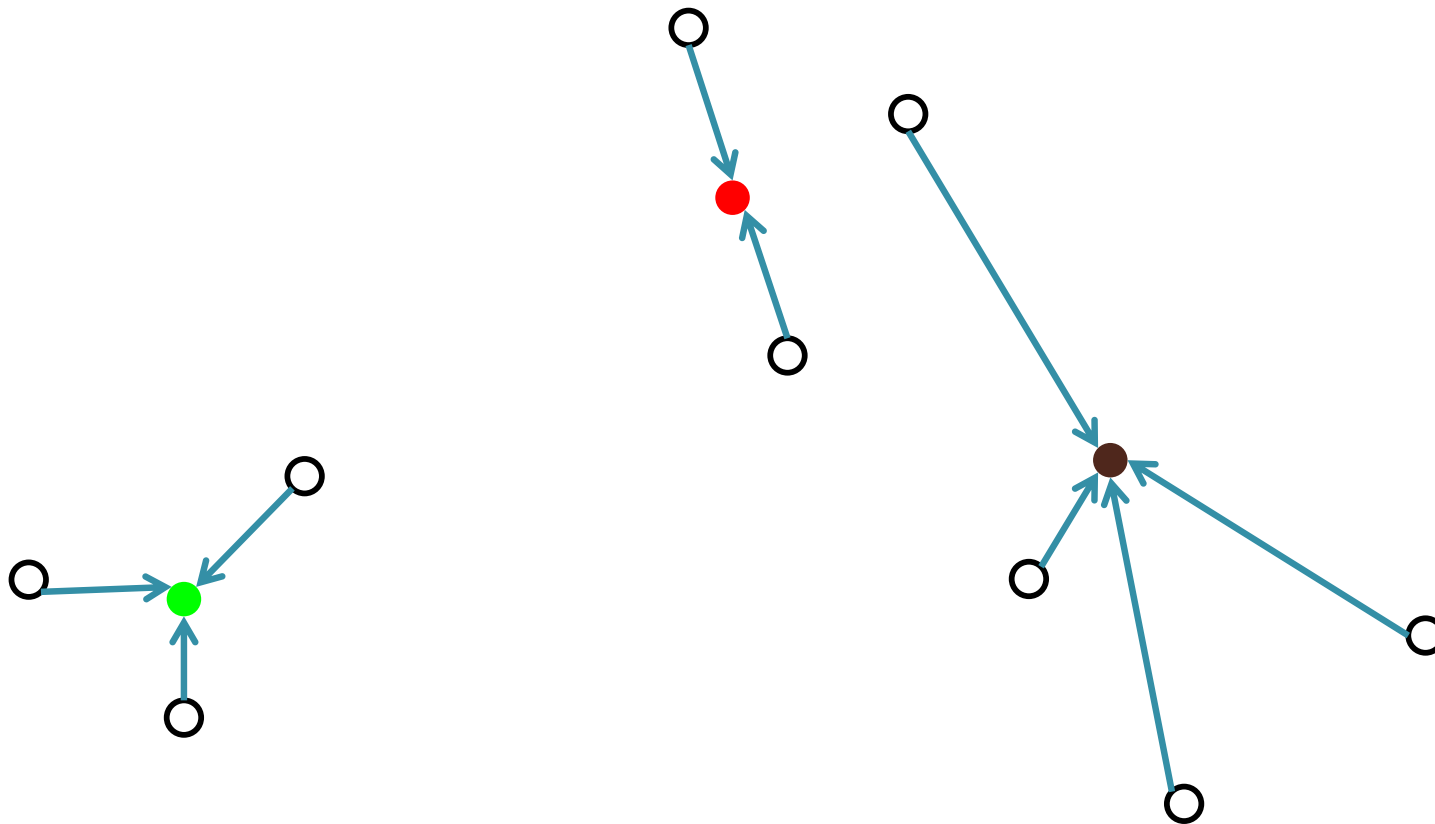
Lloyd's method: Random Initialization

Assign each point to its nearest center



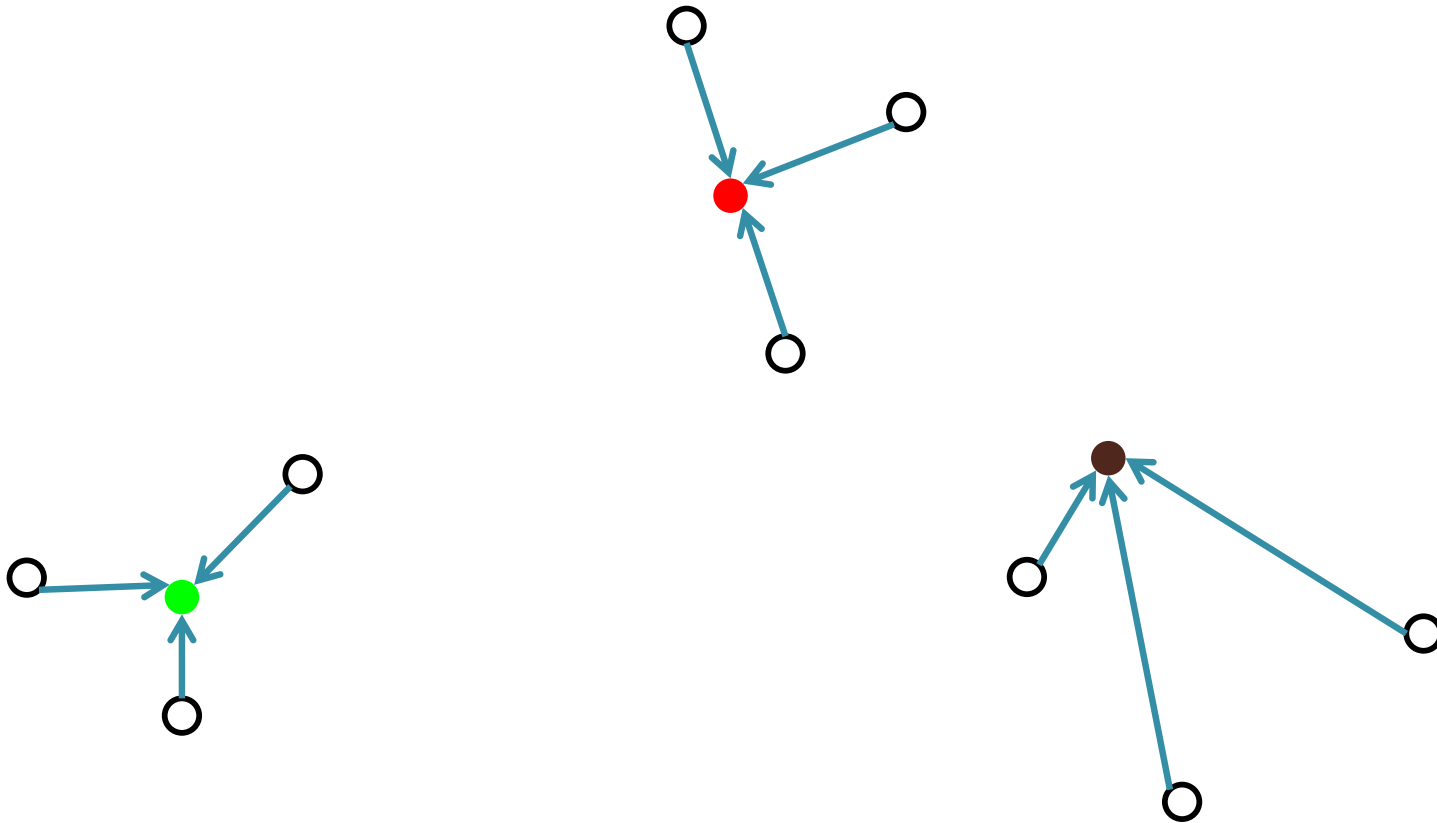
Lloyd's method: Random Initialization

Recompute optimal centers given a fixed clustering



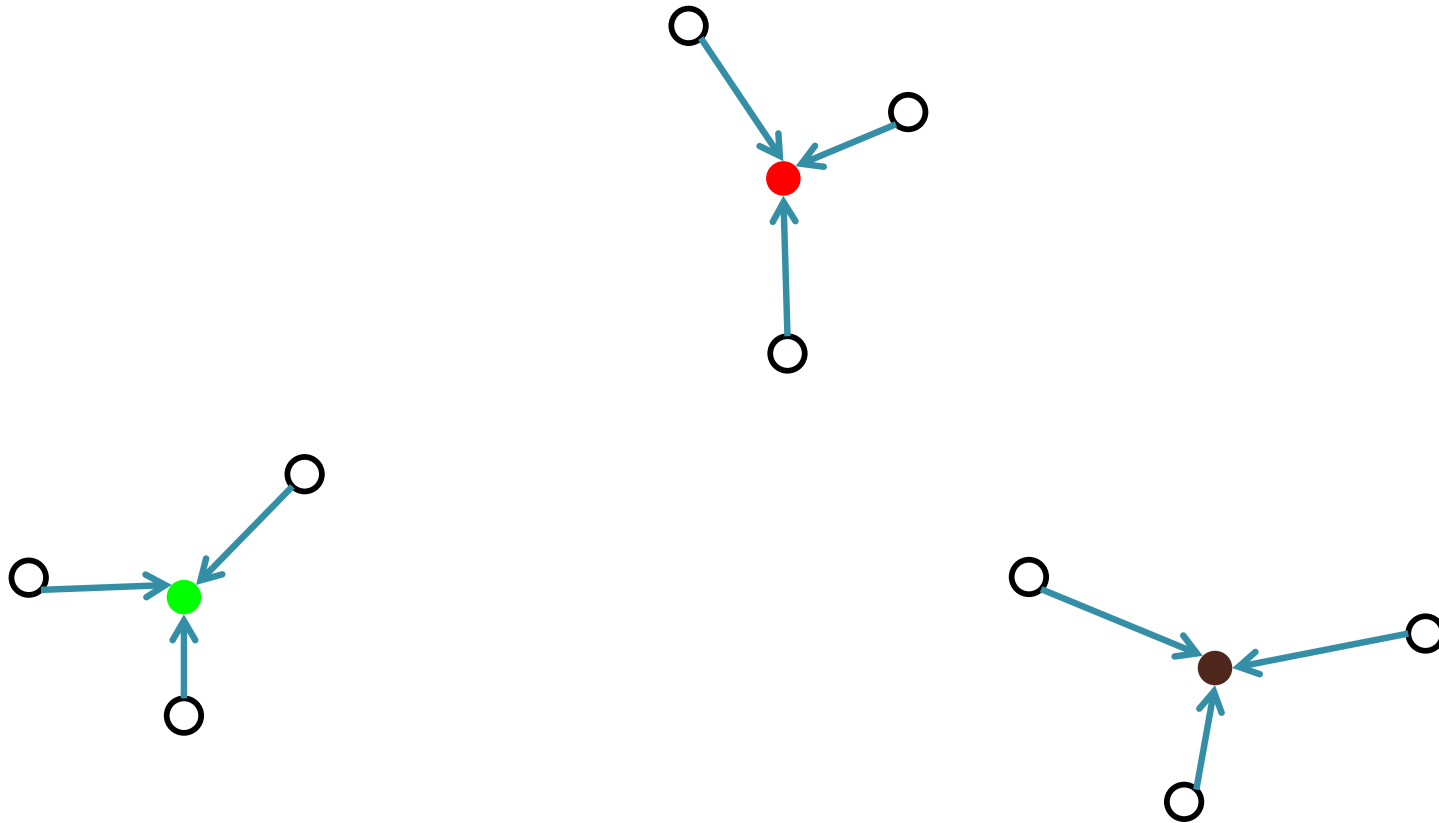
Lloyd's method: Random Initialization

Assign each point to its nearest center



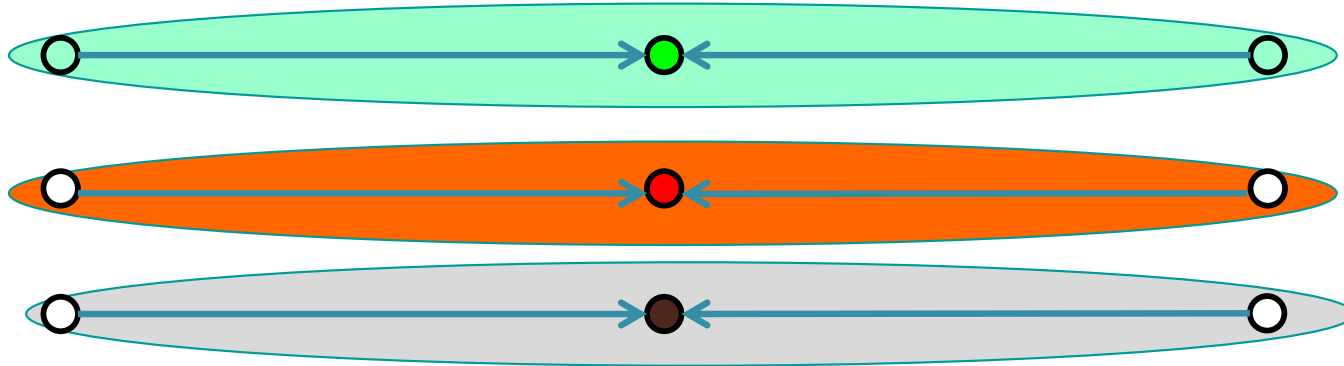
Lloyd's method: Random Initialization

Recompute optimal centers given a fixed clustering



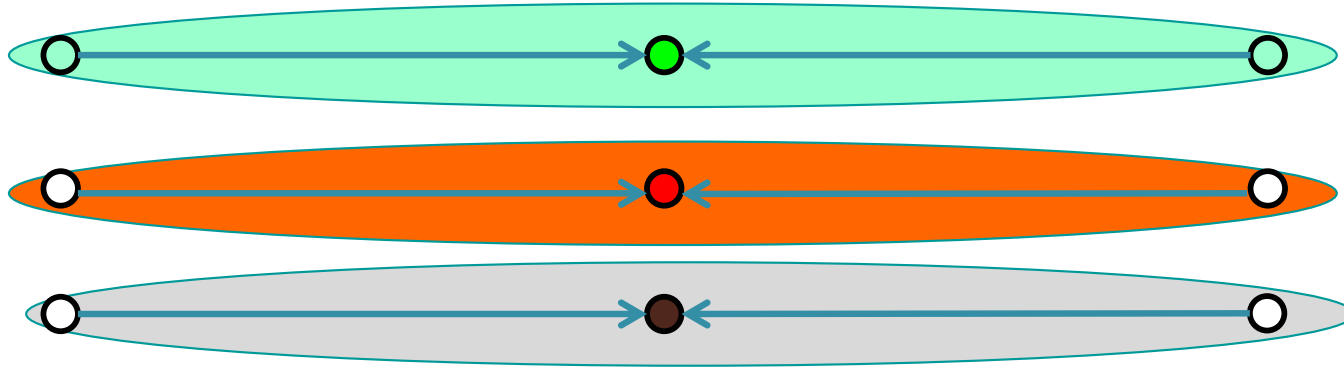
Get a good quality solution in this example.

Lloyd's method: Performance



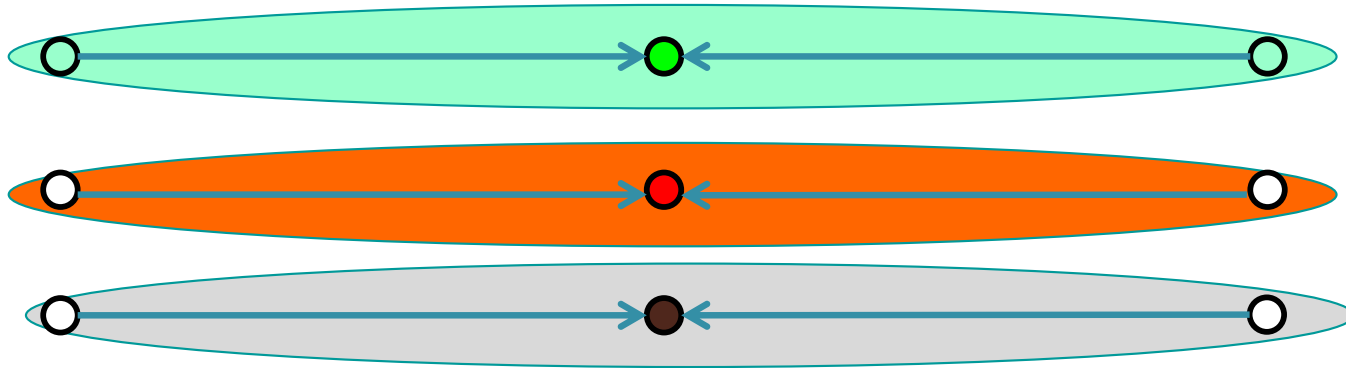
It always converges, but it may converge at a local optimum that is different from the global optimum, and in fact could be arbitrarily worse in terms of its score.

Lloyd's method: Performance

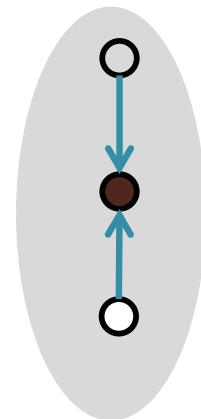
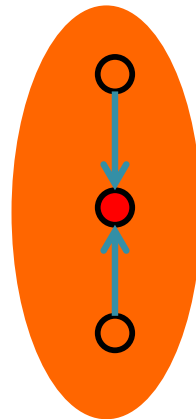
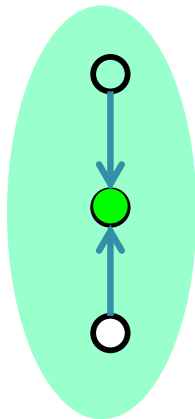


Local optimum: every point is assigned to its nearest center and every center is the mean value of its points.

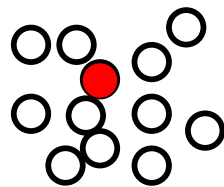
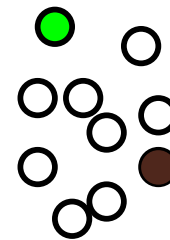
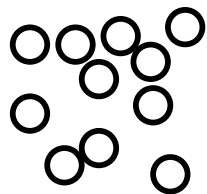
Lloyd's method: Performance



.It is arbitrarily worse than optimum solution....

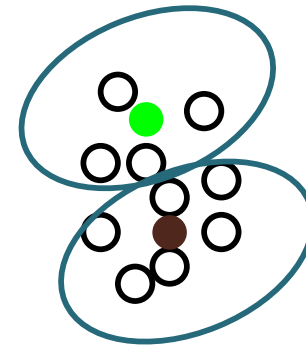
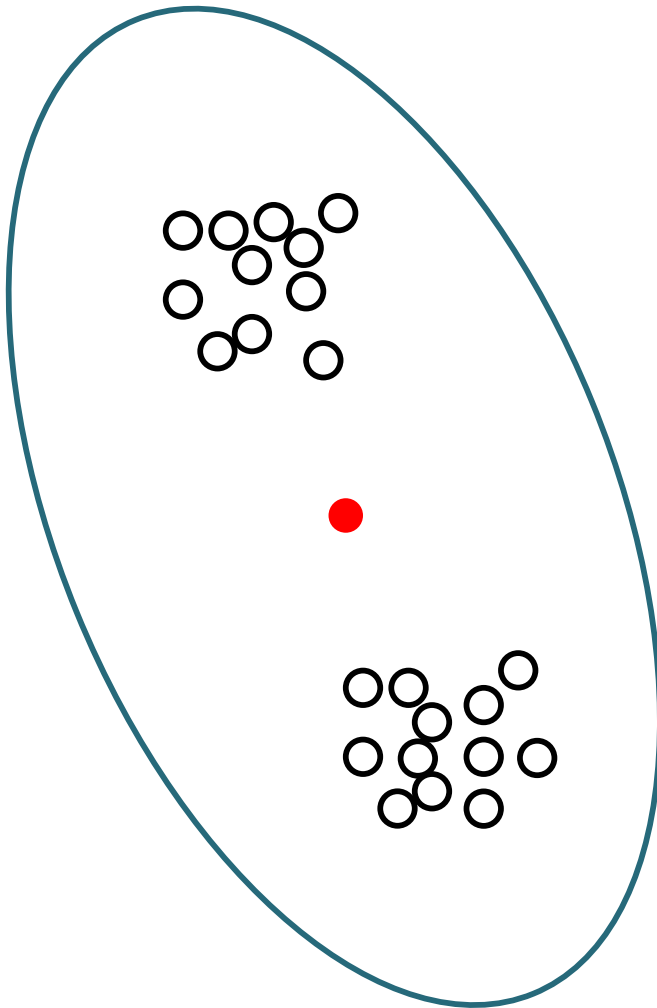


Lloyd's method: Performance



This bad performance, can happen even with well separated Gaussian clusters.

Lloyd's method: Performance



This bad performance, can happen even with well separated Gaussian clusters.

Some Gaussian are combined.....



Lloyd's method: Performance

- If we do random initialization, as k increases, it becomes more likely we won't have perfectly picked one center per Gaussian in our initialization (so Lloyd's method will output a bad solution).
 - For k equal-sized Gaussians, $\Pr[\text{each initial center is in a different Gaussian}] \approx \frac{k!}{k^k} \approx \frac{1}{e^k}$
 - Becomes unlikely as k gets large.

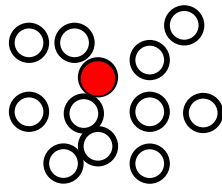
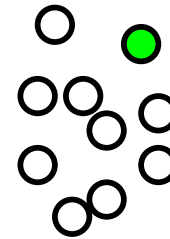
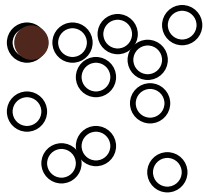
Another Initialization Idea: Furthest Point Heuristic

Choose $\mathbf{c}_1$ arbitrarily (or at random).

- For $j = 2, \dots, k$
 - Pick $\mathbf{c}_j$ among datapoints $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^n$ that is farthest from previously chosen $\mathbf{c}_1, \mathbf{c}_2, \dots, \mathbf{c}_{j-1}$

Fixes the Gaussian problem. But it can be thrown off by outliers....

Furthest point heuristic does well on previous example



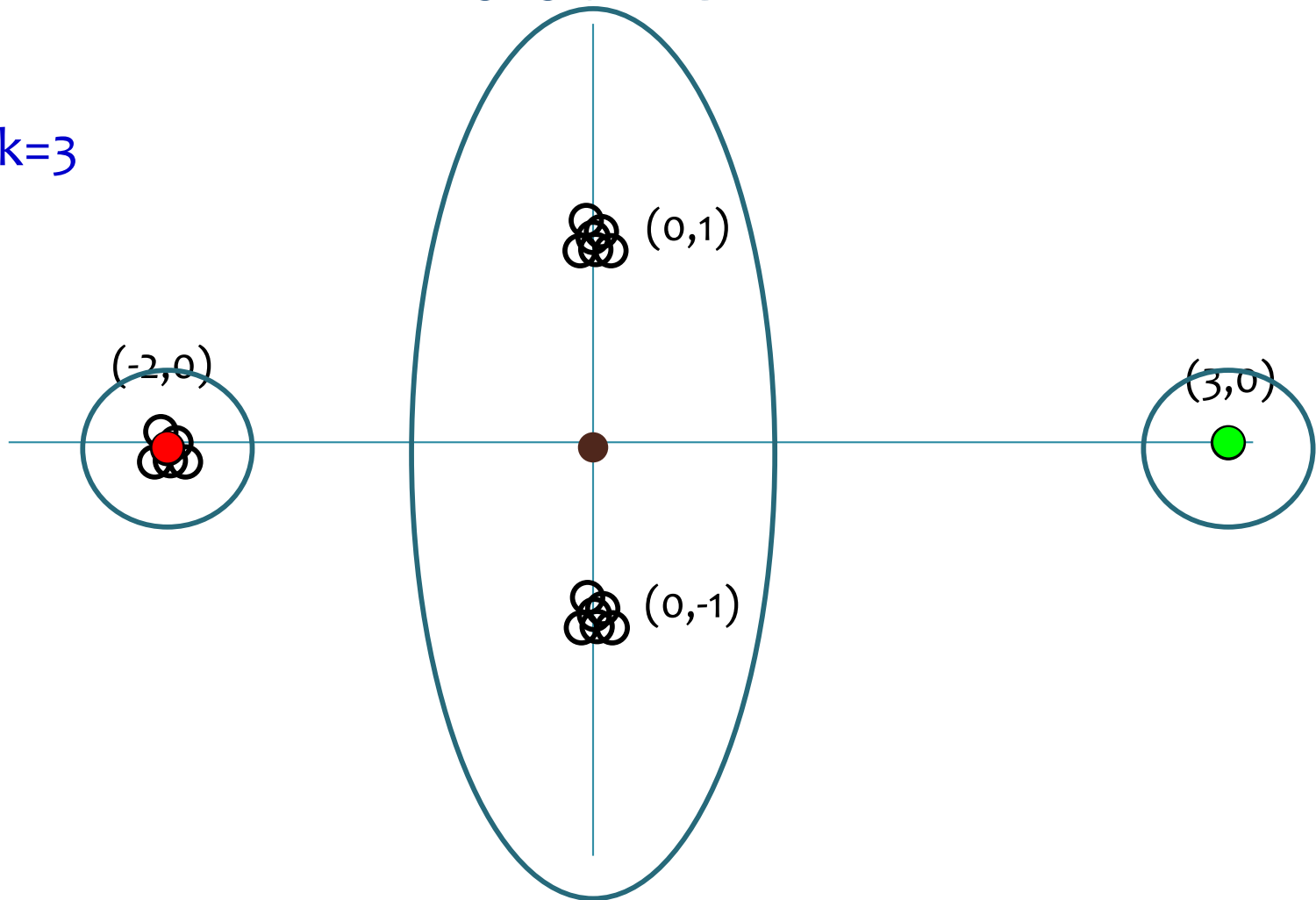
Furthest point initialization heuristic sensitive to outliers

Assume $k=3$



Furthest point initialization heuristic sensitive to outliers

Assume $k=3$



K-means++ Initialization: D^2 sampling [AV07]

- Interpolate between random and furthest point initialization
- Let $D(\mathbf{x})$ be the distance between a point x and its nearest center. Chose the next center proportional to $D^2(\mathbf{x})$.

- Choose $\mathbf{c}_1$ at random.
- For $j = 2, \dots, k$
 - Pick $\mathbf{c}_j$ among $\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^n$ according to the distribution

$$\Pr(\mathbf{c}_j = \mathbf{x}^i) \propto \min_{j' < j} \|\mathbf{x}^i - \mathbf{c}_{j'}\|^2 D^2(\mathbf{x}^i)$$

Theorem: K-means++ always attains an $O(\log k)$ approximation to optimal k-means solution in expectation.

Running Lloyd's can only further improve the cost.

K-means++ Idea: D^2 sampling

- Interpolate between random and furthest point initialization
- Let $D(\mathbf{x})$ be the distance between a point x and its nearest center. Chose the next center proportional to $D^\alpha(\mathbf{x})$.
 - $\alpha = 0$, random sampling
 - $\alpha = \infty$, furthest point (Side note: it actually works well for k-center)
 - $\alpha = 2$, k-means++

Side note: $\alpha = 1$, works well for k-median

K-means ++ Fix



K-means++/ Lloyd's Running Time

- K-means ++ initialization: $O(nd)$ and one pass over data to select next center. So $O(nkd)$ time in total.
- Lloyd's method

Repeat until there is no change in the cost.

- For each j : $C_j \leftarrow \{x \in S \text{ whose closest center is } c_j\}$
- For each j : $c_j \leftarrow \text{mean of } C_j$

Each round takes time $O(nkd)$.

- Exponential # of rounds in the worst case [AV07].
- Expected polynomial time in the smoothed analysis (non worst-case) model!

K-means++/ Lloyd's Summary

- Running Lloyd's can only further improve the cost.
- Exponential # of rounds in the worst case [AV07].
- Expected polynomial time in the smoothed analysis model!
- Does well in practice.
- K-means++ always attains an $O(\log k)$ approximation to optimal k-means solution in expectation.

What value of k ???

- Heuristic: Find large gap between $k-1$ -means cost and k -means cost.
- Hold-out validation/cross-validation on auxiliary task (e.g., supervised learning task).
- Try hierarchical clustering.