



Lecture Slides for

INTRODUCTION TO

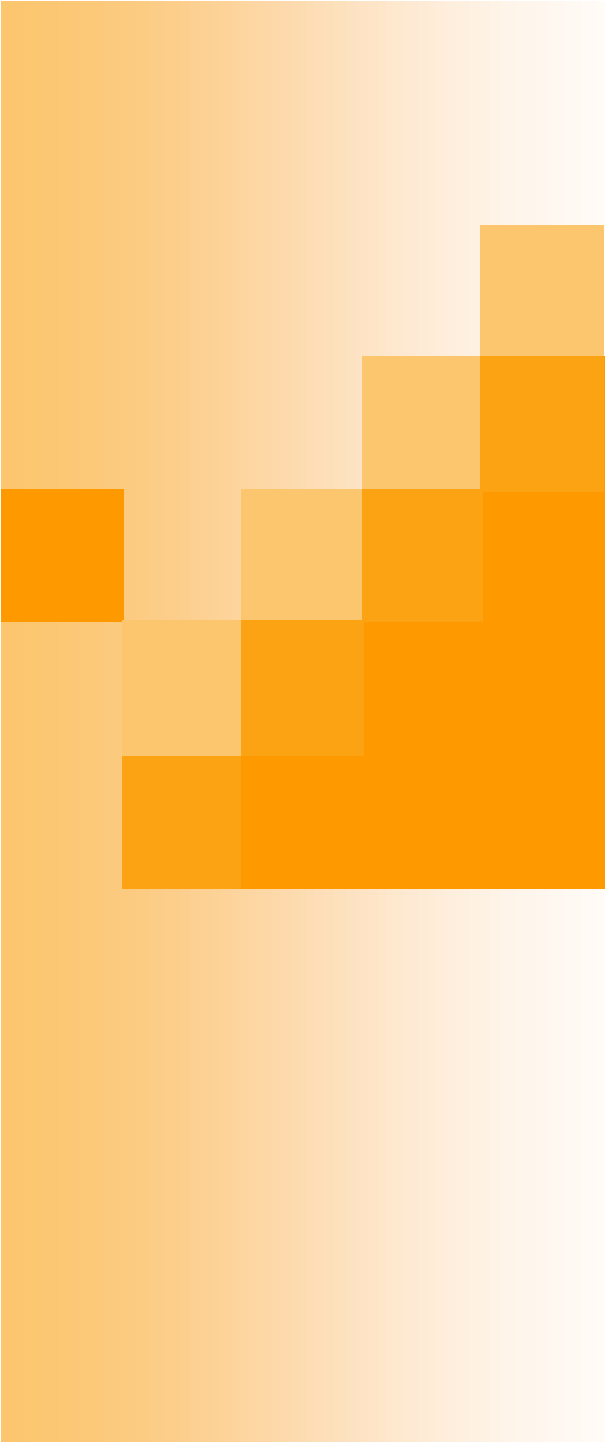
Machine Learning

ETHEM ALPAYDIN

© The MIT Press, 2004

alpaydin@boun.edu.tr

<http://www.cmpe.boun.edu.tr/~ethem/i2ml>



CHAPTER 7:

Clustering



Semiparametric Density Estimation

- **Parametric:** Assume a single model for $p(\mathbf{x} | C_i)$ (Chapter 4 and 5)
- **Semiparametric:** $p(\mathbf{x} | C_i)$ is a **mixture** of densities
Multiple possible explanations/prototypes:
Different handwriting styles, accents in speech
- **Nonparametric:** No model; data speaks for itself (Chapter 8)



Mixture Densities

$$p(\mathbf{x}) = \sum_{i=1}^k p(\mathbf{x} | G_i) P(G_i)$$

where G_i the components/groups/clusters,
 $P(G_i)$ mixture proportions (priors),
 $p(\mathbf{x} | G_i)$ component densities

Gaussian mixture where $p(\mathbf{x} | G_i) \sim \mathcal{N}(\boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)$
parameters $\Phi = \{P(G_i), \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i\}_{i=1}^k$
unlabeled sample $\mathcal{X} = \{\mathbf{x}^t\}_t$ (unsupervised learning)

Classes vs. Clusters

- Supervised $\mathcal{X} = \{\mathbf{x}^t, \mathbf{r}^t\}_{t=1}^N$
- Classes $C_i, i=1, \dots, K$
- Unsupervised $\mathcal{X} = \{\mathbf{x}^t\}_t$
- Clusters $G_i, i=1, \dots, k$

$$p(\mathbf{x}) = \sum_{i=1}^K p(\mathbf{x}|C_i)P(C_i)$$

where $p(\mathbf{x}|C_i) \sim \mathcal{N}(\boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)$

- $\Phi = \{P(C_i), \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i\}_{i=1}^K$

$$\hat{P}(C_i) = \frac{\sum_t r_i^t}{N}$$

$$\mathbf{m}_i = \frac{\sum_t r_i^t \mathbf{x}^t}{\sum_t r_i^t}$$

$$\mathbf{S}_i = \frac{\sum_t r_i^t (\mathbf{x}^t - \mathbf{m}_i)(\mathbf{x}^t - \mathbf{m}_i)^T}{\sum_t r_i^t}$$

$$p(\mathbf{x}) = \sum_{i=1}^k p(\mathbf{x}|G_i)P(G_i)$$

where $p(\mathbf{x}|G_i) \sim \mathcal{N}(\boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)$

- $\Phi = \{P(G_i), \boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i\}_{i=1}^k$

Labels, $\mathbf{r}^t_i$?



k-Means Clustering

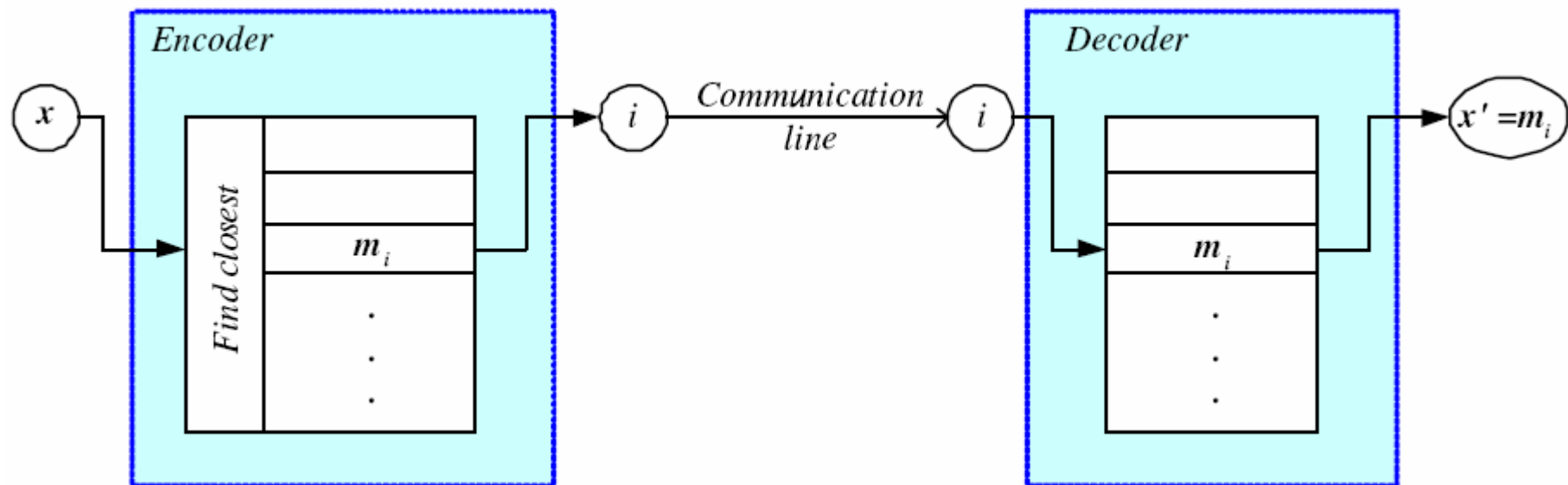
- Find k reference vectors (prototypes/codebook vectors/codewords) which best represent data
- Reference vectors, $\mathbf{m}_j, j=1,\dots,k$
- Use nearest (most similar) reference:

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

- Reconstruction error

$$E(\{\mathbf{m}_i\}_{i=1}^k | \mathcal{X}) = \sum_t \sum_i b_i^t \|\mathbf{x}^t - \mathbf{m}_i\|^2$$
$$b_i^t = \begin{cases} 1 & \text{if } \|\mathbf{x}^t - \mathbf{m}_i\| = \min_j \|\mathbf{x}^t - \mathbf{m}_j\| \\ 0 & \text{otherwise} \end{cases}$$

Encoding/Decoding



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

k-means Clustering

Initialize $\mathbf{m}_i, i = 1, \dots, k$, for example, to k random $\mathbf{x}^t$

Repeat

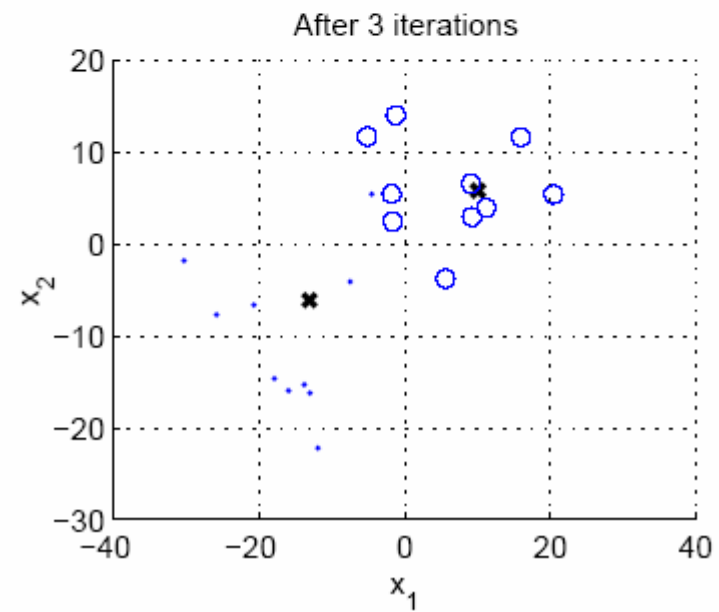
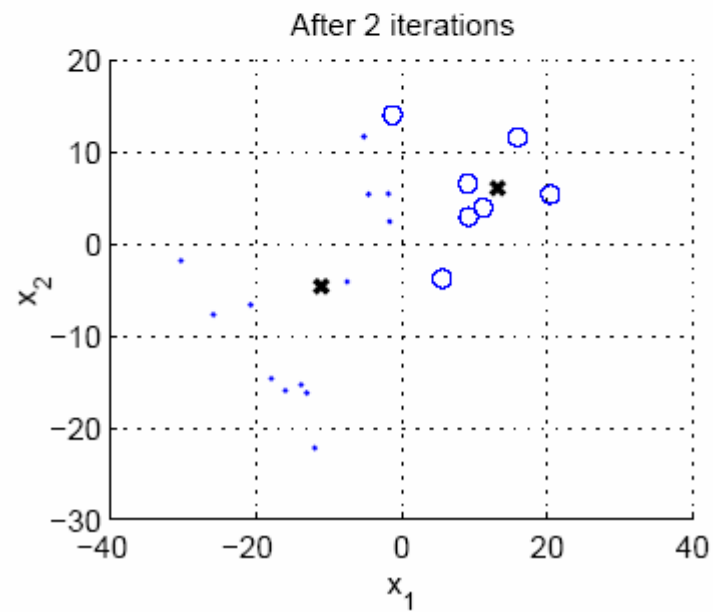
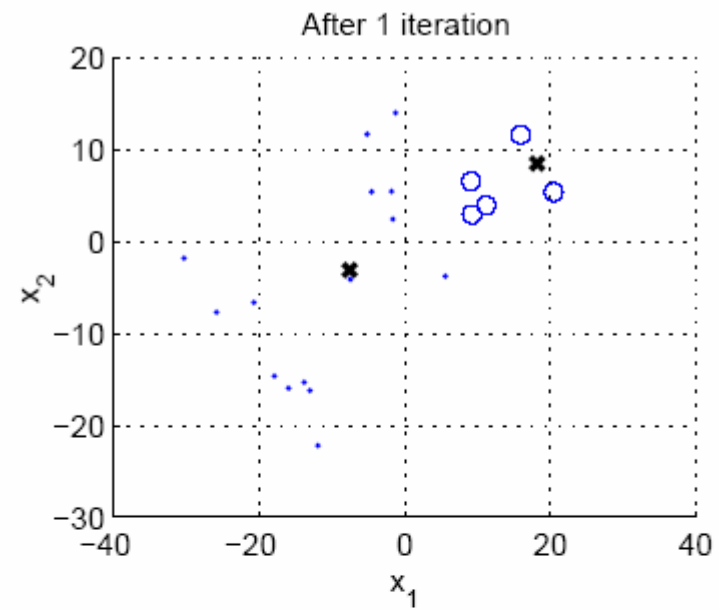
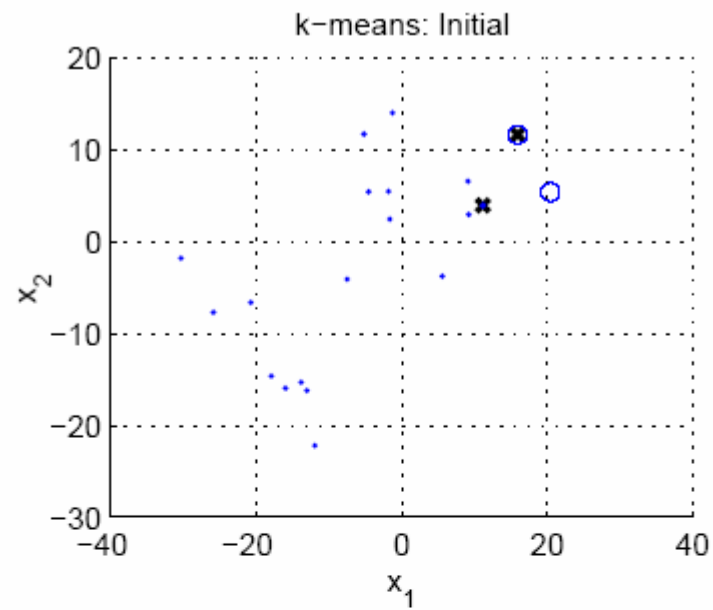
For all $\mathbf{x}^t \in \mathcal{X}$

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

For all $\mathbf{m}_i, i = 1, \dots, k$

$$\mathbf{m}_i \leftarrow \sum_t b_i^t \mathbf{x}^t / \sum_t b_i^t$$

Until $\mathbf{m}_i$ converge





Expectation-Maximization (EM)

- Log likelihood with a mixture model

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

- Assume hidden variables $\mathbf{z}$, which when known, make optimization much simpler
- Complete likelihood, $\mathcal{L}_c(\Phi|\mathcal{X},\mathcal{Z})$, in terms of $\mathbf{x}$ and $\mathbf{z}$
- Incomplete likelihood, $\mathcal{L}(\Phi|\mathcal{X})$, in terms of $\mathbf{x}$



E- and M-steps

- Iterate the two steps
 1. E-step: Estimate z given $\mathcal{X}$ and current Φ
 2. M-step: Find new Φ' given z , $\mathcal{X}$, and old Φ .

$$\text{E-step} \quad : \quad \mathcal{Q}(\Phi|\Phi^l) = E[\mathcal{L}_c(\Phi|\mathcal{X}, \mathcal{Z})|\mathcal{X}, \Phi^l]$$

$$\text{M-step} \quad : \quad \Phi^{l+1} = \arg \max_{\Phi} \mathcal{Q}(\Phi|\Phi^l)$$

An increase in $\mathcal{Q}$ increases incomplete likelihood

$$\mathcal{L}(\Phi^{l+1}|\mathcal{X}) \geq \mathcal{L}(\Phi^l|\mathcal{X})$$

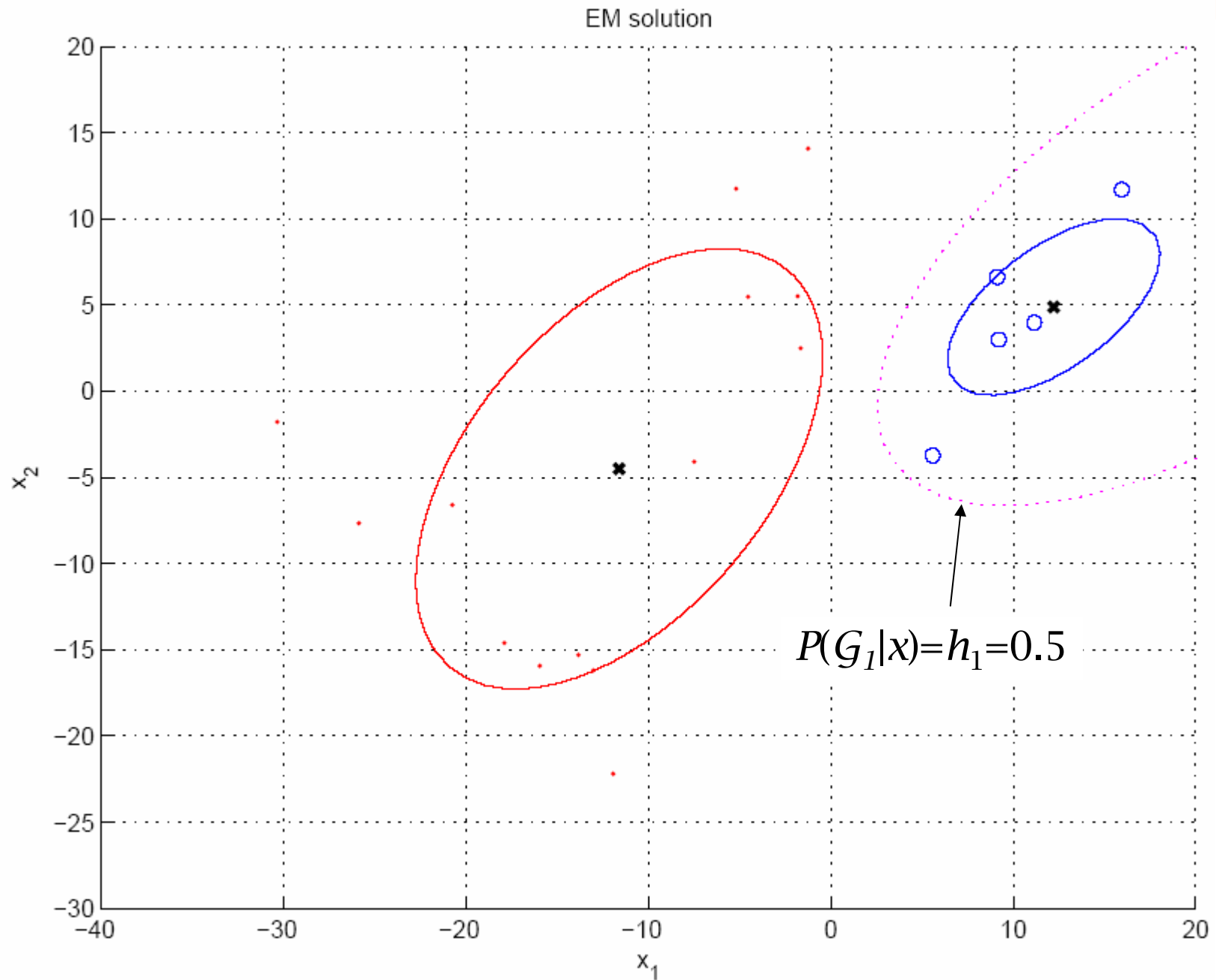
EM in Gaussian Mixtures

- $z_i^t = 1$ if $\mathbf{x}^t$ belongs to G_i , 0 otherwise (labels $\mathbf{r}^t_i$ of supervised learning); assume $p(\mathbf{x}|G_i) \sim \mathcal{N}(\boldsymbol{\mu}_i, \boldsymbol{\Sigma}_i)$

- E-step:
$$E[z_i^t | \mathcal{X}, \Phi^l] = \frac{p(\mathbf{x}^t | G_i, \Phi^l) P(G_i)}{\sum_j p(\mathbf{x}^t | G_j, \Phi^l) P(G_j)}$$
$$= P(G_i | \mathbf{x}^t, \Phi^l) \equiv h_i^t$$

- M-step:
$$P(G_i) = \frac{\sum_t h_i^t}{N} \quad \mathbf{m}_i^{l+1} = \frac{\sum_t h_i^t \mathbf{x}^t}{\sum_t h_i^t}$$
$$\mathbf{S}_i^{l+1} = \frac{\sum_t h_i^t (\mathbf{x}^t - \mathbf{m}_i^{l+1})(\mathbf{x}^t - \mathbf{m}_i^{l+1})^T}{\sum_t h_i^t}$$

Use estimated labels in place of unknown labels





Mixtures of Latent Variable Models

- Regularize clusters
 1. Assume shared/diagonal covariance matrices
 2. Use PCA/FA to decrease dimensionality: Mixtures of PCA/FA

$$p(\mathbf{x}^t | G_i) \sim \mathcal{N}(\mathbf{m}_i, \mathbf{V}_i \mathbf{V}_i^T + \Psi_i)$$

Can use EM to learn $\mathbf{V}_i$ (Ghahramani and Hinton, 1997; Tipping and Bishop, 1999)



After Clustering

- Dimensionality reduction methods find correlations between features and group features
- Clustering methods find similarities between instances and group instances
- Allows knowledge extraction through
 - number of clusters,
 - prior probabilities,
 - cluster parameters, i.e., center, range of features.Example: CRM, customer segmentation



Clustering as Preprocessing

- Estimated group labels h_j (soft) or b_j (hard) may be seen as the dimensions of a new k dimensional space, where we can then learn our discriminant or regressor.
- **Local representation** (only one b_j is 1, all others are 0; only few h_j are nonzero) vs **Distributed representation** (After PCA; all z_j are nonzero)



Mixture of Mixtures

- In classification, the input comes from a mixture of classes (supervised).
- If each class is also a mixture, e.g., of Gaussians, (unsupervised), we have a mixture of mixtures:

$$p(\mathbf{x}|C_i) = \sum_{j=1}^{k_i} p(\mathbf{x}|G_{ij})P(G_{ij})$$

$$p(\mathbf{x}) = \sum_{i=1}^K p(\mathbf{x}|C_i)P(C_i)$$



Hierarchical Clustering

- Cluster based on similarities/distances
- Distance measure between instances $\mathbf{x}^r$ and $\mathbf{x}^s$
Minkowski (L_p) (Euclidean for $p = 2$)

$$d_m(\mathbf{x}^r, \mathbf{x}^s) = \left[\sum_{j=1}^d (x_j^r - x_j^s)^p \right]^{1/p}$$

City-block distance

$$d_{cb}(\mathbf{x}^r, \mathbf{x}^s) = \sum_{j=1}^d |x_j^r - x_j^s|$$



Agglomerative Clustering

- Start with N groups each with one instance and merge two closest groups at each iteration
- Distance between two groups G_i and G_j :

- Single-link:

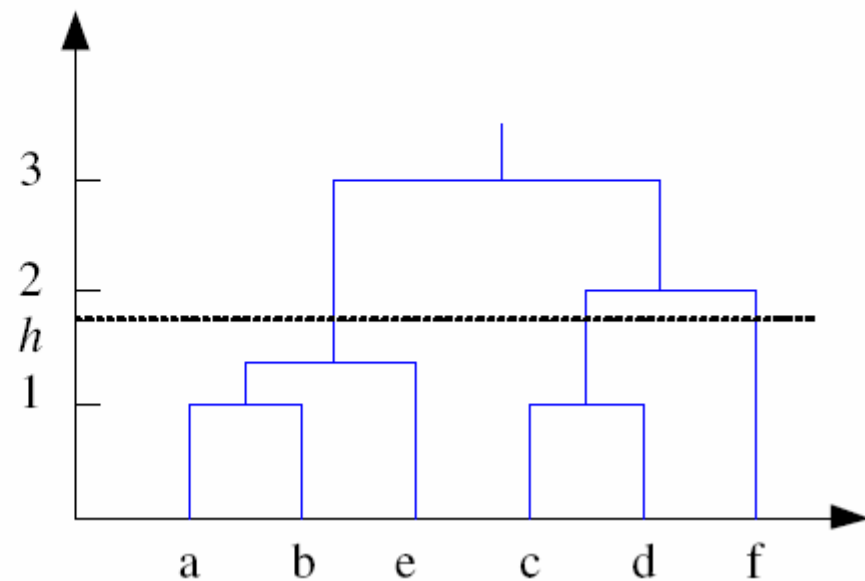
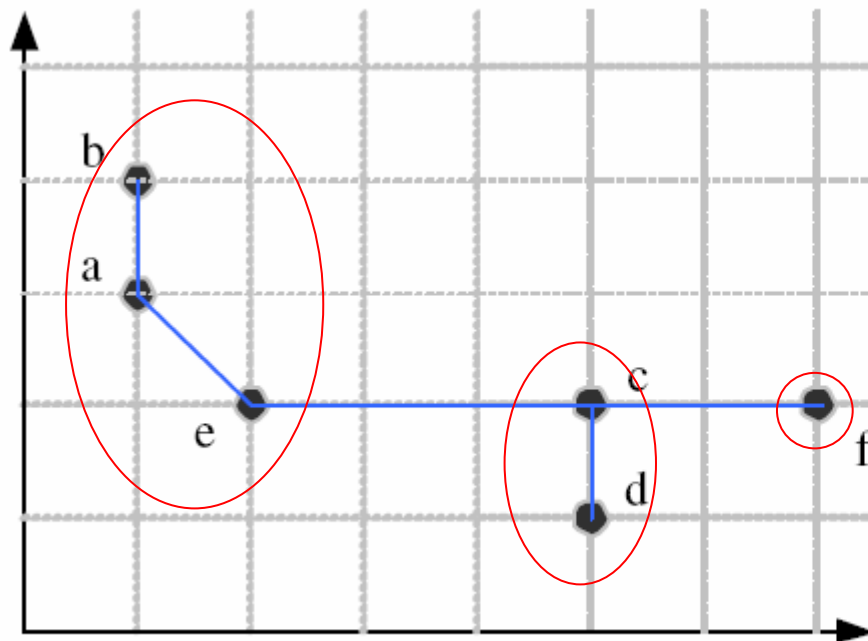
$$d(G_i, G_j) = \min_{\mathbf{x}^r \in G_i, \mathbf{x}^s \in G_j} d(\mathbf{x}^r, \mathbf{x}^s)$$

- Complete-link:

$$d(G_i, G_j) = \max_{\mathbf{x}^r \in G_i, \mathbf{x}^s \in G_j} d(\mathbf{x}^r, \mathbf{x}^s)$$

- Average-link, centroid

Example: Single-Link Clustering



Dendrogram



Choosing k

- Defined by the application, e.g., image quantization
- Plot data (after PCA) and check for clusters
- Incremental (leader-cluster) algorithm: Add one at a time until “elbow” (reconstruction error/log likelihood/intergroup distances)
- Manual check for meaning