

Lecture 5

Mixture models

$$p(\vec{y}|\vec{\theta}) = \sum_{k=1}^K \pi_k \underbrace{p_k(\vec{y})}_{\text{prob. distribution of the } k^{\text{th}} \text{ mixture component}},$$

$$\begin{cases} 0 \leq \pi_k \leq 1, \\ \sum_{k=1}^K \pi_k = 1. \end{cases}$$

prob. distribution
of the k^{th} mixture
component

One can define $p_k(\vec{y}) = p(\vec{y}|\vec{\theta}_k) =$
 $= p(\vec{y}|z=k, \vec{\theta})$, where

$z = \{1, 2, \dots, K\}$ is a latent
variable which labels mixture
components and

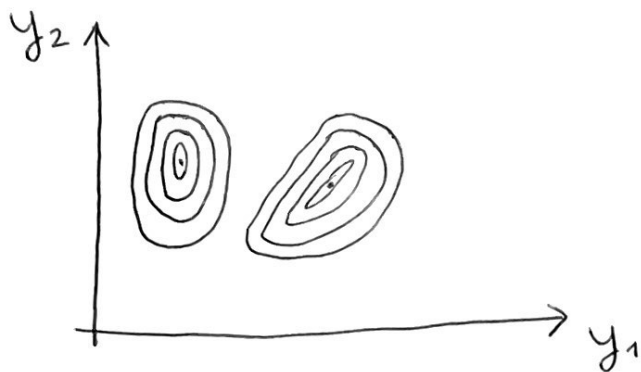
$\vec{\theta} = \{\pi_1, \dots, \pi_K, \vec{\theta}_1, \dots, \vec{\theta}_K\}$ is a
vector of all model prms.

$$\text{Now, } p(\vec{y}|\vec{\theta}) = \sum_{k=1}^K p(\vec{y}|z=k, \vec{\theta}) \underbrace{p(z=k|\vec{\theta})}_{\text{Cat}(z|\vec{\theta}) = \pi_k} \quad \textcircled{=}$$

$$\textcircled{=} \sum_k \pi_k p_k(\vec{y}), \quad \text{as above}$$

Gaussian mixture models (GMM):

$$p(\vec{y}|\vec{\theta}) = \sum_{k=1}^K \pi_k \mathcal{N}(\vec{y}|\vec{\mu}_k, \Sigma_k)$$



mixture of 2
gaussians in 2D

Bernoulli mixture models (BMM):

Let $y_i = \{0, 1\}$ and $\vec{y} = \{y_1, y_2, \dots, y_D\}$
is a set of D binary observations,

$$\left\{ \begin{aligned} p(\vec{y} | z=k, \vec{\theta}) &= \prod_{d=1}^D \text{Ber}(y_d | \mu_{d,k}) = \\ &= \prod_{d=1}^D \mu_{d,k}^{y_d} (1 - \mu_{d,k})^{1-y_d}, \\ p(z=k | \vec{\theta}) &= \text{Cat}(z=k | \vec{\theta}) = \pi_k \text{ as above.} \end{aligned} \right.$$

Here, $\vec{\theta} = \{\pi_1, \dots, \pi_K, \vec{\mu}_1, \dots, \vec{\mu}_K\}$, where

each $\vec{\mu}_j = \{\mu_{1,j}, \mu_{2,j}, \dots, \mu_{D,j}\}$
 $j=1, \dots, K$ is a D -dim vector

BMMs can be used to fit black-and-white
images, e.g. 0-9 digits in the MNIST
dataset.

dominates $\log p(\vec{\theta})$, leading to

$$\vec{\theta}_{\text{MAP}} \rightarrow \vec{\theta}^*.$$

This is equiv. to $p(\vec{\theta} | \mathcal{D}) = \delta(\vec{\theta} - \vec{\theta}_{\text{MAP}})$.

For unsupervised learning, we

consider $\mathcal{I}(\vec{\theta}) = \sum_{n=1}^N \log p(\vec{y}_n | \vec{\theta})$,

$$\vec{\theta}^* = \underset{\vec{\theta}}{\text{argmax}} \mathcal{I}(\vec{\theta}), \text{ etc.}$$

Now, consider the empirical distribution of the data: (unsupervised case)

$$p_{\mathcal{D}}(\vec{y}) = \frac{1}{N} \sum_{n=1}^N \delta(\vec{y} - \vec{y}_n)$$

Define Kullback-Leibler (KL) divergence:

$$D_{\text{KL}}(p || q) = \sum_{\vec{y}} p(\vec{y}) \log \frac{p(\vec{y})}{q(\vec{y})}$$

$$D_{\text{KL}}(p || q) \geq 0, \quad = 0 \text{ iff } p = q.$$

$$\mathcal{H} \int q(\vec{y}) \rightarrow p(\vec{y} | \vec{\theta}), \Rightarrow D_{\text{KL}}(p || q) \Leftrightarrow$$
$$\begin{cases} p(\vec{y}) \rightarrow p_{\mathcal{D}}(\vec{y}) \end{cases}$$

$$\Leftrightarrow \underbrace{\sum_{\vec{y}} p_{\mathcal{D}} \log p_{\mathcal{D}}}_{\text{-entropy of } p_{\mathcal{D}}, \text{ const}(\vec{\theta})} - \underbrace{\sum_{\vec{y}} p_{\mathcal{D}} \log p(\vec{y} | \vec{\theta})}_{\text{-cross-entropy between } p_{\mathcal{D}} \text{ \& } p(\vec{y} | \vec{\theta})} \diamond$$

$$\diamond \text{const}(\vec{\theta}) - \underbrace{\frac{1}{N} \sum_{n=1}^N \log p(\vec{y}_n | \vec{\theta})}_{\mathcal{I}(\vec{\theta})}.$$

Thus, maximizing $\mathcal{I}(\vec{\theta})$ is equiv. to minimizing the KL divergence.

—

Likewise, consider

$$p_D(\vec{x}, \vec{y}) = p_D(\vec{y} | \vec{x}) p_D(\vec{x}) = \frac{1}{N} \sum_{n=1}^N \delta(\vec{x} - \vec{x}_n) \delta(\vec{y} - \vec{y}_n)$$

for supervised learning

$$\begin{aligned} \text{Then } E_{p_D(\vec{x})} [D_{KL}(p_D(\vec{y} | \vec{x}) \| p(\vec{y} | \vec{x}, \vec{\theta}))] &= \\ &= \sum_{\vec{x}} p_D(\vec{x}) \left[\sum_{\vec{y}} p_D(\vec{y} | \vec{x}) \log \frac{p_D(\vec{y} | \vec{x})}{p(\vec{y} | \vec{x}, \vec{\theta})} \right] = \\ &= \text{const}(\vec{\theta}) - \sum_{\vec{x}, \vec{y}} p_D(\vec{x}, \vec{y}) \log p(\vec{y} | \vec{x}, \vec{\theta}) = \\ &= \text{const}(\vec{\theta}) - \underbrace{\frac{1}{N} \sum_{n=1}^N \log p(\vec{y}_n | \vec{x}_n, \vec{\theta})}_{\mathcal{I}(\vec{\theta})}. \end{aligned}$$

Thus, maximizing $\mathcal{I}(\vec{\theta})$ is equiv. to minimizing the average KL divergence.

MLE, Bernoulli:

$$\mathcal{I}(\theta) = \sum_{n=1}^N \log p(y_n | \theta), \text{ where}$$

$$y_n \in \{0, 1\}, \quad \theta = p(y=1).$$

$n=1, \dots, N$

$$\begin{aligned}
 \text{Then } \mathcal{I}(\vec{\theta}) &= \sum_{n=1}^N \log [\theta^{\mathbb{I}(y_n=1)} (1-\theta)^{\mathbb{I}(y_n=0)}] = \\
 &= \underbrace{\left(\sum_{n=1}^N \mathbb{I}(y_n=1) \right)}_{N_1} \log \theta + \underbrace{\left(\sum_{n=1}^N \mathbb{I}(y_n=0) \right)}_{N_0} \log(1-\theta) .
 \end{aligned}$$

$N_0 + N_1 = N$

$$\text{Finally, } \left. \frac{d}{d\theta} \mathcal{I}(\vec{\theta}) \right|_{\theta^*} = \frac{N_1}{\theta^*} - \frac{N_0}{1-\theta^*} = 0 \text{ yields}$$

$$N_1(1-\theta^*) - N_0\theta^* = 0,$$

$$\theta^* = \frac{N_1}{N_0 + N_1}, \text{ fraction of } y_n=1 \text{ observations}$$

MLE, categorical:

$$\begin{aligned}
 \mathcal{I}(\vec{\theta}) &= \sum_{k=1}^K N_k \log \theta_k, \\
 \sum_{k=1}^K N_k &= N; \quad \sum_{k=1}^K \theta_k = 1.
 \end{aligned}$$

$$\text{Consider } \mathcal{I}(\vec{\theta}, \lambda) = \sum_k N_k \log \theta_k + \lambda (1 - \sum_k \theta_k).$$

$$\text{Then } \left. \frac{\partial \mathcal{I}}{\partial \theta_i} \right|_{\theta_i^*} = \frac{N_i}{\theta_i^*} - \lambda = 0 \Rightarrow N_i = \lambda \theta_i^*.$$

Find λ by normalization:

$$\sum_i N_i = N = \lambda \sum_i \theta_i^* \Rightarrow \lambda = N.$$

Thus, $\theta_i^* = \frac{N_i}{N}$, $i=1, \dots, K$.
 fraction of event i occurring

MLE, 1D gaussian:

$$Y \sim N(\mu, \sigma^2)$$

$$\mathcal{D} = \{y_1, \dots, y_N\}$$

$$\begin{aligned} \mathcal{L}(\mu, \sigma^2) &= \sum_{n=1}^N \log \left[(2\pi\sigma^2)^{-1/2} e^{-\frac{(y_n - \mu)^2}{2\sigma^2}} \right] = \\ &= N \left(-\frac{1}{2}\right) \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum_n (y_n - \mu)^2. \end{aligned}$$

$$\text{Consider } \begin{cases} \frac{\partial}{\partial \mu} \mathcal{L}(\mu, \sigma^2) = 0, \\ \frac{\partial}{\partial \sigma^2} \mathcal{L}(\mu, \sigma^2) = 0. \end{cases}$$

$$\text{we have } \left. \frac{\partial}{\partial \mu} \sum_n (y_n - \mu)^2 \right|_{\mu^*} = 2(N\mu^* - \sum_n y_n) = 0, \quad \text{or}$$

$$\left[\mu^* = \frac{1}{N} \sum_n y_n \equiv \bar{y} \right] \text{ sample average}$$

$$\frac{N}{2} \frac{1}{2\pi\sigma_*^2} 2\pi - \frac{1}{2(\sigma_*^2)^{3/2}} \sum_n (y_n - \bar{y})^2 = 0, \quad \text{"}\mu^*\text{" or}$$

$$\frac{1}{\sigma_*^2} \sum_n (y_n - \bar{y})^2 = N, \quad \text{sample variance}$$

$$\left[\sigma_*^2 = \frac{1}{N} \sum_n (y_n - \bar{y})^2 \right]$$

MLE, multi-D Gaussian: