

Mixtures - E.G. 7 and 7

$$p(\underline{x}) = \sum_{t=1}^K p(\underline{x} | G_t) P(G_t)$$

G_t are the mixture components

$p(\underline{x} | G_t)$ component densities.

$P(G_t)$ mixture proportion. Different classes

If component densities are Gaussian we have $p(\underline{x} | G_t) \sim \mathcal{N}(\underline{\mu}_t, \underline{\Sigma}_t)$

$\underline{\theta} = \{ P(G_t), \underline{\mu}_t, \underline{\Sigma}_t \}_{t=1}^K$, one parameter that need to be estimated.

In supervised case, we know the classes, and the assignment of instances to classes.

$$r_i^t = \begin{cases} 1, & \text{if } \underline{x}^t \in C_i \\ 0, & \text{otherwise} \end{cases}$$

In this case, it is straight forward to learn the parameters

$$\hat{P}(G_t) = \sum_i r_i^t / N$$

$$\underline{m}_t = \sum_i r_i^t \underline{x}^t / \sum_i r_i^t$$

$$\underline{\Sigma}_t = \sum_i r_i^t (\underline{x}^t - \underline{m}_t)(\underline{x}^t - \underline{m}_t)^T / \sum_i r_i^t$$

(2)

Mixture Densities (cont)

But, usually we do not know the classes.

This is called unsupervised learning

i.e. r^t is not known.

We need to estimate r^t — later, we show that the EM algorithm can be used.

k-means clustering.

Vector quantization:

Sample $X = \{x^t\}_{t=1}^N$

k reference vectors $\underline{m}_j \quad j=1 \text{ to } k.$

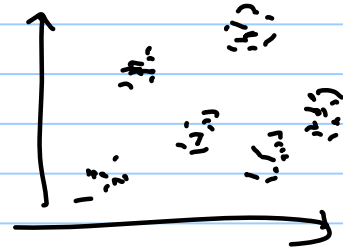
Represent x^t by the most similar entry m_i

$$\|x^t - \underline{m}_i\| = \min_j \|x^t - \underline{m}_j\|$$

$\underline{m}_i$ — codebook vectors.

Compression — results from quantization.

How to calculate the best $\{\underline{m}_j\}$?



(3)

K-means (cont)

Reconstruction Error

$$E(\{\underline{m}_i\}_{i=1}^k | X) = \sum_t \sum_i b_i^t \|\underline{x}^t - \underline{m}_i\|^2.$$

$$\text{where } b_i^t = \begin{cases} 1, & \text{if } \|\underline{x}^t - \underline{m}_i\| = \min_j \|\underline{x}^t - \underline{m}_j\| \\ 0, & \text{otherwise} \end{cases}$$

k-means is an iterative procedure that seeks to minimize $E(\cdot)$.

Start with random $\underline{m}_i$.

Then compute b_i^t

$$\text{re-estimate the means } \underline{m}_i = \frac{\sum_t b_i^t \underline{x}^t}{\sum_t b_i^t}.$$

Repeat.

This procedure is guaranteed to converge (almost always).

But the final results depend on the initialization.

Alternative Initialization Strategies:

- Take random k instances as initial $\underline{m}_i$
- Calculate the means of all the data, then add small perturbation
- Calculate the principal components, divide its range into k-intervals, etc.

(4)

Expectation-Maximization (EM) Algorithm.

$$\begin{aligned} \mathcal{L}(\phi | X) &= \log \prod_t p(x^t | \phi) \\ &= \sum_t \log \sum_{i=1}^k p(x^t | \sigma_i) P(\sigma_i) \end{aligned}$$

EM Introduce hidden variable Z which associate the data to the class.

$$Z^t = (z_1^t, \dots, z_k^t), \quad z_i^t = \begin{cases} 1 & \text{if } x^t \text{ belongs to class } \sigma_i \\ 0 & \text{otherwise.} \end{cases}$$

$$P(Z^t) = \prod_{i=1}^k \pi_i^{z_i^t}$$

$$P(x^t | Z^t) = \prod_{i=1}^k p_i(x^t)^{z_i^t}$$

$p_i(x^t)$ short for $p(x^t | \sigma_i)$

$$P(x^t, Z^t) = P(Z^t) P(x^t | Z^t)$$

$$\mathcal{L}_c(\phi | X, Z) = \log \prod_t p(x^t, z^t | \phi)$$

$$= \sum_t \log p(x^t, z^t | \phi)$$

$$= \sum_t (\log p(z^t | \phi) + \log p(x^t | z^t, \phi))$$

$$= \sum_t \sum_i z_i^t (\log \pi_i + \log p_i(x^t | \phi))$$

(5) EM with

E-step

$$\begin{aligned} Q(\phi | \phi^e) &= E[\log p(x, z) | x, \phi^e] \\ &= E[\sum_i \log p_i(x, z) | x, \phi^e] \\ &= \sum_i E[\log p_i(x, z) | x, \phi^e] \\ &\quad \{ \log \pi_i + \log p_i(x^t | \phi^e) \} \end{aligned}$$

where $E[\log p_i(x, z) | x, \phi^e]$

$$\begin{aligned} &= E[\log p_i(x^t, z^t) | x^t, \phi^e] \\ &= \log p(z^t = 1 | x^t, \phi^e) \\ &= \log p(x^t | z^t = 1, \phi^e) p(z^t = 1 | \phi^e) \\ &\quad \frac{p(x^t | \phi^e)}{p(x^t | \phi^e)} \quad \text{Bayes rule} \\ &= \frac{p_i(x^t | \phi^e) \pi_i}{\sum_j p_j(x^t | \phi^e) \pi_j} \\ &= \frac{p(x^t | G_i, \phi^e) p(G_i)}{\sum_j p(x^t | G_j, \phi^e) p(G_j)} \\ &= p(G_i | x^t, \phi^e) = h_i^t \end{aligned}$$

Expected value of hidden variable, $E[z^t]$
is the posterior probability that x^t is generated by G_i .
"Soft label"

(6) EM (Cont)

M-step: maximize to get next of parameter $\phi^{l+1} = \arg \max_{\phi} Q(\phi | \phi^l)$

$$Q(\phi | \phi^l) = \sum_t \sum_i h_i^t \{ \log \pi_i + \log p_i(x^t | \phi^l) \}$$
$$= \sum_t \sum_i h_i^t \log \pi_i + \sum_t \sum_i h_i^t \log p_i(x^t | \phi^l)$$

solve to get $\pi_i = \sum_t h_i^t / N$

If we assume Gaussian components
 $p_j(x^t | \phi) \sim N(\underline{m}_j, \underline{S}_j)$

M-step is $\underline{m}_i^{l+1} = \frac{\sum_t h_i^t x^t}{\sum_t h_i^t}$

$$\underline{S}_i^{l+1} = \frac{\sum_t h_i^t (x^t - \underline{m}_i^{l+1})(x^t - \underline{m}_i^{l+1})^T}{\sum_t h_i^t}$$

where $h_i^t = \frac{\pi_i |\underline{S}_i|^{-k/2} \exp\{-\frac{1}{2}(x^t - \underline{m}_i)^T \underline{S}_i^{-1} (x^t - \underline{m}_i)\}}{\sum_j \pi_j |\underline{S}_j|^{-k/2} \exp\{-\frac{1}{2}(x^t - \underline{m}_j)^T \underline{S}_j^{-1} (x^t - \underline{m}_j)\}}$

(7)

Supervised Learning after Clustering.

Clustering can be used for data exploration by grouping instances.

If such groups are found, they can be named and their attributes defined.

Clustering is frequently used as a pre-processing stage.

(8) Modern View of EM.

Task estimate parameters ϕ from $p(x, z | \phi)$.

Problem we can't observe z

Proposed Solution: Estimate ϕ from $p(x | \phi)$
where $p(x | \phi) = \sum_z p(x, z | \phi)$

Use iterative algorithm, EM, to estimate ϕ .

Formulate minimize $-\log p(x | \phi)$
w.r.t. ϕ .

Equivalent to minimize

$$F[\phi, q(z)] = -\log p(x | \phi) + \sum_z q(z) \log \frac{q(z)}{p(z | x, \phi)}$$

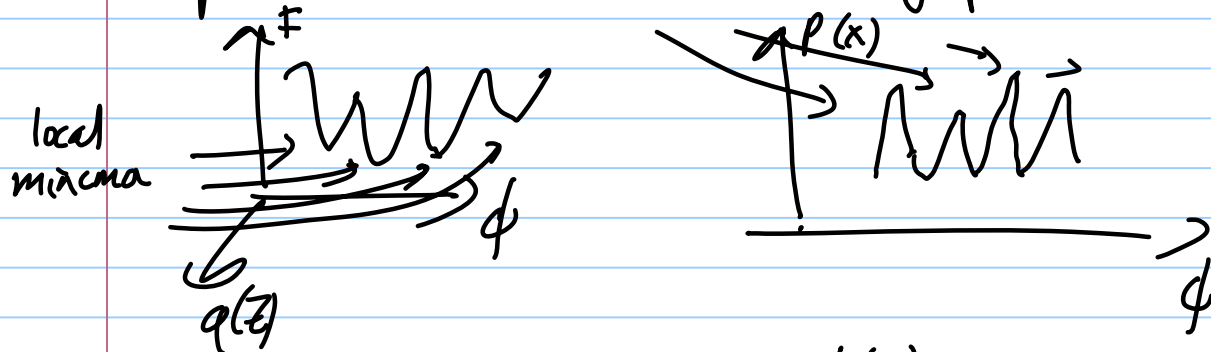
w.r.t. ϕ & $q(z)$.

Why $\rightarrow$ because the second term is ≥ 0 and $= 0$ only if $q(z) = p(z | x, \phi)$

Kullback-Leibler divergence.

(9) Algorithm (EM) minimize $F[\phi, q(z)]$ w.r.t. ϕ & $q(z)$ alternately.

This is guaranteed to always decrease F and converge to a local minimum. This corresponds to a local maximum of $p(x|\phi)$



Initialize the parameters $\phi^{(0)}$,
then repeat the two steps : at time t .

E-step : minimize F w.r.t. $q(z)$ gives
 $q^{t+1}(z) = p(z|x, \phi^t)$

M-step : minimize F w.r.t. ϕ gives
 $\phi^{t+1} = \underset{\phi}{\text{arg min}} - \sum_z q^{t+1}(z) \log p(x, z|\phi)$

see next page ...

(10) E-step follows from standard expression of $F[\phi, q(z)] = -\log p(x|\phi) + \sum_z q(z) \log \left\{ \frac{q(z)}{p(z|x, \phi)} \right\}$

M-step follows by re-expressing $q(z)$ ^{minimize by setting} $q(z) = p(z|x, \phi)$.

$$F = -\log p(x|\phi) + \sum_z q(z) \log q(z) - \sum_z q(z) \log p(z|x, \phi)$$

$$F = \sum_z q(z) \log q(z) - \sum_z q(z) \log p(x|\phi) - \sum_z q(z) \log p(z|x, \phi)$$

using $\sum_z q(z) = 1$

$$= -\sum_z q(z) \log \{p(x|\phi) p(z|x, \phi)\} - \sum_z q(z) \log p(x, z|\phi)$$

This is the "modern" formulation of EM. It gives a proof of convergence. It shows where the two steps come from.

It also leads to variants of EM (beyond scope of the course).