

Mixtures of Gaussians

- Uses:
- Clustering
 - Model non-Gaussian distributions



Model:

Same as Gaussian Bayes classifier

Except "labels" are hidden or "latent" variables

Cluster choice $z^{(n)} \sim \text{Discrete}(\underline{\pi})$

$$z^{(n)} \in \{1, 2, 3 \dots K\}$$

↑
+ve vector,
sums to 1

Features

$$\underline{x}^{(n)} | z^{(n)} = k \sim N(\underline{x}^{(n)}; \underline{\mu}^{(k)}, \Sigma^{(k)})$$

Likelihood

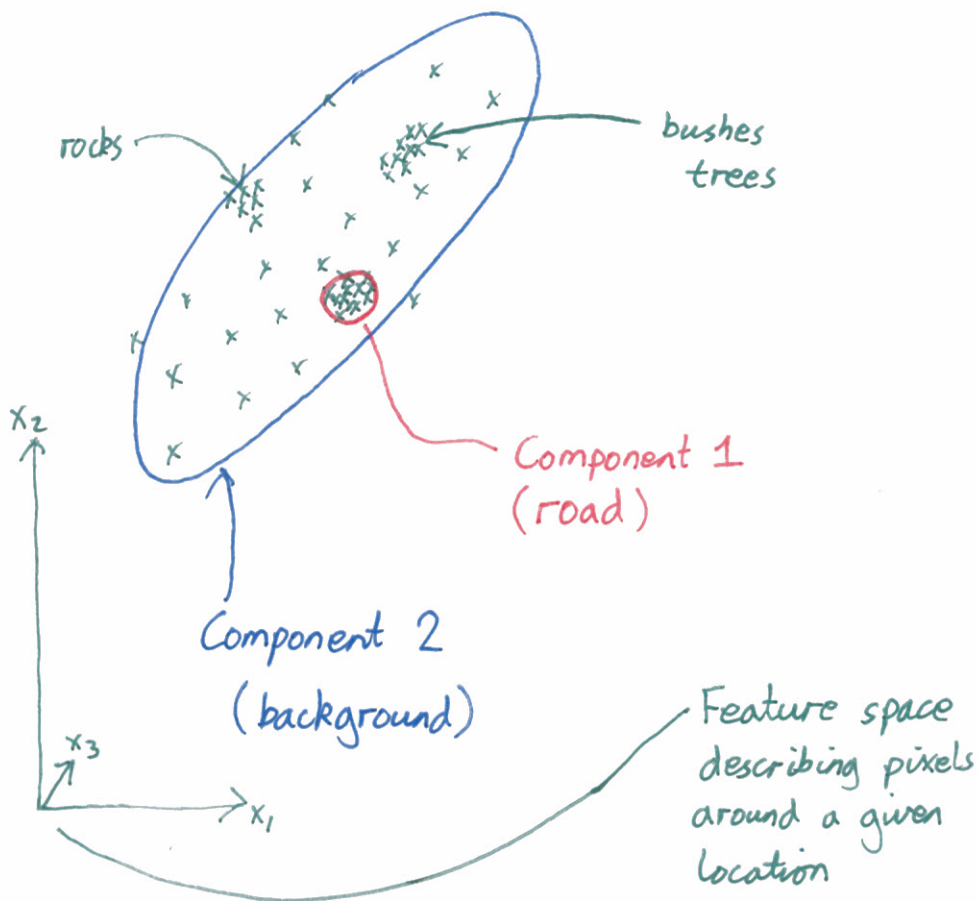
Given one observation: $p(\underline{x}^{(n)} | \theta = (\underline{\pi}, \{\underline{\mu}^{(k)}, \Sigma^{(k)}\}))$

$$= \sum_{k=1}^K p(\underline{x}^{(n)}, z^{(n)} = k | \theta)$$

$$= \sum_k p(\underline{x}^{(n)} | z^{(n)} = k, \theta) p(z^{(n)} = k | \theta)$$

$$= \sum_k N(\underline{x}^{(n)}; \underline{\mu}^{(k)}, \Sigma^{(k)}) \pi_k$$

Use in "Stanley" car



Gradient-based fitting

$$\text{Cost} = - \sum_{n=1}^N \log(p(\underline{x}^{(n)} | \theta))$$

Write as function of unconstrained parameters

$\tilde{L}^{(k)}$ lower-triangular matrix



$$L^{(k)} = \begin{cases} L_{ij} = e^{\tilde{L}_{ij}} & i=j \\ L_{ij} = \tilde{L}_{ij} & i > j \\ L_{ij} = 0 & i < j \end{cases}$$



$$\Sigma^{(k)} = L L^{(k)\top}$$



neg. log likelihood

$k=1, \dots, K$

$\underline{c}$ vector k params

$$\underline{\Pi} = \text{softmax}(\underline{c})$$

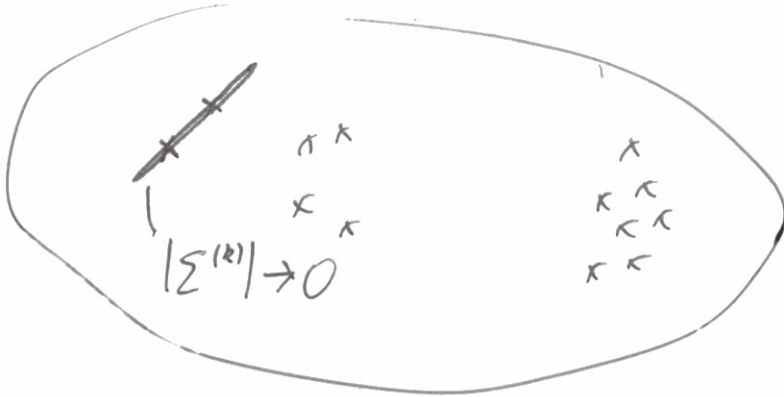
Initialization matters

$$\underline{M}^{(k)} = \underline{0}$$

$$\underline{\Sigma}^{(k)} = \underline{I}$$

$$\underline{M}^{(k)} \leftarrow \underline{M}^{(k)} - \eta \nabla_{\underline{w}} NLL$$

Maximum Likelihood is ∞



The alternative EM

Idea: Pretend we observe $\{z^{(n)}\}$

Responsibility initialization:

$$\Gamma_R^{(n)} = \begin{cases} 1 & \text{if } z^{(n)} = k \\ 0 & \text{otherwise} \end{cases}$$

Maximum likelihood parameters:

$$\pi_R = \frac{\Gamma_R}{N} \quad \Gamma_k = \sum_{n=1}^N \Gamma_k^{(n)}$$

$$\underline{\mu}^{(k)} = \frac{1}{\Gamma_k} \sum_{n=1}^N \Gamma_k^{(n)} \underline{x}^{(n)}$$

$$\underline{\Sigma}^{(k)} = \frac{1}{\Gamma_k} \sum_n \Gamma_k^{(n)} \underline{x}^{(n)} \underline{x}^{(n)T} - \underline{\mu}^{(k)} \underline{\mu}^{(k)T}$$

EM Algorithm

0) Initialize params θ

1) E-step

$$r_k^{(n)} = P(z^{(n)} = k \mid \underline{x}^{(n)}, \theta)$$

2) M-step

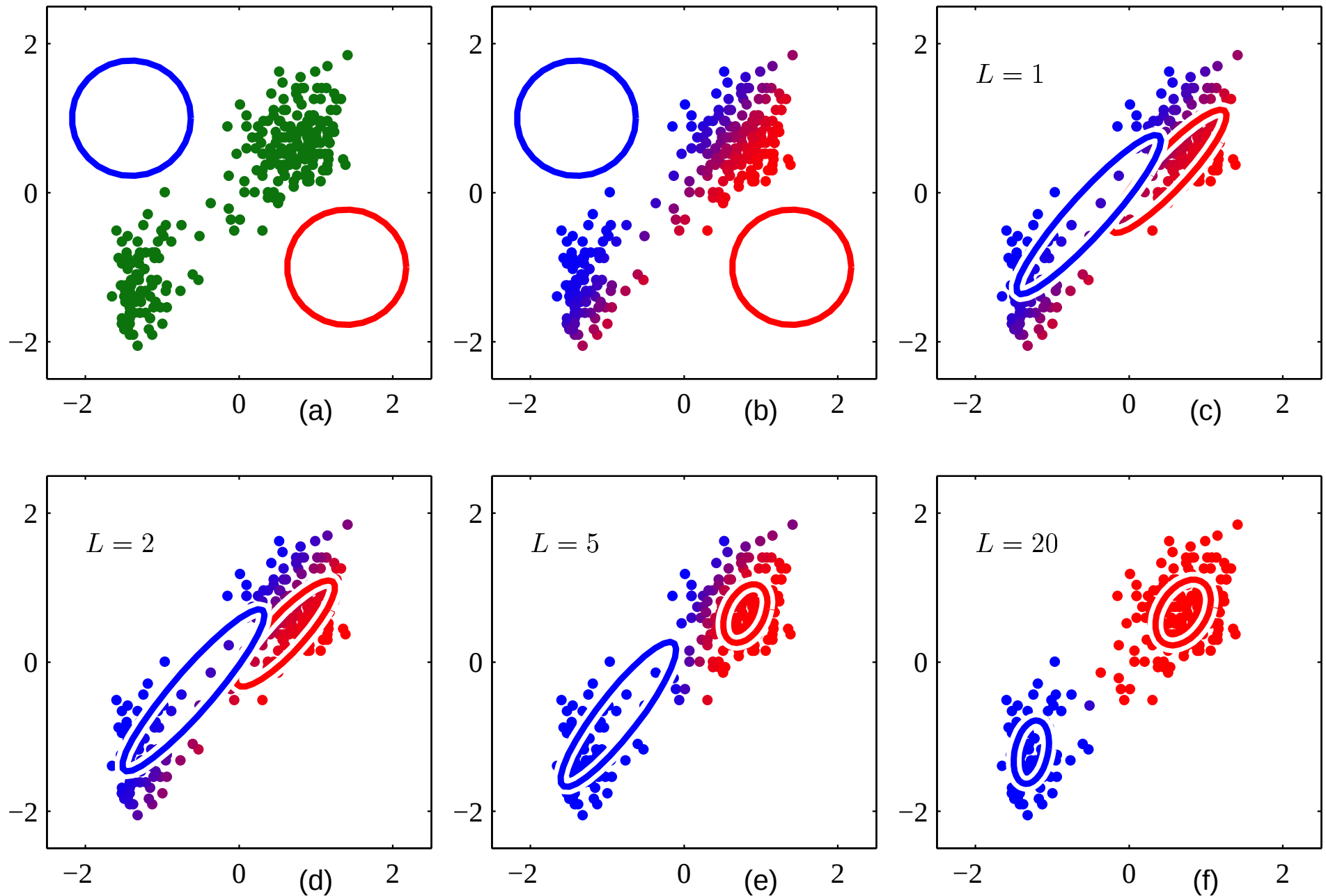
Uses these real-valued $r_k^{(n)}$

In these equations to fit θ .



3) Go To 1) if not converged.

EM algorithm for Gaussian mixtures



Interpretation

$$\Gamma_k^{(n)} = Q(z^{(n)} = k) = P(z^{(n)} = k | \underline{x}^{(n)}, \theta^{(old)})$$

Compare distributions

$$D_{KL}(Q \parallel P) = \sum_{\underline{z}} Q(\underline{z}) \log \frac{Q(\underline{z})}{P(\underline{z} | \underline{x}, \theta)} \geq 0$$

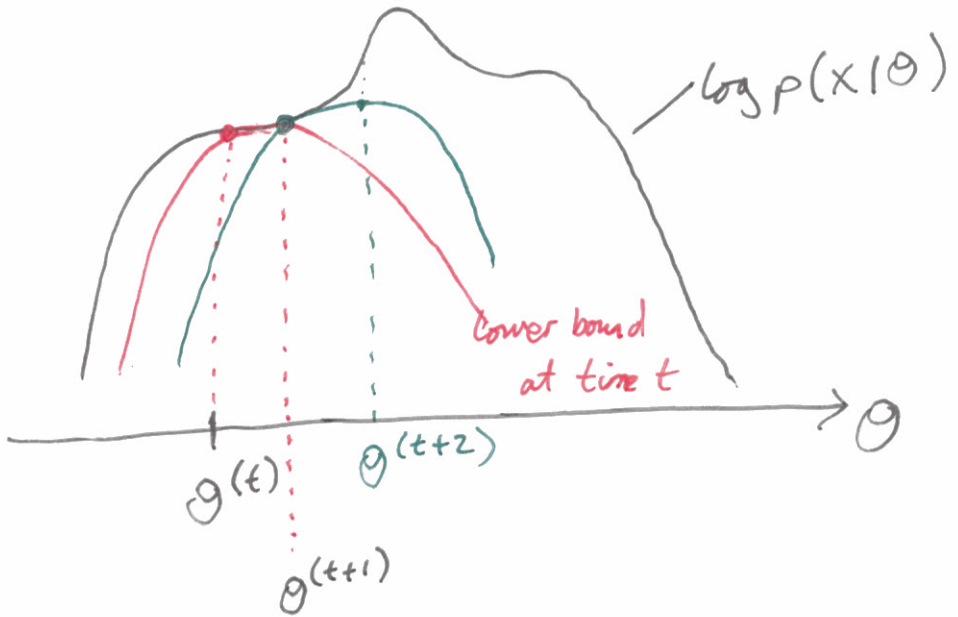
$$\left[P(\underline{z} | \underline{x}, \theta) = \frac{P(\underline{z}, \underline{x} | \theta)}{P(\underline{x} | \theta)} \leftarrow \text{likelihood} \right]$$

$$\Rightarrow \sum_{\underline{z}} Q(\underline{z}) \log \frac{Q(\underline{z})}{P(\underline{z}, \underline{x} | \theta)} \geq -\log P(\underline{x} | \theta)$$

Bound on log-likelihood

Tight if $\theta = \theta^{(old)}$

Bound-based optimizer



If loose bound:

