

# Announcements

## Assignments

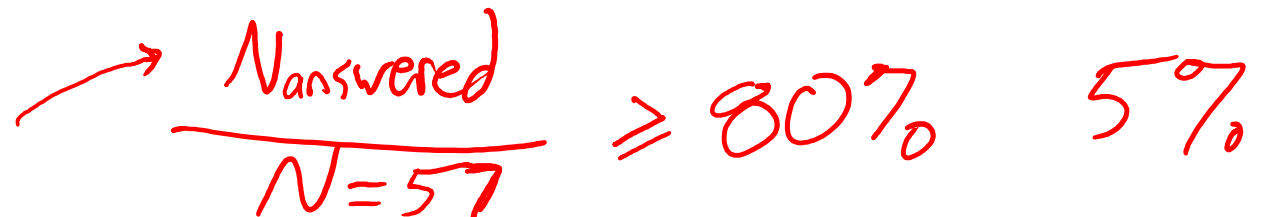
- HW9 (online)
  - Due Thu 4/16, 11:59 pm

## TA Applications:

- Please apply to TA with us! Link for MLD applications will be in piazza.

## Participation points

- Starting now, we're capping the denominator (57 polls) in the participation points calculation



Handwritten red formula:  $\frac{N_{\text{answered}}}{N=57} \geq 80\% \quad 5\%$

# Introduction to Machine Learning

## Clustering GMM and EM

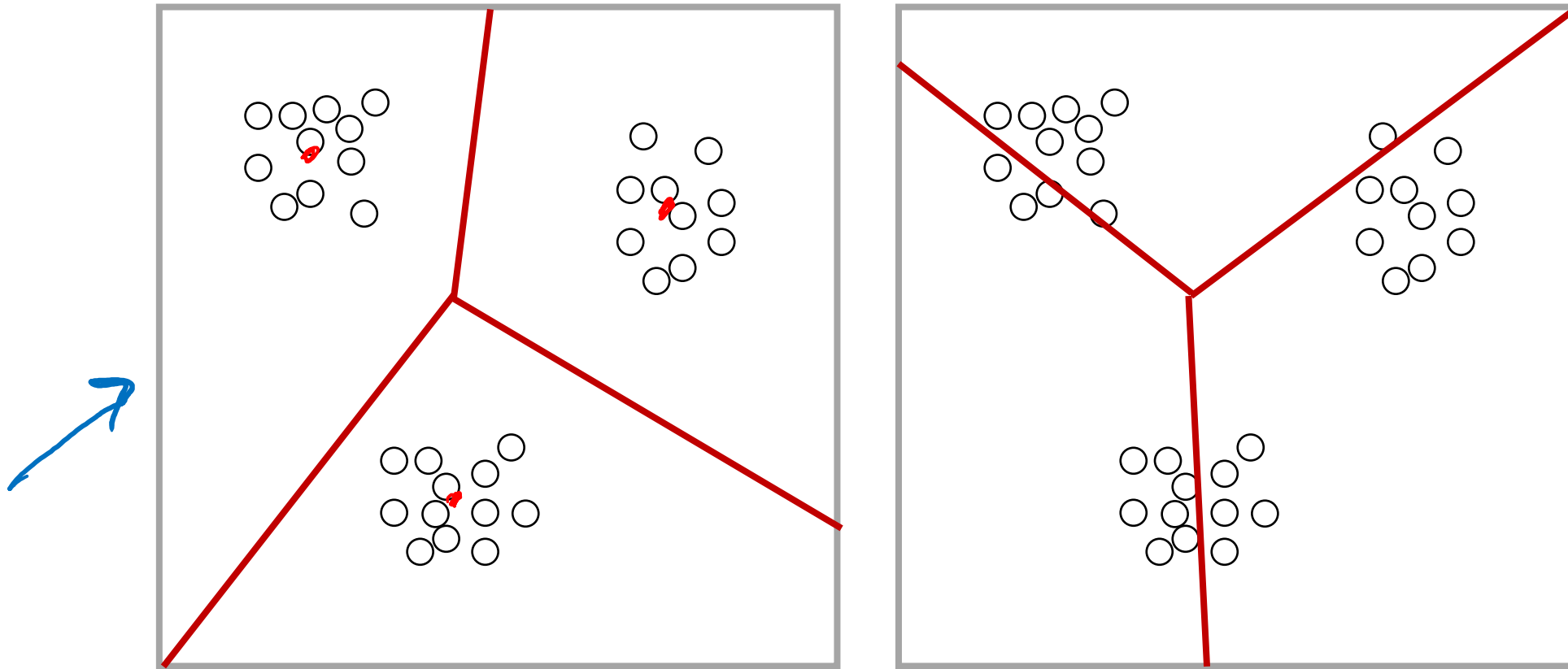
Instructor: Pat Virtue

Slide credits: Aarti Singh, Eric Xing, Carlos Guestrin

# K-means Optimization

Question: Which of these partitions is “better”?

*min*



# K-means Optimization

Input:  $\vec{x}^{(1)}, \dots, \vec{x}^{(N)}$   $K$   $\vec{x}^{(i)} \in \mathbb{R}^M$   
 $\rightarrow \vec{z}^{(1)}, \dots, \vec{z}^{(N)}$

Output:  $\{\vec{c}_1, \dots, \vec{c}_k\}$   $\vec{c}_k \in \mathbb{R}^M$

$$C = \operatorname{argmin}_C \sum_{i=1}^N \min_{j \in \{1 \dots K\}} \|\vec{x}^{(i)} - \vec{c}_j\|_2^2$$

$$C, \vec{z} = \operatorname{argmin}_C \sum_{i=1}^N \min_{z^{(i)} \in \{1 \dots K\}} \|\vec{x}^{(i)} - \vec{c}_{z^{(i)}}\|_2^2$$

$$= \operatorname{argmin}_{C, \vec{z}} \sum_{i=1}^N \|\vec{x}^{(i)} - \vec{c}_{z^{(i)}}\|_2^2$$

$$z^{(i)} \in \{1 \dots K\}$$

# K-means Optimization

Alternating minimization

$$a) \vec{z} = \operatorname{argmin}_{\vec{z}} \sum_{i=1}^N \|x^i - \underline{c}_z\|_2^2$$

$$b) C = \operatorname{argmin}_C \sum_{i=1}^N \|x^i - c_z\|_2^2$$

---

$$\begin{array}{l|l} a) z^{(1)} = \operatorname{argmin}_k \|x^{(1)} - c_k\|_2^2 & b) \vec{c}_1 = \operatorname{argmin}_{c_1} \sum_{i=z^{(1)}=1} \|x^{(i)} - c_1\|_2^2 \\ \vdots & \vdots \\ z^{(N)} = " & \vec{c}_k = \operatorname{argmin}_{c_k} \sum_{i=z^{(i)}=k} \|x^{(i)} - c_k\|_2^2 \end{array}$$

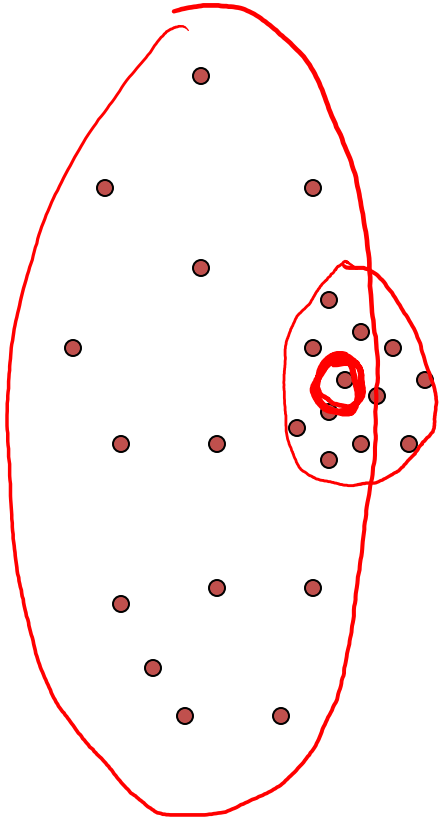
# Piazza Poll 1

[True/False] The alternating minimization algorithm will find the global minimum of the k-means objective.

$$\mathbf{C}, \mathbf{z} = \operatorname{argmin}_{\mathbf{C}, \mathbf{z}} \sum_{i=1}^N \|\mathbf{x}^{(i)} - \mathbf{c}_{z^{(i)}}\|_2^2$$

- A. Calamity
- B. True
- ☒ C. False

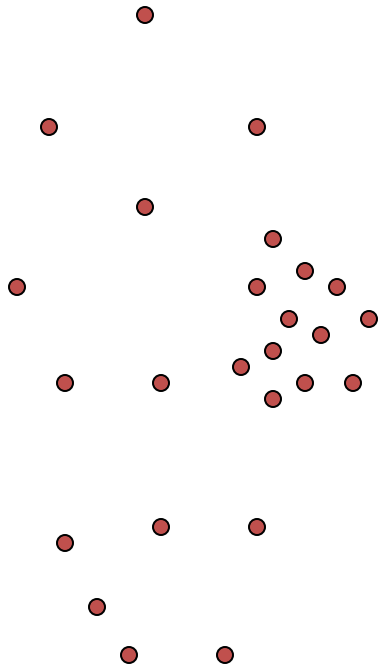
# (One) bad case for K-means



- Clusters may overlap
- Some clusters may be “wider” than others
- Clusters may not be linearly separable

$$p(z_{n,k}=1 | x_n)$$

# (One) bad case for K-means



- Clusters may overlap
- Some clusters may be “wider” than others
- Clusters may not be linearly separable

# Partitioning Algorithms

- K-means
  - **hard assignment**: each object belongs to only one cluster
- Mixture modeling
  - **soft assignment**: probability that an object belongs to a cluster

Generative approach

$$p(x) = \sum_y p(x, y)$$

$$= \sum_y p(x|y) p(y)$$

# Gaussian Mixture Model

Mixture of K Gaussian distributions: (Multi-modal distribution)

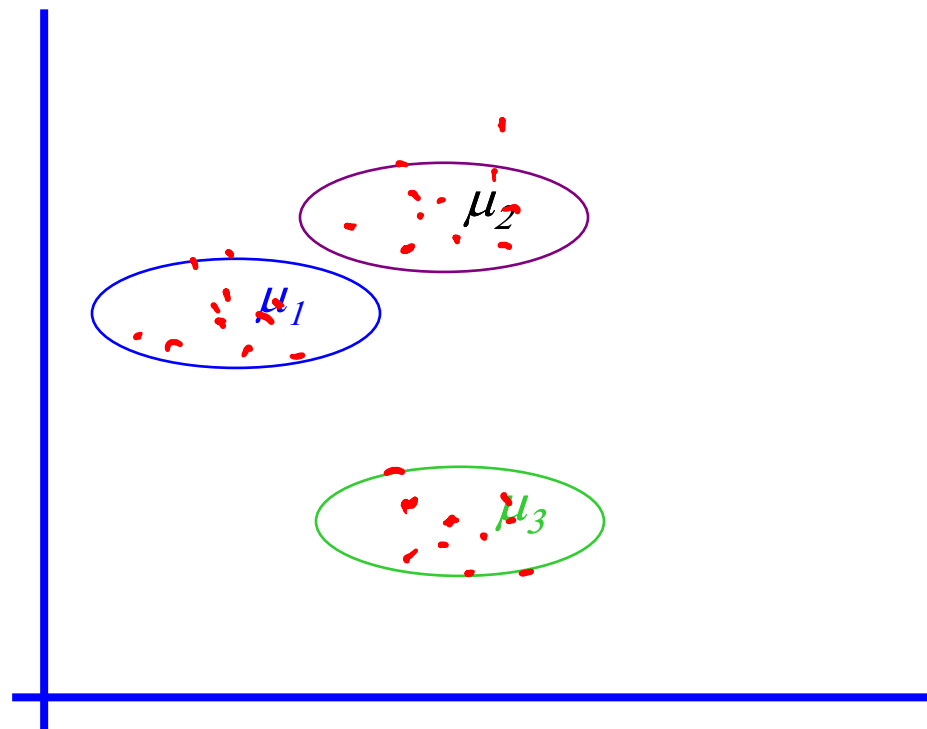
$$p(x|y=i) \sim N(\mu_i, \sigma^2 I)$$

$$\underline{p(x)} = \sum_i p(x|y=i) P(y=i)$$

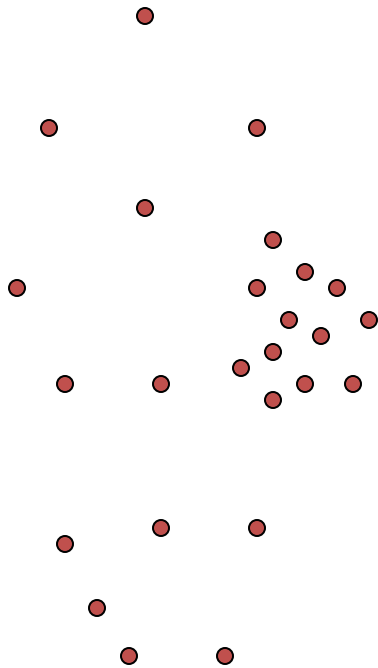
↓  
Mixture  
component

↓  
Mixture  
proportion

$$x \rightarrow x \sim p(x|y)$$



# (One) bad case for K-means



- Clusters may overlap
- Some clusters may be “wider” than others
- Clusters may not be linearly separable

# General GMM

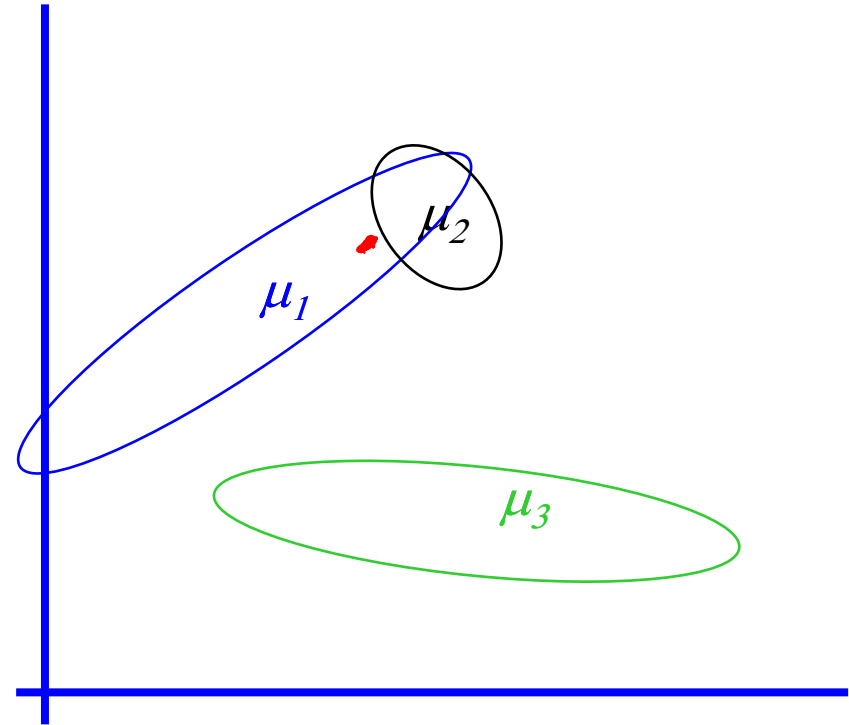
GMM – Gaussian Mixture Model (Multi-modal distribution)

$$p(x/y=i) \sim N(\mu_i, \Sigma_i)$$

$\rightarrow p(x) = \sum_i p(x/y=i) P(y=i)$

$\downarrow$   $\downarrow$

**Mixture component** **Mixture proportion**



# General GMM

GMM – Gaussian Mixture Model (Multi-modal distribution)

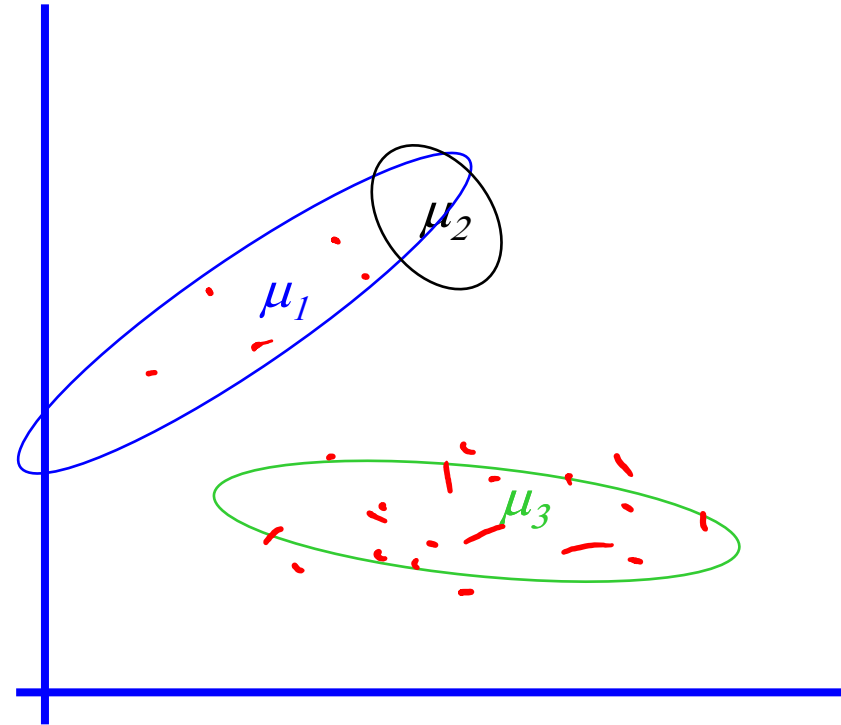
- There are  $k$  components
- Component  $i$  has an associated mean vector  $\mu_i$
- Each component generates data from a Gaussian with mean  $\mu_i$  and covariance matrix  $\Sigma_i$

Each data point is generated according to the following recipe:

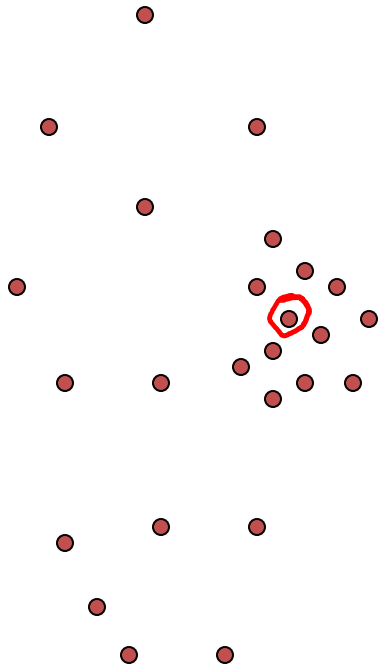
1) Pick a component at random:

→ Choose component  $i$  with probability  $P(y=i)$  ←

→ 2) Datapoint  $x \sim N(\mu_i, \Sigma_i)$



# (One) bad case for K-means



- Clusters may overlap
- Some clusters may be “wider” than others
- Clusters may not be linearly separable

$$p(y=1|x)$$

$$p(y=2|x)$$

$$p(y=3|x)$$

$$p(y=1)$$

# General GMM

GMM – Gaussian Mixture Model (Multi-modal distribution)

$$p(x|y=i) \sim N(\mu_i, \Sigma_i)$$

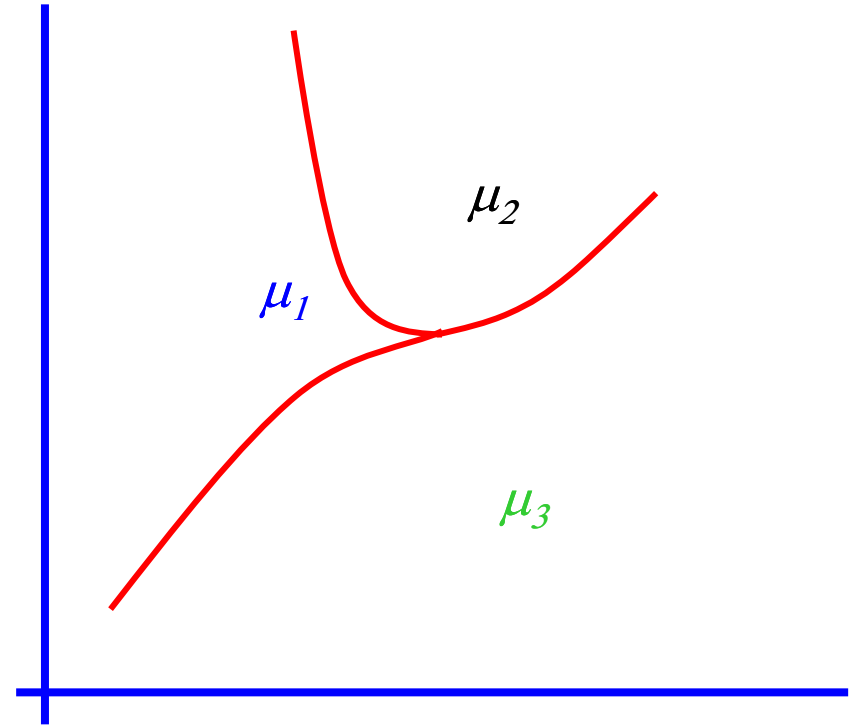
Decision boundary when probabilities are equal:

$$\log \frac{P(y=i|x)}{P(y=j|x)}$$

$$= \log \frac{p(x|y=i)P(y=i)}{p(x|y=j)P(y=j)}$$

$$= x^T \mathbf{W}x + \mathbf{w}^T x$$

Depend on  $\mu_1, \mu_2, \dots, \mu_K, \Sigma_1, \Sigma_2, \dots, \Sigma_K, P(y=1), \dots, P(y=K)$



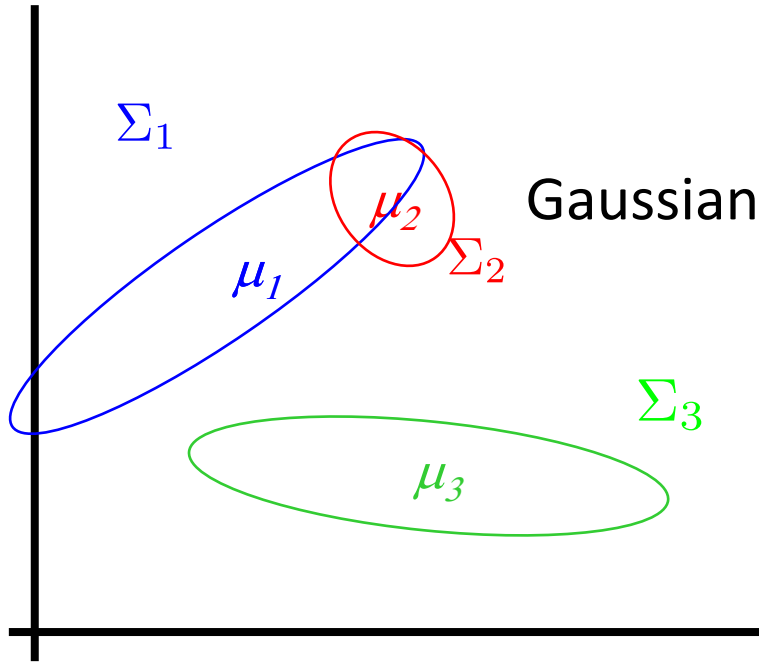
**“Quadratic Decision boundary”** – second-order terms don’t cancel out

# Learning General GMM

$$x_1, \dots, x_m \sim p(x) = \sum_{i=1}^k p(x|Y = i) P(Y = i)$$

↓  
**Mixture  
component**

↓  
**Mixture  
proportion,  $p_i$**



Gaussian mixture model

$$p(x|Y = i) \sim \mathcal{N}(\mu_i, \Sigma_i)$$

**Parameters:**  $\{p_i, \mu_i, \Sigma_i\}_{i=1}^K$

- How to estimate parameters? Max Likelihood  
But don't know labels  $Y$

# Learning General GMM

$$P(D|\theta)$$
$$p(x) \rightarrow p(x|\theta)$$

Maximize marginal likelihood:

$$\begin{aligned} \operatorname{argmax} \prod_j P(x_j) &= \operatorname{argmax} \prod_j \sum_{i=1}^K P(y_j=i, x_j) \\ &= \operatorname{argmax} \prod_j \sum_{i=1}^K \underline{P(y_j=i)} p(x_j | y_j=i) \end{aligned}$$

$$p(x) = \sum_y p(x, y)$$
$$p(x, y) = p(y) p(x|y)$$

$P(y_j=i) = P(y=i)$  Mixture component  $i$  is chosen with prob  $P(y = i)$

$$= \operatorname{argmax} \prod_{j=1}^m \sum_{i=1}^k P(y = i) \underbrace{\frac{1}{\sqrt{\det(\Sigma_i)}} \exp\left[-\frac{1}{2}(x_j - \mu_i)^T \Sigma_i^{-1} (x_j - \mu_i)\right]}$$

How do we find the  $\mu_i, \Sigma_i$  s and  $P(y=i)$ s which give max. marginal likelihood?

\* Set  $\frac{\partial}{\partial \mu_i} \log \text{Prob}(\dots) = 0$  and solve for  $\mu_i$ 's. Non-linear not-analytically solvable

\* Use gradient descent: Doable, but often slow

# MLE Optimization

$$p(D|\theta)$$

$$\begin{aligned} \text{MLE } x \\ \hat{\theta}_{ML} &= \underset{\theta}{\operatorname{argmax}} p(x|\theta) \\ &= \underset{\theta}{\operatorname{argmax}} \prod_{i=1}^N p(x^{(i)}|\theta) \\ &= \underset{\theta}{\operatorname{argmax}} \log \prod_{i=1}^N p(x^{(i)}|\theta) \\ &= \underset{\theta}{\operatorname{argmax}} \sum_{i=1}^N \log \underline{p(x^{(i)}|\theta)} \end{aligned}$$

$$\begin{aligned} \text{MLE } y|x \\ \hat{\theta}_{ML} &= \underset{\theta}{\operatorname{argmax}} \prod_{i=1}^N p(y^{(i)}|x^{(i)}|\theta) \\ &= \underset{\theta}{\operatorname{argmax}} \prod_{i=1}^N p(x^{(i)}|y^{(i)}|\theta) p(y^{(i)}|\theta) \\ &= \underset{\theta}{\operatorname{argmax}} \sum_{i=1}^N \log [p(x|y|\theta) p(y|\theta)] \\ \hat{\theta}_{ML} &= \underset{\theta}{\operatorname{argmax}} \sum_{i=1}^N [\log p(x^i|y^i|\theta) + \log p(y^i|\theta)] \end{aligned}$$

2

# MLE Optimization

$$\begin{aligned}\hat{\theta}_{ML} &= \operatorname{argmax} \prod p(x) \\ &= \operatorname{argmax} \prod_{i=1}^N \sum_z p(x, z) \\ &= \log \prod_i \sum_z p(x, z) \\ &= \sum_{i=1}^N \log \sum_z \underbrace{p(x, z)}_{p(z)p(x|z)}\end{aligned}$$

# GMM vs. k-means

Maximize marginal likelihood:

$$\begin{aligned}\operatorname{argmax} \prod_j P(x_j) &= \operatorname{argmax} \prod_j \sum_{i=1}^K P(y_j=i, x_j) \\ &= \operatorname{argmax} \prod_j \sum_{i=1}^K P(y_j=i) p(x_j | y_j=i)\end{aligned}$$

- What happens if we assume **Hard assignment**?

$$\begin{aligned}P(y_j = i) &= 1 \text{ if } i = C(j) \\ &= 0 \text{ otherwise}\end{aligned}$$

$$\operatorname{argmax} \prod_j P(x_j) = \operatorname{argmax} \prod_j p(x_j | y_j=C(j))$$

Same as  
k-means!

$$= \operatorname{argmax} \prod_{j=1}^n \exp\left(\frac{-1}{2\sigma^2} \|x_j - \mu_{C(j)}\|^2\right)$$

$$= \operatorname{argmin} \sum_{j=1}^n \|x_j - \mu_{C(j)}\|^2 = \operatorname{argmin}_{\mu, C} F(\mu, C)$$

C z  
- z

# Expectation-Maximization (EM)

A general algorithm to deal with hidden data, but we will study it in the context of unsupervised learning (hidden labels) first

- No need to choose step size as in Gradient methods.

- EM is an Iterative algorithm with two linked steps:

→ E-step: fill-in hidden data (Y) using inference  $E[Y|x] \leftarrow \theta^t$

M-step: apply standard MLE/MAP method to estimate parameters

$$\{p_i, \mu_i, \Sigma_i\}_{i=1}^k \quad \theta^{t+1} \leftarrow E[Y|x]$$

- We will see that this procedure monotonically improves the likelihood (or leaves it unchanged). Thus it always converges to a local optimum of the likelihood.

$$p(x|\theta^t) \leq p(x|\theta^{(t+1)}) \leq p(x|\theta^{(t+2)})$$

# EM for spherical, same variance GMMs

## E-step

Compute “expected” classes of all datapoints for each class

$$p(y_k^{(i)} = 1 \mid x, \theta^{(t)}) \leftarrow \theta^{(t)}$$

## M-step

Compute Max. like  $\mu$  given our data’s class membership distributions (weights)

$$\theta^{(t+1)} \leftarrow p(y_k^{(i)} = 1 \mid x, \theta^{(t)})$$

Iterate.

# EM for spherical, same variance GMMs

## E-step

Compute “expected” classes of all datapoints for each class

$$\rightarrow P(y = i | x_j, \mu_1 \dots \mu_k) \propto \exp\left(-\frac{1}{2\sigma^2} \|x_j - \mu_i\|^2\right) P(y = i)$$

In K-means “E-step”  
we do hard assignment

EM does soft assignment

## M-step

Compute Max. like  $\mu$  given our data’s class membership distributions (weights)

$$\mu_i = \frac{\sum_{j=1}^m P(y = i | x_j) x_j}{\sum_{j=1}^m P(y = i | x_j)}$$

$$\frac{\sum_{j=1}^m z_i^{(j)} x_j}{\sum_{j=1}^m z_i^{(j)}} \leftarrow \frac{\text{sum } x \text{ in } i}{N \text{ points in } i}$$

Exactly same as MLE with weighted data

Iterate.

# EM for general GMMs

Iterate. On iteration  $t$  let our estimates be

$$\lambda_t = \{ \underline{\mu_1^{(t)}}, \underline{\mu_2^{(t)}} \dots \underline{\mu_k^{(t)}}, \underline{\Sigma_1^{(t)}}, \underline{\Sigma_2^{(t)}} \dots \underline{\Sigma_k^{(t)}}, \underline{p_1^{(t)}}, \underline{p_2^{(t)}} \dots \underline{p_k^{(t)}} \}$$

$p_i^{(t)}$  is shorthand for estimate of  $P(y=i)$  on  $t$ 'th iteration

~~$P(x|x)$~~

## E-step

Compute “expected” classes of all datapoints for each class

$$P(y = i | x_j, \lambda_t) \propto p_i^{(t)} P(x_j | \mu_i^{(t)}, \Sigma_i^{(t)})$$

Just evaluate a Gaussian at  $x_j$

## M-step

Compute MLEs given our data's class membership distributions (weights)

$$\mu_i^{(t+1)} = \frac{\sum_j P(y = i | x_j, \lambda_t) x_j}{\sum_j P(y = i | x_j, \lambda_t)} \quad \Sigma_i^{(t+1)} = \frac{\sum_j P(y = i | x_j, \lambda_t) (x_j - \mu_i^{(t+1)}) (x_j - \mu_i^{(t+1)})^T}{\sum_j P(y = i | x_j, \lambda_t)}$$

$$p_i^{(t+1)} = \frac{\sum_j P(y = i | x_j, \lambda_t)}{m}$$

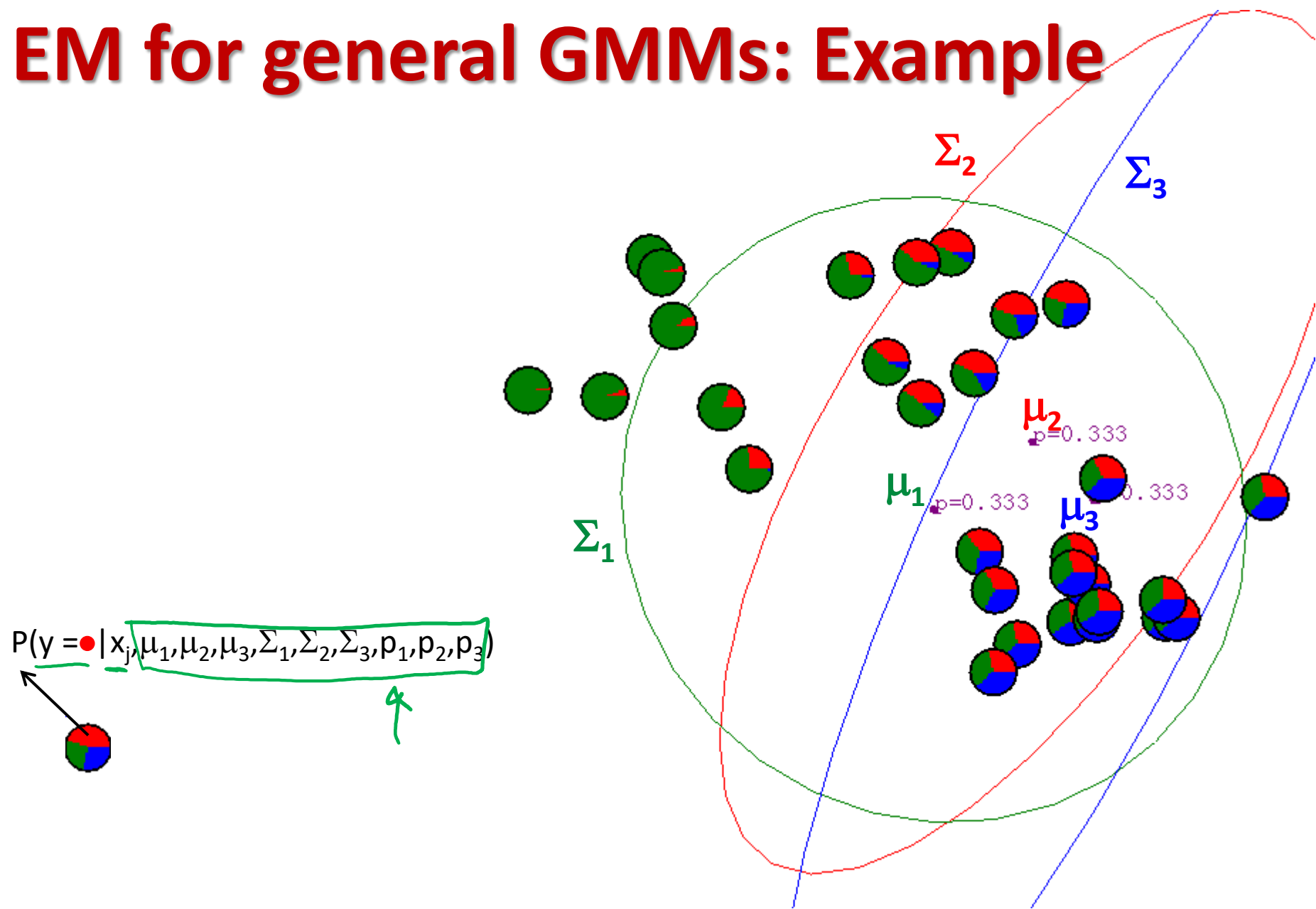
$m = \text{\#data points}$

# EM Convergence

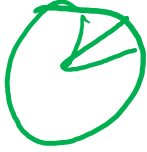
$$p(x | \lambda^{(t)}) \leq p(x | \lambda^{(t+1)})$$

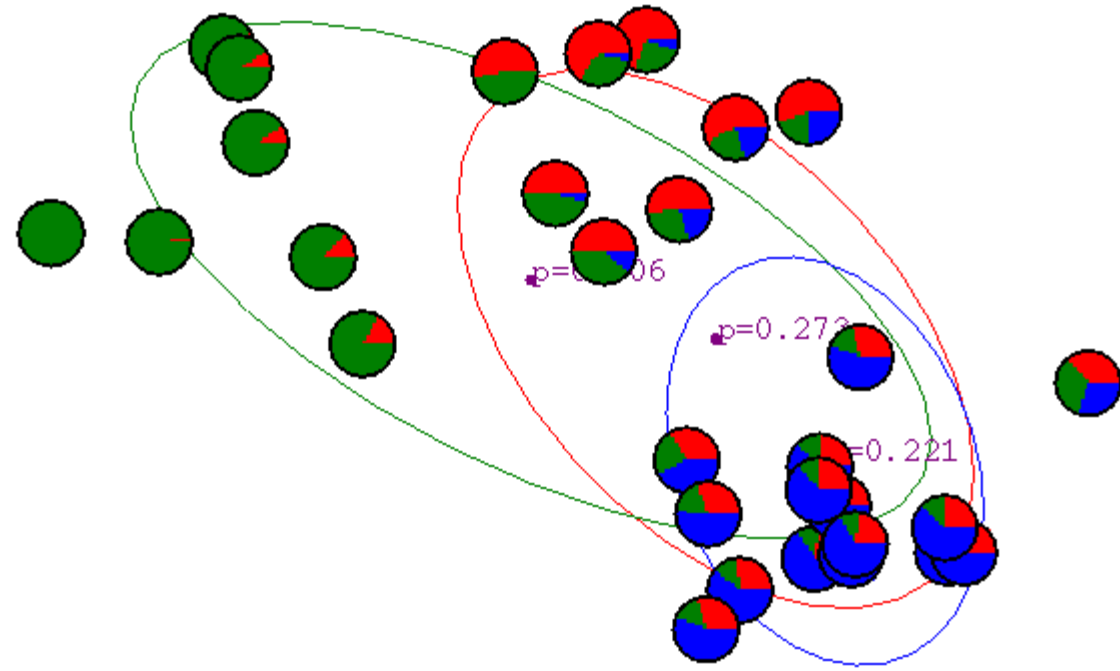
$$p(x | \lambda^*)$$

# EM for general GMMs: Example

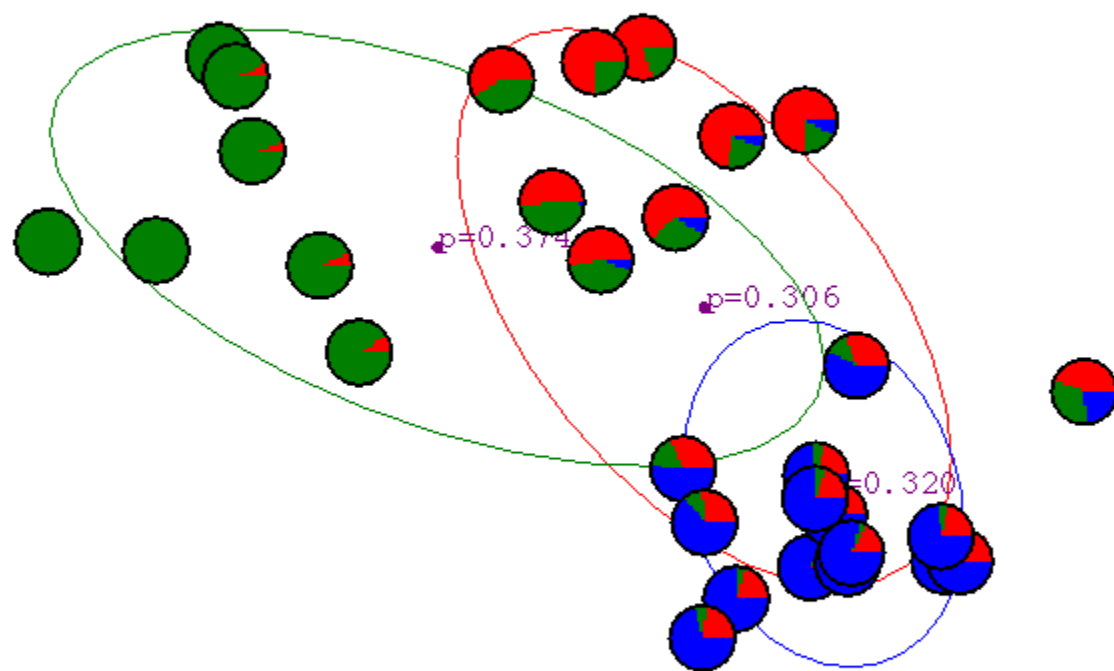


# After 1<sup>st</sup> iteration

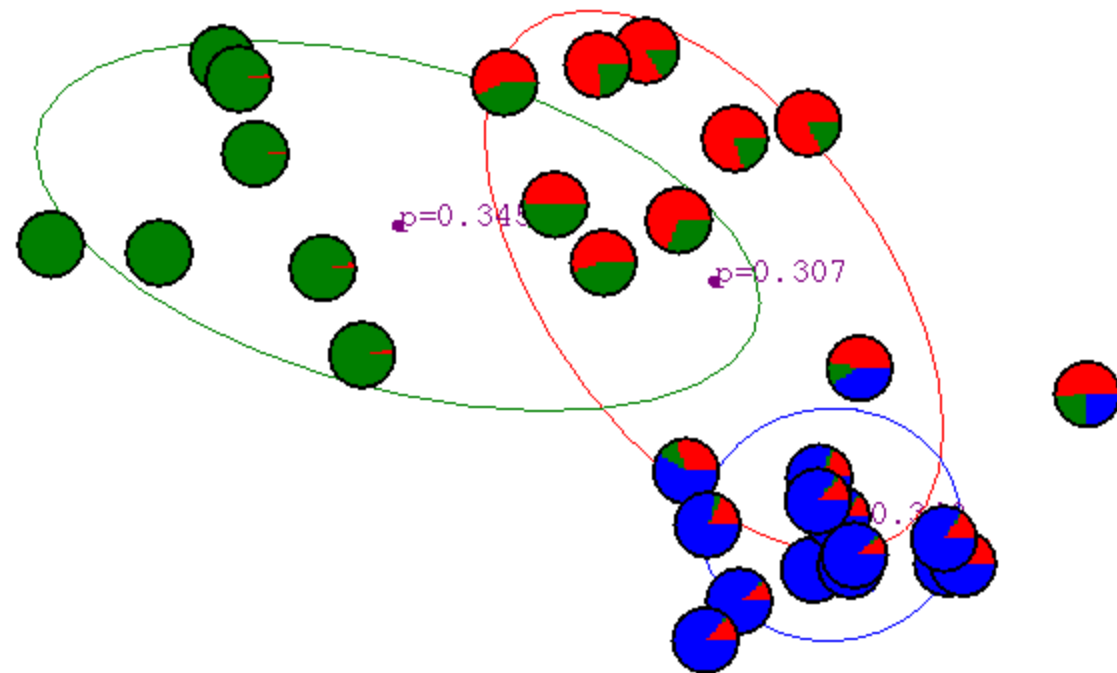
$$p(y|x, \lambda)$$




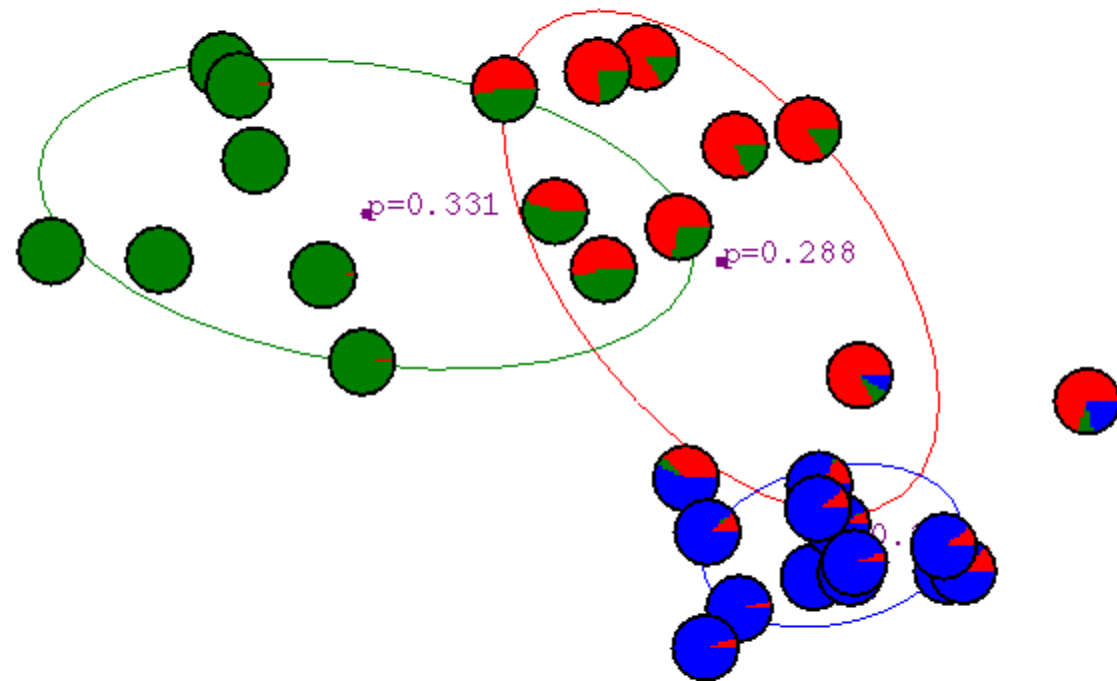
## After 2<sup>nd</sup> iteration



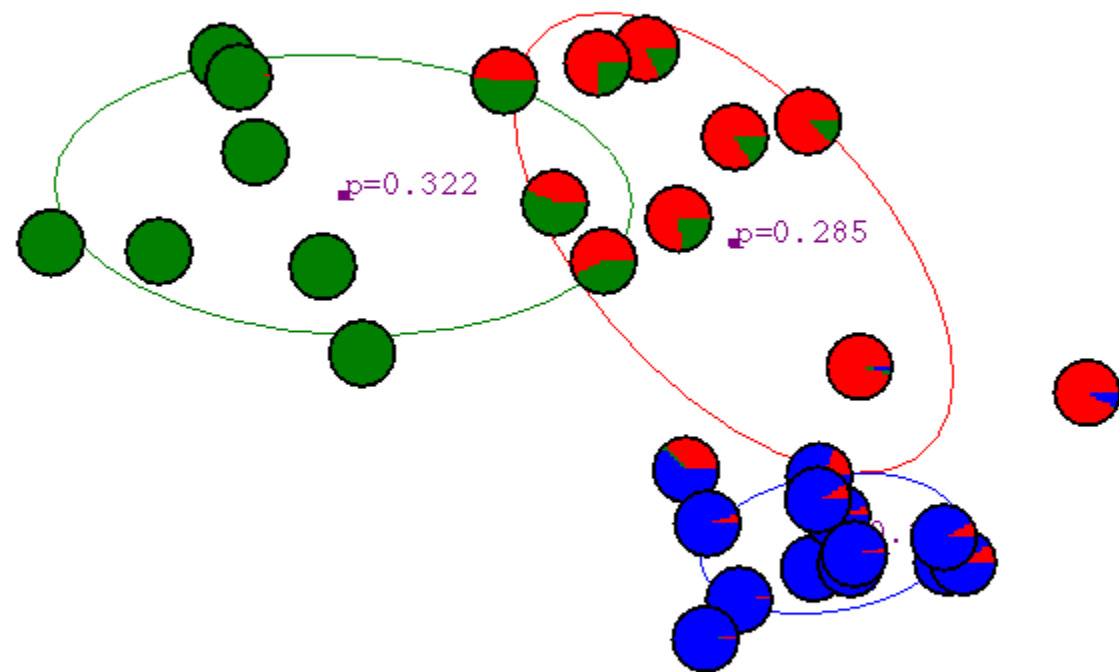
# After 3<sup>rd</sup> iteration



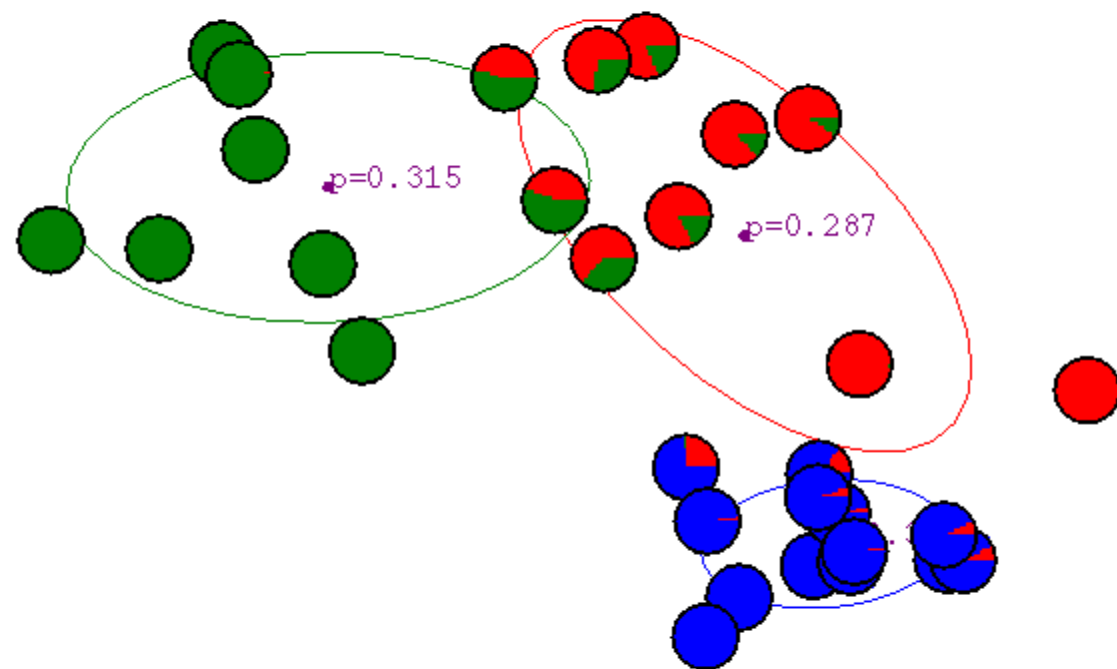
# After 4<sup>th</sup> iteration



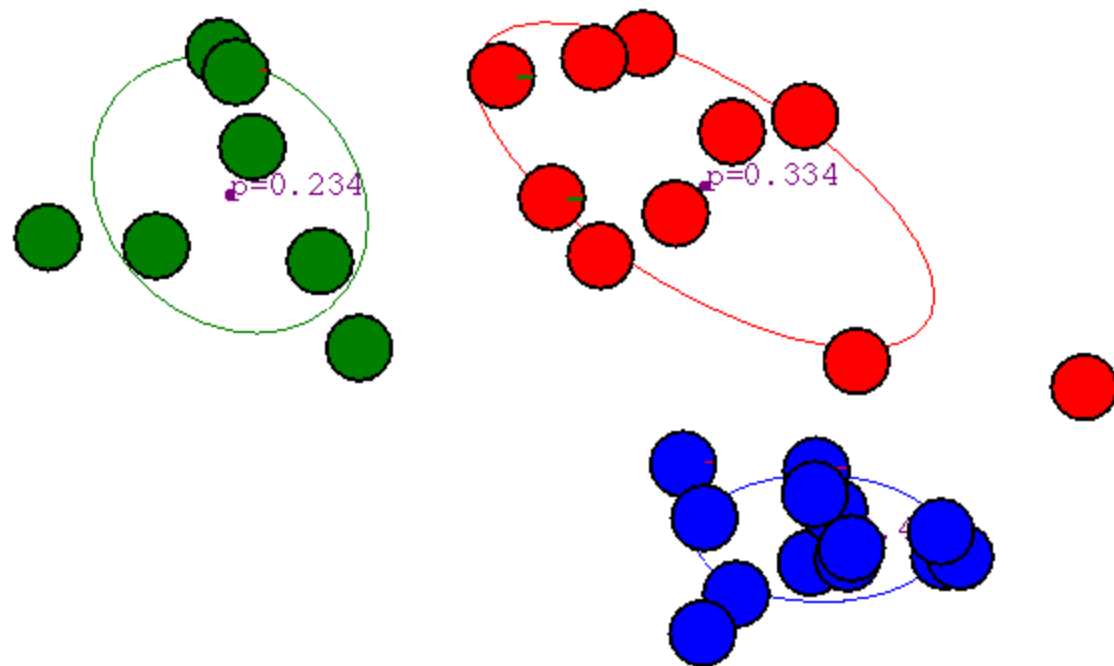
# After 5<sup>th</sup> iteration



# After 6<sup>th</sup> iteration



# After 20<sup>th</sup> iteration



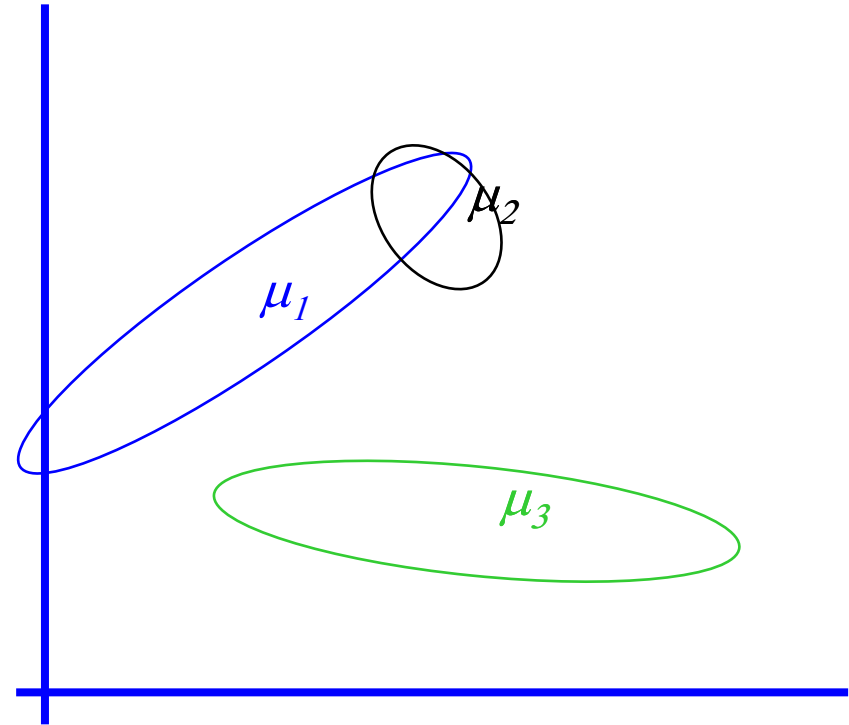
# General GMM

GMM – Gaussian Mixture Model (Multi-modal distribution)

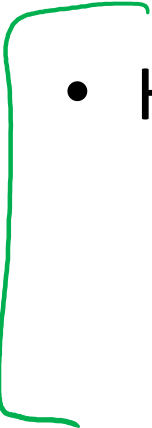
$$p(x) = \sum_i p(x|y=i) P(y=i)$$

↓                      ↓  
**Mixture**            **Mixture**  
**component**       **proportion**

$$p(x|y=i) \sim N(\mu_i, \Sigma_i)$$



# What you need to know...

- 
- Hierarchical clustering algorithms
    - Single-linkage
    - Complete-linkage
    - Centroid-linkage
    - Average-linkage
  - Partition based clustering algorithms
    - K-means
      - Coordinate descent
      - Seeding
      - Choosing K
    - Mixture models
      - EM algorithm

# General EM algorithm

Marginal likelihood –  $\mathbf{x}$  is observed,  $\mathbf{z}$  is missing:

$$\begin{aligned}\log P(\mathbf{D}; \theta) &= \log \prod_{j=1}^m P(\mathbf{x}_j \mid \theta) \\ &= \sum_{j=1}^m \log P(\mathbf{x}_j \mid \theta) \\ &= \sum_{j=1}^m \log \sum_{\mathbf{z}} P(\mathbf{x}_j, \mathbf{z} \mid \theta)\end{aligned}$$

$$\mathbf{D} = \{\mathbf{x}_j\}_{j=1}^m$$

$\theta$  - model parameter(s)

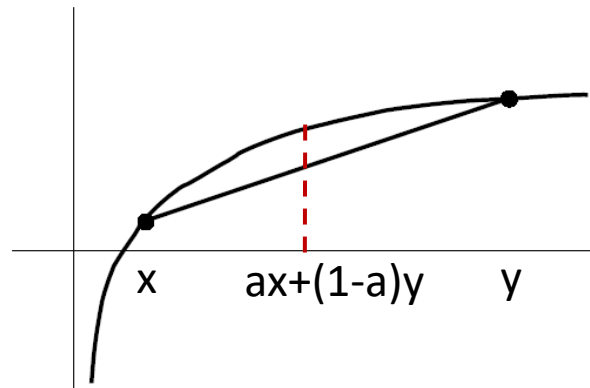
How to maximize marginal likelihood using EM?

# Lower-bound on marginal likelihood

$$\begin{aligned}\log P(D; \theta) &= \sum_{j=1}^m \log \sum_{\mathbf{z}} P(\mathbf{x}_j, \mathbf{z} \mid \theta) \\ &= \sum_{j=1}^m \log \underbrace{\sum_{\mathbf{z}} Q(\mathbf{z} \mid \mathbf{x}_j)}_{P(\mathbf{z})} \underbrace{\frac{P(\mathbf{z}, \mathbf{x}_j \mid \theta)}{Q(\mathbf{z} \mid \mathbf{x}_j)}}_{f(\mathbf{z})}\end{aligned}$$

Variational  
approach

Jensen's inequality:  $\log \sum_{\mathbf{z}} P(\mathbf{z}) f(\mathbf{z}) \geq \sum_{\mathbf{z}} P(\mathbf{z}) \log f(\mathbf{z})$



$\log$ : concave function

$$\log(ax+(1-a)y) \geq a \log(x) + (1-a) \log(y)$$

# Lower-bound on marginal likelihood

$$\begin{aligned}\log P(D; \theta) &= \sum_{j=1}^m \log \sum_{\mathbf{z}} P(\mathbf{x}_j, \mathbf{z} \mid \theta) \\ &= \sum_{j=1}^m \log \sum_{\mathbf{z}} \underbrace{Q(\mathbf{z} \mid \mathbf{x}_j)}_{P(\mathbf{z})} \underbrace{\frac{P(\mathbf{z}, \mathbf{x}_j \mid \theta)}{Q(\mathbf{z} \mid \mathbf{x}_j)}}_{f(\mathbf{z})}\end{aligned}$$

**Jensen's inequality:**  $\log \sum_{\mathbf{z}} P(\mathbf{z}) f(\mathbf{z}) \geq \sum_{\mathbf{z}} P(\mathbf{z}) \log f(\mathbf{z})$

$$\geq \sum_{j=1}^m \sum_{\mathbf{z}} Q(\mathbf{z} \mid \mathbf{x}_j) \log \frac{P(\mathbf{z}, \mathbf{x}_j \mid \theta)}{Q(\mathbf{z} \mid \mathbf{x}_j)} =: \underbrace{F(\theta, Q)}$$

# EM as Coordinate Ascent

$$\log P(D; \theta) \geq F(\theta, Q)$$

**E-step:** Fix  $\theta$ , maximize  $F$  over  $Q$

$$Q^{t+1} = \arg \max_Q F(\theta^t, Q)$$

**M-step:** Fix  $Q$ , maximize  $F$  over  $\theta$

$$\theta^{t+1} = \arg \max_{\theta} F(\theta, Q^{t+1})$$

# Convergence of EM

$$\underbrace{\log P(D; \theta)} \geq \underbrace{F(\theta, Q)}$$

**E-step:** Fix  $\theta$ , maximize  $F$  over  $Q$

$$Q^{t+1} = \arg \max_Q F(\theta^t, Q)$$

E-step maximizes lower bound  $F$  on marginal likelihood => doesn't decrease the marginal likelihood

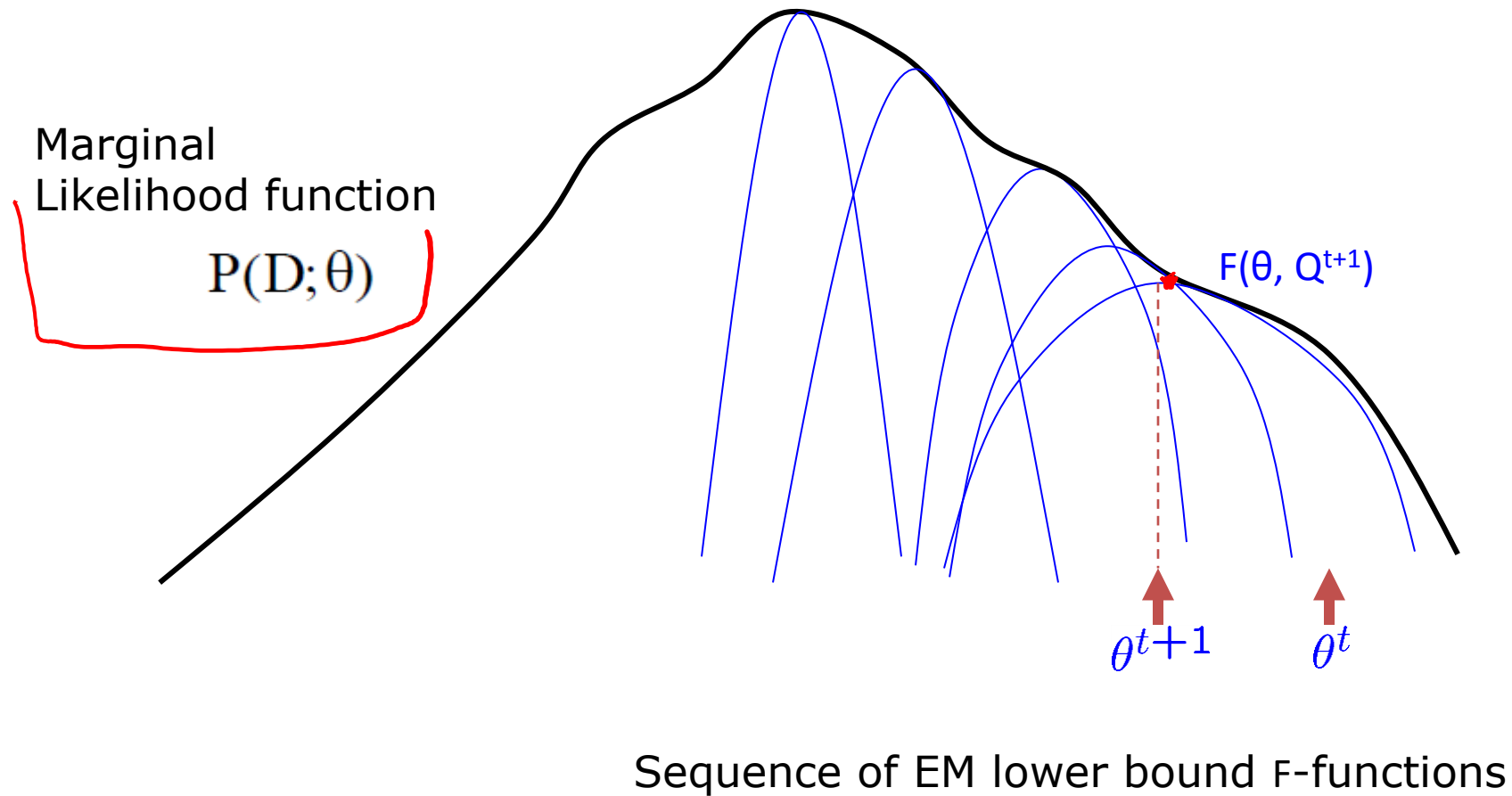
**M-step:** Fix  $Q$ , maximize  $F$  over  $\theta$

$$\theta^{t+1} = \arg \max_{\theta} F(\theta, Q^{t+1})$$

M-step maximizes lower bound  $F$  on marginal likelihood => doesn't decrease the marginal likelihood

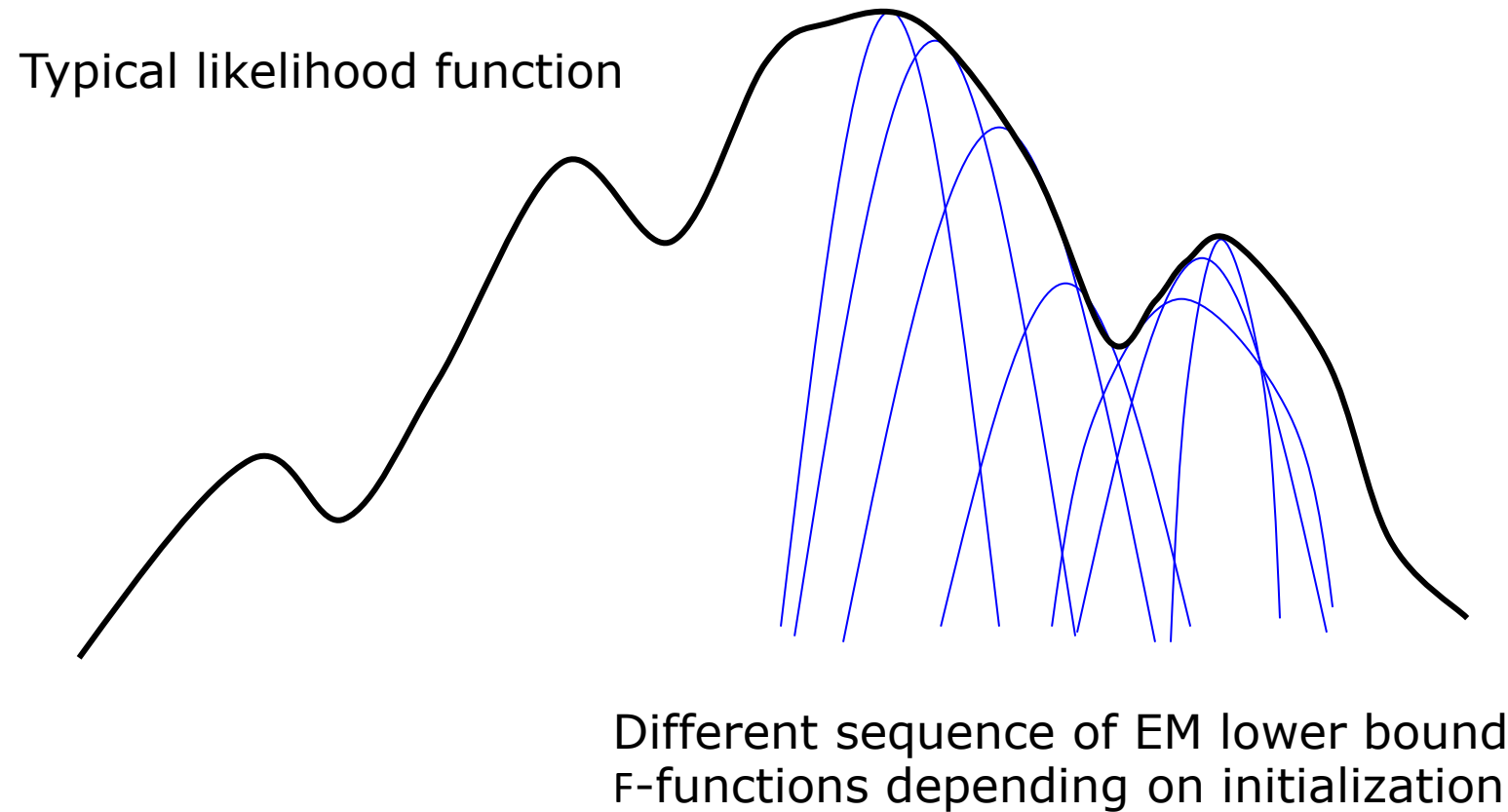
Since marginal likelihood is bounded, convergence follows!

# Convergence of EM



**EM monotonically converges to a local maximum of likelihood !**

# EM & Local Maxima



**Use multiple, randomized initializations in practice**

# EM as Coordinate Ascent

$$\log P(D; \theta) \geq F(\theta, Q)$$

**E-step:** Fix  $\theta$ , maximize  $F$  over  $Q$

$$Q^{t+1} = \arg \max_Q F(\theta^t, Q)$$

**M-step:** Fix  $Q$ , maximize  $F$  over  $\theta$

$$\theta^{t+1} = \arg \max_{\theta} F(\theta, Q^{t+1})$$

# E step

$$\log P(\mathbf{D}; \theta) \geq F(\theta, Q)$$

**E-step:** Fix  $\theta$ , maximize  $F$  over  $Q$

$$\begin{aligned} \log P(\mathbf{D}; \theta^{(t)}) &\geq F(\theta^{(t)}, Q) = \sum_{j=1}^m \sum_{\mathbf{z}} Q(\mathbf{z} | \mathbf{x}_j) \log \frac{P(\mathbf{z}, \mathbf{x}_j | \theta^{(t)})}{Q(\mathbf{z} | \mathbf{x}_j)} \\ &= \sum_{j=1}^m \sum_{\mathbf{z}} Q(\mathbf{z} | \mathbf{x}_j) \log \frac{P(\mathbf{z} | \mathbf{x}_j, \theta^{(t)}) P(\mathbf{x}_j | \theta^{(t)})}{Q(\mathbf{z} | \mathbf{x}_j)} \\ &= \underbrace{\sum_{j=1}^m \sum_{\mathbf{z}} Q(\mathbf{z} | \mathbf{x}_j) \log \frac{P(\mathbf{z} | \mathbf{x}_j, \theta^{(t)})}{Q(\mathbf{z} | \mathbf{x}_j)}}_{-KL(Q(\mathbf{z}|\mathbf{x}_j), P(\mathbf{z}|\mathbf{x}_j, \theta^{(t)}))} + \underbrace{\sum_{j=1}^m \sum_{\mathbf{z}} Q(\mathbf{z} | \mathbf{x}_j) \log P(\mathbf{x}_j | \theta^{(t)})}_{\log P(\mathbf{D}; \theta^{(t)})} \end{aligned}$$

**KL divergence between two distributions**

# E step

$$\log P(\mathbf{D}; \theta) \geq F(\theta, Q)$$

**E-step:** Fix  $\theta$ , maximize  $F$  over  $Q$

$$\begin{aligned} \log P(\mathbf{D}; \theta^{(t)}) &\geq F(\theta^{(t)}, Q) = \sum_{j=1}^m \sum_{\mathbf{z}} Q(\mathbf{z} | \mathbf{x}_j) \log \frac{P(\mathbf{z}, \mathbf{x}_j | \theta^{(t)})}{Q(\mathbf{z} | \mathbf{x}_j)} \\ &= \sum_{j=1}^m -KL(Q(\mathbf{z} | \mathbf{x}_j), P(\mathbf{z} | \mathbf{x}_j, \theta^{(t)})) + \log P(\mathbf{D}; \theta^{(t)}) \end{aligned}$$

**KL  $\geq 0$ , above expression is maximized if KL divergence = 0**

**KL(Q,P) = 0 iff Q = P**

**Therefore,**

$$\textbf{E step: } Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) = P(\mathbf{z} | \mathbf{x}_j, \theta^{(t)})$$

# E step

$$\log \mathbf{P}(\mathbf{D}; \theta) \geq F(\theta, Q)$$

**E-step:** Fix  $\theta$ , maximize  $F$  over  $Q$

$$\log \mathbf{P}(\mathbf{D}; \theta^{(t)}) \geq F(\theta^{(t)}, Q) = \sum_{j=1}^m -KL(Q(\mathbf{z}|\mathbf{x}_j), P(\mathbf{z}|\mathbf{x}_j, \theta^{(t)})) + \log \mathbf{P}(\mathbf{D}; \theta^{(t)})$$

$$\rightarrow Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) = P(\mathbf{z} | \mathbf{x}_j, \theta^{(t)})$$

Compute probability of missing data  $\mathbf{z}$  given current choice of  $\theta$

Re-aligns  $F$  with marginal likelihood !!

$$F(\theta^{(t)}, Q^{(t+1)}) = \log \mathbf{P}(\mathbf{D}; \theta^{(t)})$$

# M step

$$\log P(D; \theta) \geq F(\theta, Q)$$

**M-step:** Fix  $Q$ , maximize  $F$  over  $\theta$

$$\log P(D; \theta) \geq F(\theta, Q^{(t+1)}) = \sum_{j=1}^m \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} \mid \mathbf{x}_j) \log \frac{P(\mathbf{z}, \mathbf{x}_j \mid \theta)}{Q^{(t+1)}(\mathbf{z} \mid \mathbf{x}_j)}$$

$$= \sum_{j=1}^m \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} \mid \mathbf{x}_j) \log P(\mathbf{z}, \mathbf{x}_j \mid \theta) + \sum_{j=1}^m \underbrace{H(Q^{(t+1)}(\mathbf{z} \mid \mathbf{x}_j))}_{\text{Fixed (Independent of } \theta \text{)}}$$

$$\sum_{\mathbf{z}} \underbrace{\sum_{j=1}^m \log P(\mathbf{z}, \mathbf{x}_j \mid \theta) Q^{(t+1)}(\mathbf{z} \mid \mathbf{x}_j)}_{\text{Expected log likelihood wrt } Q}$$

Log likelihood if  $\mathbf{z}$  was known

# M step

$$\log P(\mathbf{D}; \theta) \geq F(\theta, Q)$$

**M-step:** Fix  $Q$ , maximize  $F$  over  $\theta$

$$\log P(\mathbf{D}; \theta) \geq F(\theta, Q^{(t+1)}) = \sum_{\mathbf{z}} \sum_{j=1}^m \log P(\mathbf{z}, \mathbf{x}_j | \theta) Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) + \sum_{j=1}^m \underbrace{H(Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j))}_{\text{Fixed (Independent of } \theta \text{)}}$$

$$\theta^{(t+1)} \leftarrow \arg \max_{\theta} \underbrace{\sum_{\mathbf{z}} \sum_{j=1}^m \log P(\mathbf{z}, \mathbf{x}_j | \theta) Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j)}_{\text{Expected log likelihood wrt } Q^{(t+1)}}$$

Use expected counts instead of counts when computing MLE:

If learning requires  $\text{Count}(\mathbf{x}, \mathbf{z})$ , Use  $E_{Q^{(t+1)}}[\text{Count}(\mathbf{x}, \mathbf{z})]$

# EM as Coordinate Ascent

$$\log P(D; \theta) \geq F(\theta, Q)$$

**E-step:** Fix  $\theta$ , maximize  $F$  over  $Q$

$$Q^{t+1} = \arg \max_Q F(\theta^t, Q)$$

$$Q^{(t+1)}(z | x_j) = P(z | x_j, \theta^{(t)})$$

E.g.,  $P(y = i | x_j, m)$

Compute probability of missing data given current choice of  $\theta$

**M-step:** Fix  $Q$ , maximize  $F$  over  $\theta$

$$\theta^{t+1} = \arg \max_{\theta} F(\theta, Q^{t+1})$$

$$\theta^{(t+1)} \leftarrow \arg \max_{\theta} \sum_z \sum_{j=1}^m \log P(z, x_j | \theta) Q^{(t+1)}(z | x_j)$$

Compute estimate of  $\theta$  by maximizing marginal likelihood using  $Q^{(t+1)}(z | x_j)$

# Summary: EM Algorithm

- A way of maximizing likelihood function for hidden variable models. Finds MLE of parameters when the original (hard) problem can be broken up into two (easy) pieces:
  1. Estimate some “missing” or “unobserved” data from observed data and current parameters.
  2. Using this “complete” data, find the maximum likelihood parameter estimates.
- Alternate between filling in the latent variables using the best guess (posterior) and updating the parameters based on this guess:
  1. E-step:  $Q^{t+1} = \arg \max_Q F(\theta^t, Q)$
  2. M-step:  $\theta^{t+1} = \arg \max_{\theta} F(\theta, Q^{t+1})$
- In the M-step we optimize a lower bound on the likelihood. In the E-step we close the gap, making bound=likelihood.
- EM performs coordinate ascent on  $F$ , can get stuck in local minima.
- BUT Extremely popular in practice.