

Lecture Notes on the Expectation-Maximization Algorithm

John C. Chao

Econ 721 Lecture Notes

November 22, 2022

- **A Motivating Example:** Suppose we want to estimate the following two-component mixture model:

Let

$$W_i = (B_i, Y_{1,i}, Y_{2,i})' \text{ for } i = 1, \dots, n.$$

Suppose that

$$\{W_i\} \equiv i.i.d.$$

and suppose that

$$\begin{aligned} Y_{1,i} &\sim N(\mu_1, \sigma_1^2), \quad Y_{2,i} \sim N(\mu_2, \sigma_2^2), \\ B_i &= \begin{cases} 1 & \text{with prob } p \\ 0 & \text{with prob } 1 - p \end{cases} \end{aligned}$$

Moreover, let

$$Y_i = (1 - B_i) Y_{1,i} + B_i Y_{2,i}$$

and assume that B_i , $Y_{1,i}$, and $Y_{2,i}$ are mutually independent and that we only observe Y_i and not B_i , $Y_{1,i}$, and $Y_{2,i}$ separately.

- **Remarks:**

- ① We can think of the underlying generating mechanism is one where a Bernoulli random variable $B_i \in \{0, 1\}$ is first generated with probability p ; and, then, depending on the outcome delivers either $Y_{1,i}$ or $Y_{2,i}$.
- ② Many economic/econometric models come in the form of a mixture (e.g. models of regime switching, unobserved heterogeneity, etc.)

- **Probability Density Function:**

The pdf of Y_i is given by

$$g_Y(y) = (1 - p) \phi_{\theta_1}(y) + p \phi_{\theta_2}(y),$$

where for $s = 1, 2$

$$\begin{aligned}\theta_s &= (\mu_s, \sigma_s^2)', \\ \phi_{\theta_s}(y) &= \frac{1}{\sigma_s \sqrt{2\pi}} \exp \left\{ -\frac{1}{2\sigma_s^2} (y - \mu_s)^2 \right\}.\end{aligned}$$

- **Additional Notations:** Further define

$$\theta = \begin{pmatrix} p \\ \theta_1 \\ \theta_2 \end{pmatrix} = \begin{pmatrix} p \\ \mu_1 \\ \sigma_1^2 \\ \mu_2 \\ \sigma_2^2 \end{pmatrix} \text{ and } Y = \begin{pmatrix} Y_1 \\ Y_2 \\ \vdots \\ Y_n \end{pmatrix}.$$

- **Log-likelihood Function:** Given the notations developed above, the log-likelihood function for this model can be written as

$$\ell(\theta, Y) = \sum_{i=1}^n \ln \left\{ (1-p) \phi_{\theta_1}(y) + p \phi_{\theta_2}(y) \right\}.$$

Note that this log-likelihood function will be difficult to maximize (say, by Newton's method) because of the sum of terms within the log function.

- **The Case Where B_i Is Observed:** To understand the idea behind the expectation-maximization algorithm, consider the case where we could observe the values of B_i (the complete data case). Then, the problem would be a lot easier, since if $B_i = 1$; then, Y_i comes from model 2; otherwise, it comes from model 1. Hence, if the values of B_i are observable, then the probability density function of the data would take the form

$$\begin{aligned} & f(B_i, Y_i | \theta) \\ = & f(Y_i | B_i, \theta_1, \theta_2) f(B_i | p) \\ = & \left\{ \left[\phi_{\theta_1}(y_i) \right]^{(1-B_i)} \left[\phi_{\theta_2}(y_i) \right]^{B_i} \right\} \left\{ [1-p]^{(1-B_i)} p^{B_i} \right\} \end{aligned}$$

- Hence, in this case, the likelihood function and the log-likelihood function can be written down, respectively, as

$$L_0(\theta, B, Y) = \prod_{i=1}^n \left[\phi_{\theta_1}(y_i) \right]^{(1-B_i)} \left[\phi_{\theta_2}(y_i) \right]^{B_i} [1-p]^{(1-B_i)} p^{B_i},$$

$$\begin{aligned} \ell_0(\theta, B, Y) &= \sum_{i=1}^n \left\{ (1-B_i) \ln \phi_{\theta_1}(y_i) + B_i \ln \phi_{\theta_2}(y_i) \right\} \\ &\quad + \sum_{i=1}^n \left\{ (1-B_i) \ln (1-p) + B_i \ln p \right\}, \end{aligned}$$

where $B = (B_1, \dots, B_n)'$.

EM Algorithm

- Since in reality the values of B_i are typically unknown, the EM algorithm proceeds in an iterative manner, substituting for each B_i with its expected value

$$\begin{aligned}\gamma_i(\theta) &= E[B_i | \theta, Y] \\ &= E[B_i | \theta, Y_i] \quad (\text{by independence}) \\ &= \Pr(B_i = 1 | \theta, Y_i) \\ &= \frac{f(B_i = 1 \cap Y_i | \theta)}{f(Y_i | \theta)} \\ &= \frac{f(B_i = 1 | p) f(Y_i | B_i = 1, \theta)}{f(Y_i | \theta)} \\ &= \frac{p \phi_{\theta_2}(y_i)}{(1 - p) \phi_{\theta_1}(y_i) + p \phi_{\theta_2}(y_i)}.\end{aligned}$$

- **Remark:** The quantity $\gamma_i(\theta) = E[B_i | \theta, Y]$ is often called the **responsibility** of model 2 for observation i .

- In addition, taking conditional expectation of the log-likelihood function, we get

$$\begin{aligned} & E \left[\ell_0 (\theta, B, Y) \mid \hat{\theta}, Y \right] \\ &= \sum_{i=1}^n \left\{ \left(1 - E \left[B_i \mid \hat{\theta}, Y_i \right] \right) \ln \phi_{\theta_1} (y_i) + E \left[B_i \mid \hat{\theta}, Y_i \right] \ln \phi_{\theta_2} (y_i) \right\} \\ &\quad + \sum_{i=1}^n \left\{ \left(1 - E \left[B_i \mid \hat{\theta}, Y_i \right] \right) \ln (1 - p) + E \left[B_i \mid \hat{\theta}, Y_i \right] \ln p \right\} \\ &= \sum_{i=1}^n \left\{ \left(1 - \gamma_i \left(\hat{\theta} \right) \right) \ln \phi_{\theta_1} (y_i) + \gamma_i \left(\hat{\theta} \right) \ln \phi_{\theta_2} (y_i) \right\} \\ &\quad + \sum_{i=1}^n \left\{ \left(1 - \gamma_i \left(\hat{\theta} \right) \right) \ln (1 - p) + \gamma_i \left(\hat{\theta} \right) \ln p \right\} \end{aligned}$$

- **Remark:** The EM algorithm iterates back and forth between an expectation step and a maximization step. Under the expectation step, we do a soft assignment of each observation to each model, i.e., the current estimates of the parameters are used to assign responsibilities according to the relative density (under each model) of the sample points. Under the maximization step these responsibilities are used to construct a weighted log-likelihood, which we then maximize to update our estimates of the parameters.

- More precisely, the EM algorithm goes as follows:

Step 1: Take initial estimates of the parameters

$\hat{\mu}_{1,0}, \hat{\sigma}_{1,0}^2, \hat{\mu}_{2,0}, \hat{\sigma}_{2,0}^2, \hat{p}_0$ (to be specified below).

Step 2: (Expectation Step) In the k^{th} step, compute the responsibilities

$$\hat{\gamma}_{i,k} = \frac{\hat{p}_{k-1} \phi_{\hat{\theta}_{2,k-1}}(y_i)}{(1 - \hat{p}_{k-1}) \phi_{\hat{\theta}_{1,k-1}}(y_i) + \hat{p}_{k-1} \phi_{\hat{\theta}_{2,k-1}}(y_i)} \text{ for } i = 1, \dots, n;$$

where $\hat{\theta}_{1,k-1} = (\hat{\mu}_{1,k-1}, \hat{\sigma}_{1,k-1}^2)'$ and $\hat{\theta}_{2,k-1} = (\hat{\mu}_{2,k-1}, \hat{\sigma}_{2,k-1}^2)$.

- **Step 3:** (Maximization Step) Compute the weighted means and variances for the $(k + 1)^{th}$ step as

$$\begin{aligned}\hat{\mu}_{1,k} &= \frac{\sum_{i=1}^n (1 - \hat{\gamma}_{i,k}) y_i}{\sum_{i=1}^n (1 - \hat{\gamma}_{i,k})}, & \hat{\sigma}_{1,k}^2 &= \frac{\sum_{i=1}^n (1 - \hat{\gamma}_{i,k}) (y_i - \hat{\mu}_{1,k})^2}{\sum_{i=1}^n (1 - \hat{\gamma}_{i,k})} \\ \hat{\mu}_{2,k} &= \frac{\sum_{i=1}^n \hat{\gamma}_{i,k} y_i}{\sum_{i=1}^n \hat{\gamma}_{i,k}}, & \hat{\sigma}_{2,k}^2 &= \frac{\sum_{i=1}^n \hat{\gamma}_{i,k} (y_i - \hat{\mu}_{2,k})^2}{\sum_{i=1}^n \hat{\gamma}_{i,k}}\end{aligned}$$

Also, compute the mixing probability as

$$\hat{p}_{k+1} = \frac{1}{n} \sum_{i=1}^n \hat{\gamma}_{i,k}.$$

Step 4: Iterate steps 2 and 3 until convergence.

- **Remark:** Initial estimates $\hat{\mu}_{1,0}$ and $\hat{\mu}_{2,0}$ could be made by simply choosing two of the y_i 's. $\hat{\sigma}_{1,0}^2$ and $\hat{\sigma}_{2,0}^2$ could both be set equal to the overall sample variance

$$\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2$$

and the initial mixing proportion $\hat{p}_0$ can be set to 0.5.

EM Algorithm - General Formulation

- More generally, EM algorithms are useful for problems for which direct maximization of the likelihood function is difficult, but could be made easier by enlarging the sample with latent (unobserved) data. Suppose we have

Y — observed data,

θ — parameter vector,

$\ell(\theta, Y)$ - log-likelihood associated with observed data

and

Z — latent or missing data,

$W = (Y, Z)$ — complete data

$\ell_0(\theta; W)$ — log-likelihood associated with complete data

EM Algorithm - General Formulation

- The more general EM algorithm goes as follows:
- **Step 1:** Take initial estimate of the parameter vector $\hat{\theta}^{(0)}$
- **Step 2:** (Expectation Step) In the k^{th} step, compute the responsibilities

$$Q\left(\theta', \hat{\theta}^{(k-1)}\right) = E\left[\ell_0\left(\theta', W\right) \mid Y, \hat{\theta}^{(k-1)}\right].$$

as a function of the dummy argument θ' .

Step 3: (Maximization Step) Determine $\hat{\theta}^{(k)}$ as

$$\hat{\theta}^{(k)} = \arg \max_{\theta'} Q\left(\theta', \hat{\theta}^{(k-1)}\right)$$

Step 4: Iterate steps 2 and 3 until convergence.

EM Algorithm - General Formulation

- **Remark:** To see why the EM algorithm works, note that

$$f(Z|Y, \theta') = \frac{f(Z, Y|\theta')}{f(Y|\theta')} = \frac{f(W|\theta')}{f(Y|\theta')},$$

so that

$$f(Y|\theta') = \frac{f(W|\theta')}{f(Z|Y, \theta')}$$

In terms of log-likelihoods, this implies that

$$\begin{aligned}\ell(\theta'; Y) &= \ell_0(\theta'; W) - \ln f(Z|Y, \theta') \\ &= \ell_0(\theta'; W) - \ell_1(Y, \theta'|Z)\end{aligned}$$

Taking the conditional expectation with respect to the distribution of $W|Y$ evaluated at the parameter value $\hat{\theta}$, we get

$$\begin{aligned}\ell(\theta'; Y) &= E[\ell(\theta'; Y) | Y, \hat{\theta}] \\ &= E[\ell_0(\theta'; W) | Y, \hat{\theta}] - E[\ell_1(Y, \theta'|Z) | Y, \hat{\theta}].\end{aligned}$$

EM Algorithm - General Formulation

- or

$$\ell(\theta'; Y) = Q(\theta', \hat{\theta}) - R(\theta', \hat{\theta}).$$

- **Claim:** $R(\theta', \hat{\theta})$ is maximized at $R(\hat{\theta}, \hat{\theta})$.

- **Proof of Claim:** Consider

$$\begin{aligned} & R(\theta', \hat{\theta}) - R(\hat{\theta}, \hat{\theta}) \\ &= E \left[\ell_1(Y, \theta' | Z) \mid Y, \hat{\theta} \right] - E \left[\ell_1(Y, \hat{\theta} | Z) \mid Y, \hat{\theta} \right] \\ &= E \left(\ln \left[\frac{f(Z | Y, \theta')}{f(Z | Y, \hat{\theta})} \right] \mid Y, \hat{\theta} \right) \\ &\leq \ln \left(E \left[\frac{f(Z | Y, \theta')}{f(Z | Y, \hat{\theta})} \right] \mid Y, \hat{\theta} \right) \\ &= \ln \int f(Z | Y, \theta') dZ = \ln 1 = 0. \end{aligned}$$

EM Algorithm - General Formulation

- Using this result, it follows that if $\hat{\theta}^{(k)}$ maximizes $Q(\theta', \hat{\theta}^{(k-1)})$; then,

$$\begin{aligned} & \ell(\theta'; Y) \big|_{\theta' = \hat{\theta}^{(k)}} - \ell(\theta'; Y) \big|_{\theta' = \hat{\theta}^{(k-1)}} \\ = & E \left[\ell(\theta'; Y) \mid Y, \hat{\theta}^{(k-1)} \right] \big|_{\theta' = \hat{\theta}^{(k)}} - E \left[\ell(\theta'; Y) \mid Y, \hat{\theta}^{(k-1)} \right] \big|_{\theta' = \hat{\theta}^{(k-1)}} \\ = & Q(\hat{\theta}^{(k)}, \hat{\theta}^{(k-1)}) - Q(\hat{\theta}^{(k-1)}, \hat{\theta}^{(k-1)}) \\ & - \left[R(\hat{\theta}^{(k)}, \hat{\theta}^{(k-1)}) - R(\hat{\theta}^{(k-1)}, \hat{\theta}^{(k-1)}) \right] \\ \geq & 0. \end{aligned}$$