

# Computational statistics

## Chapter 3: EM algorithm

Thierry Denœux  
Université de technologie de Compiègne

August 2026



- An iterative optimization strategy useful when maximizing the likelihood is difficult, but:
  - ▶ There are **missing** (non-observed) data
  - ▶ If the missing data were observed, maximizing the likelihood would be easy.
- Many applications in statistics and econometrics.
- Can be very simple to implement. Can reliably find an optimum through stable, uphill steps.



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Notation

$\mathbf{Y}$  : Observed data

$\mathbf{Z}$  : Missing data or latent variables

$\mathbf{X}$  : Complete data  $\mathbf{X} = (\mathbf{Y}, \mathbf{Z})$

$\theta$  : Unknown parameter

$L(\theta)$  : observed-data likelihood, short for  $L(\theta; \mathbf{y}) = f(\mathbf{y}; \theta)$

$L_c(\theta)$  : complete-data likelihood, short for  $L(\theta; \mathbf{x}) = f(\mathbf{x}; \theta)$

$\ell(\theta), \ell_c(\theta)$  : observed and complete-data log-likelihoods



# Q function

- Suppose we seek to maximize  $L(\theta)$  with respect to  $\theta$ .
- Define  $Q(\theta, \theta^{(t)})$  to be the expectation of the complete-data log-likelihood, conditional on the observed data  $\mathbf{Y} = \mathbf{y}$ . Namely

$$\begin{aligned} Q(\theta, \theta^{(t)}) &= \mathbb{E}_{\theta^{(t)}} \{ \ell_c(\theta) \mid \mathbf{y} \} \\ &= \mathbb{E}_{\theta^{(t)}} \{ \log f(\mathbf{X}; \theta) \mid \mathbf{y} \} \\ &= \int [\log f(\mathbf{x}; \theta)] f(\mathbf{z} \mid \mathbf{y}; \theta^{(t)}) d\mathbf{z}. \end{aligned}$$

$(f(\mathbf{x} \mid \mathbf{y}; \theta^{(t)}) = f(\mathbf{z} \mid \mathbf{y}; \theta^{(t)})$  because  $\mathbf{Z}$  is the only random part of  $\mathbf{X}$  once we are given  $\mathbf{Y} = \mathbf{y}$ .)



# The EM Algorithm

Start with  $\theta^{(0)}$ . Then

- 1 **E step:** Compute  $Q(\theta, \theta^{(t)})$ .
- 2 **M step:** Maximize  $Q(\theta, \theta^{(t)})$  with respect to  $\theta$ . Set  $\theta^{(t+1)}$  equal to the maximizer of  $Q$ .
- 3 Increment  $t$  and return to the E step unless a stopping criterion has been met; e.g.,

$$\ell(\theta^{(t+1)}) - \ell(\theta^{(t)}) \leq \epsilon$$

or

$$\|\theta^{(t+1)} - \theta^{(t)}\| \leq \epsilon.$$



# Convergence of the EM Algorithm

- It can be proved that  $L(\theta)$  increases after each EM iteration, i.e.,  $L(\theta^{(t+1)}) \geq L(\theta^{(t)})$  for  $t = 0, 1, \dots$
- Consequently, the algorithm converges to a **local maximum** of  $L(\theta)$  if the likelihood function is bounded above.
- Typically, we run the algorithm several times with random initial conditions, and we keep the results of the best run.



## Example: mixture of normal and uniform distributions

- Let  $\mathbf{Y} = (Y_1, \dots, Y_n)$  be an i.i.d. sample from a mixture of a normal distribution  $\mathcal{N}(\mu, \sigma)$  and a uniform distribution  $\mathcal{U}([-a, a])$ , with pdf

$$f(y; \boldsymbol{\theta}) = \pi \phi(y; \mu, \sigma) + (1 - \pi)c \quad (1)$$

where  $\phi(\cdot; \mu, \sigma)$  is the normal pdf,  $c = (2a)^{-1}$  is a known constant,  $\pi$  is the proportion of the normal distribution in the mixture and  $\boldsymbol{\theta} = (\mu, \sigma, \pi)^T$  is the vector of parameters.

- Typically, the uniform distribution corresponds to outliers in the data. The proportion of outliers in the population is then  $1 - \pi$ .
- We want to estimate  $\boldsymbol{\theta}$ .



# Observed and complete-data likelihoods

- Let  $Z_i = 1$  if observation  $i$  is not an outlier,  $Z_i = 0$  otherwise. We have  $Z_i \sim \mathcal{B}(\pi)$ .
- The vector  $\mathbf{Z} = (Z_1, \dots, Z_n)$  is the missing data.
- Observed-data likelihood:

$$L(\boldsymbol{\theta}) = \prod_{i=1}^n f(y_i; \boldsymbol{\theta}) = \prod_{i=1}^n [\pi \phi(y_i; \mu, \sigma) + (1 - \pi)c]$$

- Complete-data likelihood:

$$\begin{aligned} L_c(\boldsymbol{\theta}) &= \prod_{i=1}^n f(y_i, z_i; \boldsymbol{\theta}) = \prod_{i=1}^n f(y_i \mid z_i; \mu, \sigma) f(z_i; \pi) \\ &= \prod_{i=1}^n [\phi(y_i; \mu, \sigma)^{z_i} c^{1-z_i} \pi^{z_i} (1 - \pi)^{1-z_i}] \end{aligned}$$



# Derivation of function $Q$

- Complete-data log-likelihood:

$$\ell_c(\boldsymbol{\theta}) = \sum_{i=1}^n z_i \log \phi(y_i; \mu, \sigma) + \left( n - \sum_{i=1}^n z_i \right) \log c + \sum_{i=1}^n (z_i \log \pi + (1 - z_i) \log(1 - \pi))$$

- It is linear in the  $z_i$ . Consequently, the  $Q$  function is simply

$$Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)}) = \sum_{i=1}^n z_i^{(t)} \log \phi(y_i; \mu, \sigma) + \left( n - \sum_{i=1}^n z_i^{(t)} \right) \log c + \sum_{i=1}^n \left( z_i^{(t)} \log \pi + (1 - z_i^{(t)}) \log(1 - \pi) \right)$$

with  $z_i^{(t)} = \mathbb{E}_{\boldsymbol{\theta}^{(t)}}[Z_i | y_i]$ .



# EM algorithm

E-step: compute

$$\begin{aligned} z_i^{(t)} &= \mathbb{E}_{\theta^{(t)}}[Z_i \mid y_i] = \mathbb{P}_{\theta^{(t)}}[Z_i = 1 \mid y_i] \\ &= \frac{\phi(y_i; \mu^{(t)}, \sigma^{(t)})\pi^{(t)}}{\phi(y_i; \mu^{(t)}, \sigma^{(t)})\pi^{(t)} + c(1 - \pi^{(t)})} \end{aligned}$$

M-step: Maximize  $Q(\theta, \theta^{(t)})$ . We get

$$\pi^{(t+1)} = \frac{1}{n} \sum_{i=1}^n z_i^{(t)}, \quad \mu^{(t+1)} = \frac{\sum_{i=1}^n z_i^{(t)} y_i}{\sum_{i=1}^n z_i^{(t)}}, \quad \sigma^{(t+1)} = \sqrt{\frac{\sum_{i=1}^n z_i^{(t)} (y_i - \mu^{(t+1)})^2}{\sum_{i=1}^n z_i^{(t)}}}$$



# Bayesian posterior mode

- Consider a **Bayesian estimation** problem with likelihood  $L(\boldsymbol{\theta})$  and prior  $f(\boldsymbol{\theta})$ .
- The posterior density is proportional to  $L(\boldsymbol{\theta})f(\boldsymbol{\theta})$ . It can also be maximized by the EM algorithm.
- The E-step requires

$$Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)}) = \mathbb{E}_{\boldsymbol{\theta}^{(t)}} \{ \ell_c(\boldsymbol{\theta}) \mid \mathbf{y} \} + \log f(\boldsymbol{\theta})$$

- The addition of the log-prior often makes it more difficult to maximize  $Q$  during the M-step.
- Some methods can be used to facilitate the M-step in difficult situations (see below).



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Why does it work?

- **Ascent:** Each M-step increases the log-likelihood.
- **Optimization transfer:**

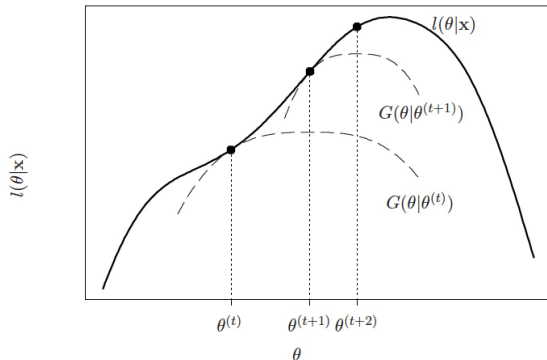
$$\ell(\boldsymbol{\theta}) \geq \underbrace{Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)}) + \ell(\boldsymbol{\theta}^{(t)}) - Q(\boldsymbol{\theta}^{(t)}, \boldsymbol{\theta}^{(t)})}_{G(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)})}$$

- The last two terms in  $G(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)})$  do not depend on  $\boldsymbol{\theta}$ , so  $Q$  and  $G$  are maximized at the same  $\boldsymbol{\theta}$ .
- Further,  $G$  is tangent to  $\ell$  at  $\boldsymbol{\theta}^{(t)}$ , and lies everywhere below  $\ell$ . We say that  $G$  is a **minorizing function** for  $\ell$  (see next slide).
- EM transfers optimization from  $\ell$  to the surrogate function  $G$ , which is more convenient to maximize.



# The nature of EM

- One-dimensional illustration of EM algorithm as a minorization or optimization transfer strategy.
- Each E step forms a minorizing function  $G$ , and each M step maximizes it to provide an uphill step.



# Proof

- We have

$$f(\mathbf{z} \mid \mathbf{y}; \boldsymbol{\theta}) = \frac{f(\mathbf{y}, \mathbf{z}; \boldsymbol{\theta})}{f(\mathbf{y}; \boldsymbol{\theta})} = \frac{f(\mathbf{x}; \boldsymbol{\theta})}{f(\mathbf{y}; \boldsymbol{\theta})} \Rightarrow f(\mathbf{y}; \boldsymbol{\theta}) = \frac{f(\mathbf{x}; \boldsymbol{\theta})}{f(\mathbf{z} \mid \mathbf{y}; \boldsymbol{\theta})}$$

- Consequently,

$$\ell(\boldsymbol{\theta}) = \log f(\mathbf{y}; \boldsymbol{\theta}) = \underbrace{\log f(\mathbf{x}; \boldsymbol{\theta})}_{\ell_c(\boldsymbol{\theta})} - \log f(\mathbf{z} \mid \mathbf{y}; \boldsymbol{\theta})$$

- Taking expectations on both sides wrt the conditional distribution of  $\mathbf{X}$  given  $\mathbf{Y} = \mathbf{y}$  and using  $\boldsymbol{\theta}^{(t)}$  for  $\boldsymbol{\theta}$ :

$$\ell(\boldsymbol{\theta}) = Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)}) - \underbrace{\mathbb{E}_{\boldsymbol{\theta}^{(t)}}[\log f(\mathbf{Z} \mid \mathbf{y}; \boldsymbol{\theta}) \mid \mathbf{y}]}_{H(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)})} \quad (2)$$



# Proof: $\theta^{(t)}$ is a maximizer of $H(\theta, \theta^{(t)})$

- Now, for all  $\theta \in \Theta$ ,

$$H(\theta, \theta^{(t)}) - H(\theta^{(t)}, \theta^{(t)}) = \mathbb{E}_{\theta^{(t)}} \left[ \log \frac{f(\mathbf{Z} | \mathbf{y}; \theta)}{f(\mathbf{Z} | \mathbf{y}; \theta^{(t)})} \mid \mathbf{y} \right] \quad (3a)$$

$$\leq \log \underbrace{\mathbb{E}_{\theta^{(t)}} \left[ \frac{f(\mathbf{Z} | \mathbf{y}; \theta)}{f(\mathbf{Z} | \mathbf{y}; \theta^{(t)})} \mid \mathbf{y} \right]}_{\int \frac{f(\mathbf{z} | \mathbf{y}; \theta)}{f(\mathbf{z} | \mathbf{y}; \theta^{(t)})} f(\mathbf{z} | \mathbf{y}; \theta^{(t)}) d\mathbf{z}} (*) \quad (3b)$$

$$\leq \log \underbrace{\int f(\mathbf{z} | \mathbf{y}; \theta) d\mathbf{z}}_1 = 0 \quad (3c)$$

(\*): from the concavity of the log and Jensen's inequality.

- Hence,  $\theta^{(t)}$  is a maximizer of  $H(\theta, \theta^{(t)})$



# Proof: $\ell(\cdot)$ dominates $G(\cdot, \theta^{(t)})$

Hence, for all  $\theta \in \Theta$ ,

$$\begin{aligned} H(\theta^{(t)}, \theta^{(t)}) &\geq H(\theta, \theta^{(t)}) \\ Q(\theta^{(t)}, \theta^{(t)}) - \ell(\theta^{(t)}) &\geq Q(\theta, \theta^{(t)}) - \ell(\theta) \\ \ell(\theta) &\geq \underbrace{Q(\theta, \theta^{(t)}) + \ell(\theta^{(t)}) - Q(\theta^{(t)}, \theta^{(t)})}_{G(\theta, \theta^{(t)})}. \end{aligned}$$



# Proof: $G$ is tangent to $\ell$ at $\theta^{(t)}$

- As  $\theta^{(t)}$  maximizes  $H(\theta, \theta^{(t)}) = Q(\theta, \theta^{(t)}) - \ell(\theta)$ , we have

$$H'(\theta, \theta^{(t)})|_{\theta=\theta^{(t)}} = Q'(\theta, \theta^{(t)})|_{\theta=\theta^{(t)}} - \ell'(\theta)|_{\theta=\theta^{(t)}} = 0,$$

so

$$Q'(\theta, \theta^{(t)})|_{\theta=\theta^{(t)}} = \ell'(\theta)|_{\theta=\theta^{(t)}}.$$

- Consequently, as  $G(\theta, \theta^{(t)}) = Q(\theta, \theta^{(t)}) + \text{cst}$ ,

$$G'(\theta, \theta^{(t)})|_{\theta=\theta^{(t)}} = Q'(\theta, \theta^{(t)})|_{\theta=\theta^{(t)}} = \ell'(\theta)|_{\theta=\theta^{(t)}}.$$



# Proof: monotonicity

- From (2),

$$\ell(\boldsymbol{\theta}^{(t+1)}) - \ell(\boldsymbol{\theta}^{(t)}) = \underbrace{Q(\boldsymbol{\theta}^{(t+1)}, \boldsymbol{\theta}^{(t)}) - Q(\boldsymbol{\theta}^{(t)}, \boldsymbol{\theta}^{(t)})}_A - \left[ \underbrace{H(\boldsymbol{\theta}^{(t+1)}, \boldsymbol{\theta}^{(t)}) - H(\boldsymbol{\theta}^{(t)}, \boldsymbol{\theta}^{(t)})}_B \right]$$

- $A \geq 0$  because  $\boldsymbol{\theta}^{(t+1)}$  is a maximizer of  $Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)})$ , and  $B \leq 0$  because, from (3),  $\boldsymbol{\theta}^{(t)}$  is a maximizer of  $H(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)})$ .
- Hence,

$$\boxed{\ell(\boldsymbol{\theta}^{(t+1)}) \geq \ell(\boldsymbol{\theta}^{(t)})}$$



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Monte Carlo EM (MCEM)

- Sometimes, the conditional expectation of  $\ell_c(\theta)$  given  $\mathbf{y}$  cannot be easily computed analytically in the E step.
- Approach: randomly generate sets of missing values according to the conditional distribution  $f(\mathbf{z}|\mathbf{y}; \theta^{(t)})$ , and replace the expectation by an average over generated data sets.



# Monte Carlo EM (MCEM)

- Replace the  $t$ -th E step with

- 1 Draw missing datasets  $\mathbf{Z}_1^{(t)}, \dots, \mathbf{Z}_{m^{(t)}}^{(t)}$  i.i.d. from  $f(\mathbf{z}|\mathbf{y}; \boldsymbol{\theta}^{(t)})$ . Each  $\mathbf{Z}_j^{(t)}$  is a vector of all the missing values needed to complete the observed dataset, so  $\mathbf{X}_j^{(t)} = (\mathbf{y}, \mathbf{Z}_j^{(t)})$  denotes a completed dataset where the missing values have been replaced by  $\mathbf{Z}_j^{(t)}$ .
- 2 Calculate

$$\hat{Q}^{(t+1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)}) = \frac{1}{m^{(t)}} \sum_{j=1}^{m^{(t)}} \log f(\mathbf{X}_j^{(t)}; \boldsymbol{\theta}).$$

- Then  $\hat{Q}^{(t+1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)})$  is a **Monte Carlo estimate** of  $Q(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)})$ .
- The M step is modified to maximize  $\hat{Q}^{(t+1)}(\boldsymbol{\theta}, \boldsymbol{\theta}^{(t)})$ .



# Remarks

- It is advised to increase  $m^{(t)}$  as iterations progress to reduce the Monte Carlo variability of  $\hat{Q}$ .
- MCEM will not converge in the same sense as ordinary EM, rather values of  $\theta^{(t)}$  will bounce around the true maximum, with a precision that depends on  $m^{(t)}$ .



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Generalized EM (GEM) algorithm

- In the original EM algorithm,  $\theta^{(t+1)}$  is a maximizer of  $Q(\theta, \theta^{(t)})$ , i.e.,

$$Q(\theta^{(t+1)}, \theta^{(t)}) \geq Q(\theta, \theta^{(t)})$$

for all  $\theta$ .

- However, to ensure convergence, we only need that

$$Q(\theta^{(t+1)}, \theta^{(t)}) \geq Q(\theta^{(t)}, \theta^{(t)})$$

- Any algorithm that chooses  $\theta^{(t+1)}$  at each iteration to guarantee the above condition (without maximizing  $Q(\theta, \theta^{(t)})$ ) is called a **Generalized EM (GEM) algorithm**.



# EM gradient algorithm

- Replace the M step with a single step of Newton's method, thereby approximating the maximum without actually solving for it exactly.
- Instead of maximizing, choose:

$$\begin{aligned}\theta^{(t+1)} &= \theta^{(t)} - \mathbf{Q}''(\theta, \theta^{(t)})^{-1} \Big|_{\theta=\theta^{(t)}} \mathbf{Q}'(\theta, \theta^{(t)}) \Big|_{\theta=\theta^{(t)}} \\ &= \theta^{(t)} - \mathbf{Q}''(\theta, \theta^{(t)})^{-1} \Big|_{\theta=\theta^{(t)}} \ell'(\theta^{(t)})\end{aligned}$$

- Ascent is ensured for canonical parameters in exponential families. Backtracking can ensure ascent in other cases; inflating steps can speed up convergence.



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

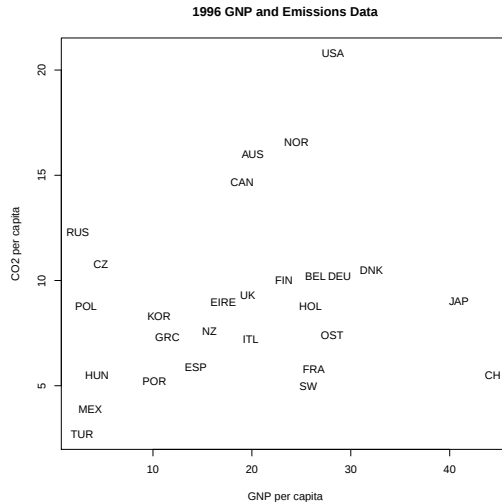
- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Introductory example



## Introductory example (continued)

- The data in the previous slide do not show any clear linear trend.
- However, there seem to be several groups for which a linear model would be a reasonable approximation.
- How to identify those groups and the corresponding linear models?



# Model

- Model: the response variable  $Y$  depends on the input variable  $X$  in different ways, depending on a latent variable  $Z$ . (Beware: we have switched back to the classical notation for regression models!)
- This model is called **mixture of regressions** or **switching regressions**. It has been widely studied in the econometrics literature.
- Model:

$$Y = \begin{cases} \beta_1^T X + \epsilon_1, \epsilon_1 \sim \mathcal{N}(0, \sigma_1) & \text{if } Z = 1, \\ \vdots \\ \beta_K^T X + \epsilon_K, \epsilon_K \sim \mathcal{N}(0, \sigma_K) & \text{if } Z = K, \end{cases}$$

with  $X = (1, X_1, \dots, X_p)$  and  $Z \sim \mathcal{M}(1, \pi_1, \dots, \pi_K)$ , so

$$f(y|X=x) = \sum_{k=1}^K \pi_k \phi(y; \beta^T x, \sigma_k)$$



# Observed and complete-data likelihoods

- Observed-data likelihood:

$$L(\theta) = \prod_{i=1}^n f(y_i | x_i; \theta) = \prod_{i=1}^n \sum_{k=1}^K \pi_k \phi(y_i; \beta_k^T x_i, \sigma_k)$$

- Complete-data likelihood:

$$\begin{aligned} L_c(\theta) &= \prod_{i=1}^n f(y_i, z_i | x_i; \theta) = \prod_{i=1}^n f(y_i | x_i, z_i; \theta) p(z_i | \pi) \\ &= \prod_{i=1}^N \prod_{k=1}^K \phi(y_i; \beta_k^T x_i, \sigma_k)^{z_{ik}} \pi_k^{z_{ik}}, \end{aligned}$$

with  $z_{ik} = 1$  if  $z_i = k$  and  $z_{ik} = 0$  otherwise.



# Derivation of function $Q$

- Complete-data log-likelihood:

$$\ell_c(\theta) = \sum_{i=1}^n \sum_{k=1}^K z_{ik} \log \phi(y_i; \beta_k^T x_i, \sigma_k) + \sum_{i=1}^N \sum_{k=1}^K z_{ik} \log \pi_k$$

- It is linear in the  $z_{ik}$ . Consequently, the  $Q$  function is simply

$$Q(\theta, \theta^{(t)}) = \sum_{i=1}^n \sum_{k=1}^K z_{ik}^{(t)} \log \phi(y_i; \beta_k^T x_i, \sigma_k) + \sum_{i=1}^n \sum_{k=1}^K z_{ik}^{(t)} \log \pi_k$$

with  $z_{ik}^{(t)} = \mathbb{E}_{\theta^{(t)}}[Z_{ik}|y_i] = \mathbb{P}_{\theta^{(t)}}[Z_i = k|y_i]$ .



# EM algorithm

- E-step: compute

$$\begin{aligned} z_{ik}^{(t)} &= \mathbb{P}_{\theta^{(t)}}[Z_i = k | y_i] \\ &= \frac{\phi(y_i; \beta_k^{(t)T} x_i, \sigma_k^{(t)}) \pi_k^{(t)}}{\sum_{\ell=1}^K \phi(y_i; \beta_\ell^{(t)T} x_i, \sigma_\ell^{(t)}) \pi_\ell^{(t)}} \end{aligned}$$

- M-step: Maximize  $Q(\theta, \theta^{(t)})$ . As before, we get

$$\pi_k^{(t+1)} = \frac{N_k^{(t)}}{n},$$

$$\text{with } N_k^{(t)} = \sum_{i=1}^N z_{ik}^{(t)}.$$



## M-step: update of the $\beta_k$ and $\sigma_k$

- In  $Q(\theta, \theta^{(t)})$ , the term depending on  $\beta_k$  is

$$SS_k = \sum_{i=1}^n z_{ik}^{(t)} (y_i - \beta_k^T \mathbf{x}_i)^2.$$

- Minimizing  $SS_k$  w.r.t.  $\beta_k$  is a **weighted least-squares** (WLS) problem. In matrix form,

$$SS_k = (\mathbf{y} - \mathbf{X}\beta_k)^T \mathbf{W}_k (\mathbf{y} - \mathbf{X}\beta_k)$$

with  $\mathbf{W}_k = \text{diag}(z_{1k}^{(t)}, \dots, z_{nk}^{(t)})$ .



## M-step: update of the $\beta_k$ and $\sigma_k$ (continued)

- The solution is the WLS estimate of  $\beta_k$ :

$$\beta_k^{(t+1)} = (\mathbf{X}^T \mathbf{W}_k \mathbf{X})^{-1} \mathbf{X}^T \mathbf{W}_k \mathbf{y}$$

- The value of  $\sigma_k$  maximizing  $Q(\theta, \theta^{(t)})$  is the weighted average of the residuals,

$$\begin{aligned} \sigma_k^{2(t+1)} &= \frac{1}{N_k^{(t)}} \sum_{i=1}^n z_{ik}^{(t)} (y_i - \beta_k^{(t+1)T} x_i)^2 \\ &= \frac{1}{N_k^{(t)}} (\mathbf{y} - \mathbf{X} \beta_k^{(t+1)})^T \mathbf{W}_k (\mathbf{y} - \mathbf{X} \beta_k^{(t+1)}) \end{aligned}$$



# Mixture of regressions using flexmix

```
library(mixtools); library(flexmix)
data("CO2data")
fit1 <- flexmix(CO2 ~ GNP, data = CO2data, k = 1)
loglik_best <- -Inf
for(i in 1:10){
  fit2 <- flexmix(CO2 ~ GNP, data = CO2data, k = 2)
  if(logLik(fit2) > loglik_best){
    loglik_best <- logLik(fit2)
    fit2best <- fit2
  }
}
fit2 <- fit2best
AIC(fit1, fit2)
> AIC(fit1, fit2)
      df      AIC
fit1   3 161.9675
fit2   7 147.9675
```

# Best solution in 10 runs

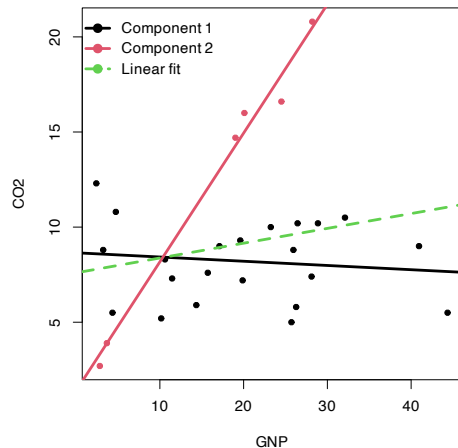
```
par1<-parameters(fit1)
par <- parameters(fit2)
cl <- clusters(fit2)

plot(CO2data$GNP, CO2data$CO2,col = cl,pch = 19,
     cex = 0.8,xlab = "CO2",ylab = "GNP")

abline(a = par["coef.(Intercept)", 1],
       b = par["coef.GNP", 1],col = 1,lwd = 3)

abline(a = par["coef.(Intercept)", 2],
       b = par["coef.GNP", 2],col = 2, lwd = 3)

abline(a = par1["coef.(Intercept)", 1],
       b = par1["coef.GNP", 1],
       col = 3,lwd = 3, lty=2)
```



# Overview

## 1 EM algorithm

- Description
- Analysis

## 2 Some variants

- Facilitating the E-step
- Facilitating the M-step

## 3 Application to Regression models

- Mixture of regressions
- Mixture of experts



# Making the mixing proportions predictor-dependent

- An interesting extension of the previous model is to assume the proportions  $\pi_k$  to be partially explained by a vector of **concomitant variables**  $W$ .
- If  $W = X$ , we can approximate the regression function by different linear functions in different regions of the predictor space.
- In ML, this method is referred to as the **mixture of experts** methods.
- A useful parametric form for  $\pi_k$  that ensures  $\pi_k \geq 0$  and  $\sum_{k=1}^K \pi_k = 1$  is the multinomial logit model

$$\pi_k(w, \alpha) = \frac{\exp(\alpha_k^T w)}{\sum_{\ell=1}^K \exp(\alpha_\ell^T w)}$$

with  $\alpha = (\alpha_1, \dots, \alpha_K)$  and  $\alpha_1 = 0$ .



# EM algorithm

- The  $Q$  function is the same as before, except that the  $\pi_k$  now depend on the  $w_i$  and parameter  $\alpha$ :

$$Q(\theta, \theta^{(t)}) = \sum_{i=1}^n \sum_{k=1}^K z_{ik}^{(t)} \log \phi(y_i; \beta_k^T x_i, \sigma_k) + \sum_{i=1}^n \sum_{k=1}^K z_{ik}^{(t)} \log \pi_k(w_i, \alpha)$$

- In the M-step, the update formula for  $\beta_k$  and  $\sigma_k$  are unchanged.
- No closed-form expression for

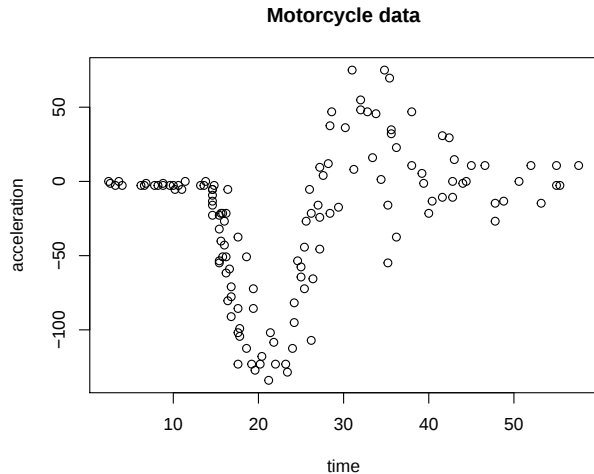
$$\hat{\alpha} = \arg \max_{\alpha} \sum_{i=1}^n \sum_{k=1}^K z_{ik}^{(t)} \log \pi_k(w_i, \alpha)$$

- **GEM algorithm**: perform a single step of Newton's algorithm. Backtracking can be used to ensure ascent.



## Example: motorcycle data

```
library("MASS")  
x<-mcycle$times  
y<-mcycle$accel  
plot(x,y)
```



# Mixture of experts using flexmix

```
library(flexmix)

K<-5

res <- flexmix(y ~ x,k=K,model=FLXMRglm(family="gaussian"),
               concomitant=FLXPmultinom(formula=~x))

beta <- parameters(res)[1:2,]
alpha <- res@concomitant@coef
```



# Plotting the posterior probabilities

```
xt <- seq(0,60,0.1)
Nt <- length(xt)
plot(x,y)

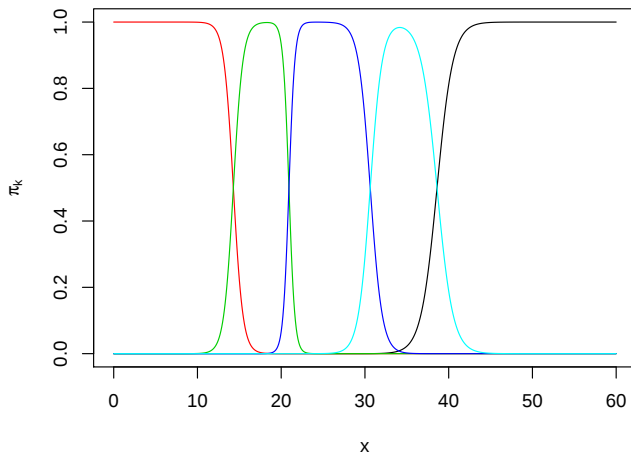
pit <- matrix(0,Nt,K)
for(k in 1:K) pit[,k] <- exp(alpha[1,k]+alpha[2,k]*xt)
pit <- pit/rowSums(pit)

plot(xt,pit[,1],type="l",col=1)
for(k in 2:K) lines(xt,pit[,k],col=k)
```



# Posterior probabilities

Motorcycle data – posterior probabilities



# Plotting the predictions

```
yhat <- rep(0,Nt)
for(k in 1:K) yhat <- yhat+pit[,k]*(beta[1,k]+beta[2,k]*xt)

plot(x,y,main="Motorcycle data",xlab="time",ylab="acceleration")

for(k in 1:K) abline(beta[1:2,k],lty=2)
lines(xt,yhat,col='red',lwd=2)
```



# Regression lines and predictions

