

Maximum Likelihood Estimation & Inference

Proposed by R. A. Fisher (1921/1925)

Controversial - contrary to Karl Pearson's "method of moments"

Consider: random data $\{x_1, \dots, x_n\}$

parameter vector

$$\underline{\theta} = (\theta_1, \dots, \theta_k)'$$

Objective: Estimate $\underline{\theta}$

Role of $\underline{\theta}$: $\underline{\theta}$ determines the x_i values through their joint p.d.f., $p(\underline{x} | \underline{\theta})$, or

$$p(x_1, \dots, x_n | \theta_1, \dots, \theta_k):$$

We can view $p(\underline{x} | \underline{\theta})$ in two ways -

(i) As a fcn. of $\{x_1, \dots, x_n\}$, given $\underline{\theta}$.

(ii) As a fcn. of $(\theta_1, \dots, \theta_k)$, given $\underline{x}$.

Latter is called the Likelihood Function.

$$L(\underline{\theta}) = L(\underline{\theta} | x_1, \dots, x_n) = p(x_1, \dots, x_n | \underline{\theta}).$$

Definition: The Maximum Likelihood

Estimator of $\underline{\theta}$ (say, $\hat{\underline{\theta}}$) is that value of $\underline{\theta}$ such that $L(\hat{\underline{\theta}}) \geq L(\underline{\theta})$, for all other $\underline{\theta}$.

Motivation: Given a particular sample of data, what value of $\underline{\theta}$ is most likely to have generated it from the population? This is $\hat{\underline{\theta}}$.

Note: (i) $\hat{\underline{\theta}}$ need not be unique.

(ii) $\hat{\underline{\theta}}$ should locate the global max of L .

(iii) If sample data are independent

$$\text{then } L(\underline{\theta} | \underline{x}) = p(\underline{x} | \underline{\theta})$$

$$= \prod_i p(x_i | \theta)$$

(iv) Any monotonic transformation of $L(\underline{\theta})$ leaves location of extrem unchanged. (e.g.: $\log L(\underline{\theta})$)

3.

Some Basic Concepts & Notation:

(i) "Gradient Vector" : $\begin{bmatrix} \frac{\partial \log L(\theta)}{\partial \theta} \end{bmatrix}$
 ("Score Vector")

(ii) "Hessian Matrix" : $\begin{bmatrix} \frac{\partial^2 \log L(\theta)}{\partial \theta \partial \theta'} \end{bmatrix}$
 (KX1) (KX1)
 (KXC)

(iii) "Likelihood Equation(s)":

$$\frac{\partial \log L(\theta)}{\partial \theta} = \mathbf{0} \quad (KX1)$$

To obtain the MLE, $\hat{\theta}$, we solve the Likelihood Equations, & then check the second-order condition(s) to make sure that we have maximized (not minimized) $L(\theta)$. If the Hessian matrix is at least negative semi-definite, then $\log L$ is concave, & this is sufficient (not necessary) for a maximum.

4

Example:

Geometric Distribution.

Independent Bernoulli trials Y_i = no. of failures before first success:

$$P_r(Y_i = y_i) = (1-p)^{y_i} p$$

where $p = P_r(\text{success})$

$$L = P(Y_1, Y_2, \dots, Y_n) = \prod_{i=1}^n P(Y_i)$$

$$= (1-p)^{\sum y_i} p^n \quad (\sum y_i \text{ sufficient})$$

$$\log L = \sum y_i \log(1-p) + n \log(p)$$

$$(\partial \log L / \partial p) = -\frac{n \bar{y}}{1-p} + n/p \quad (i)$$

$$(\partial^2 \log L / \partial p^2) = -\frac{n \bar{y}}{(1-p)^2} - n/p^2 \quad (ii)$$

$$\text{From (i): } -\frac{n \bar{y}}{1-p} + n/p = 0$$

$$\Rightarrow \bar{p} = \left(\frac{1}{1 + \bar{y}} \right) \quad \#$$

From (ii): $(\partial^2 \log L / \partial p^2) < 0$, everywhere.

[Show that $E(Y_i) = (1-p)/p$. Note that $(1-\bar{p})/\bar{p} = \bar{y}$, so $\bar{y}$ is a natural estimator of $E(Y_i)$ here. Relate this to "invariance" later.

Example:

$$y_i = \beta x_i + \varepsilon_i \quad ; \quad \varepsilon_i \sim \text{iid } N(0, \sigma^2)$$

$$L(\beta, \sigma^2) = p(y_1, \dots, y_n | \beta, \sigma^2)$$

$$= \prod_{i=1}^n p(y_i | \beta, \sigma^2)$$

$y_i \sim \text{i.i.d. } N(\beta x_i, \sigma^2).$

$$\text{So, } L = \prod_{i=1}^n \left(\frac{1}{\sigma \sqrt{2\pi}} \right) \exp \left[-\frac{1}{2\sigma^2} (y_i - \beta x_i)^2 \right]$$

$$= (2\pi \sigma^2)^{-n/2} \exp \left[-\frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - \beta x_i)^2 \right]$$

$$\log L = -n/2 \log(2\pi) - n/2 \log(\sigma^2)$$

$$- \frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - \beta x_i)^2.$$

$$(i) \left(\frac{\partial \log L}{\partial \beta} \right) = \frac{1}{\sigma^2} \sum_i (x_i y_i - \beta x_i^2) = 0$$

$$(ii) \left(\frac{\partial \log L}{\partial \sigma^2} \right) = -\left(\frac{n}{2\sigma^2} \right) + \left(\frac{1}{2\sigma^4} \right) \sum_i (y_i - \beta x_i)^2 = 0$$

From (i): $\sum_i x_i y_i = \beta \sum_i x_i^2$

or,

$$\boxed{\tilde{\beta} = \left(\sum_i x_i y_i \right) / \left(\sum_i x_i^2 \right)}.$$

Substituting in (ii):

$$\left(\frac{n}{2\sigma^2} \right) = \left(\frac{1}{2\sigma^4} \right) \sum_i (y_i - \tilde{\beta} x_i)^2$$

or,

$$\boxed{\tilde{\sigma}^2 = \frac{1}{n} \sum_i (y_i - \tilde{\beta} x_i)^2}$$

Now, consider 2nd-order condition:

$$(i) \left(\frac{\partial^2 \log L}{\partial \beta^2} \right) = -\frac{1}{\sigma^2} \sum_i x_i^2$$

$$(ii) \left(\frac{\partial^2 \log L}{\partial \beta \partial \sigma^2} \right) = -\frac{1}{\sigma^4} \sum_i (x_i y_i - \beta x_i^2)$$

$$= \left(\frac{\partial^2 \log L}{\partial \sigma^2 \partial \beta} \right)$$

$$(iii) \left(\frac{\partial^2 \log L}{\partial \sigma^2 \partial \sigma^2} \right) = \left(\frac{n}{2\sigma^4} \right) - \left(\frac{1}{\sigma^6} \right) \sum_i (y_i - \beta x_i)^2$$

Substituting $\tilde{\beta}$ for β , $\tilde{\sigma}^2$ for σ^2 :

$$(a) \left(\frac{\partial^2 \log L}{\partial \beta^2} \right) \Big|_{\tilde{\beta}, \tilde{\sigma}^2} = -\frac{1}{\tilde{\sigma}^2} \sum_i x_i^2$$

$$(b) \left(\frac{\partial^2 \log L}{\partial \beta \partial \sigma^2} \right) \Big|_{\tilde{\beta}, \tilde{\sigma}^2} = -\frac{1}{\tilde{\sigma}^4} \sum_i (x_i y_i - \tilde{\beta} x_i^2)$$

$$= -\frac{1}{\tilde{\sigma}^4} \left[\sum_i x_i y_i - \sum_i x_i y_i \right] = 0.$$

$$\begin{aligned}
 (b) \quad \left(\frac{\partial^2 \log L}{\partial \sigma^2 \partial \sigma^2} \right) \Big|_{\beta, \sigma^2} &= \left(\frac{n}{2\sigma^4} \right) - \frac{1}{\sigma^6} \sum_i (y_i - \hat{\beta} x_i)^2 \\
 &= (n/2\sigma^4) - (n\sigma^2/\sigma^6) = - (n/2\sigma^4).
 \end{aligned}$$

So, the Hessian Matrix is:

$$H(\beta, \sigma^2) = \begin{bmatrix} -\sum_i x_i^2 / \sigma^2 & 0 \\ 0 & -n/2\sigma^4 \end{bmatrix}$$

$$|a_{11}| = a_{11} = -\sum_i x_i^2 / \sigma^2 < 0$$

$$|H| = (-\sum_i x_i^2 / \sigma^2) (-n/2\sigma^4) > 0$$

So, H is negative-definite.

Obviously generalizes to multiple regression model:

$$\hat{\beta} = (X'X)^{-1} X'y$$

$$\hat{\sigma}^2 = (y - X\hat{\beta})'(y - X\hat{\beta}) / n.$$

Sometimes the information we have is about some random variable, $\underline{z}$, not the observed $\underline{x}$. We need to derive $p(\underline{x})$ from $p(\underline{z})$ in order to form $L(\theta | \underline{x})$.

Example — we derived density of y_i from assumed density of ε_i in regression example.

Definition: Let $\underline{z}$ and $\underline{x}$ be functionally related random vectors, each $(n \times 1)$.

Then, to derive $p(\underline{x})$ from $p(\underline{z})$, we use:

$$p(\underline{x}) = p(\underline{z}) J_+$$

$$\text{where } J_+ = \text{abs.} \left| \frac{\partial \underline{z}_i}{\partial \underline{x}_j} \right| \quad (n \times n)$$

is the Jacobian of the mapping from $\underline{z}$ to $\underline{x}$

(In regression example $\partial \varepsilon_i / \partial y_i = 1$;

a $\partial \varepsilon_i / \partial y_j = 0$ ($i \neq j$), so $J_+ = 1$.)

9.

Example:

$$p(z) = \frac{1}{2} ; \quad 0 \leq z \leq 2$$

= 0 ; otherwise

$$\text{Let } x = 4z ; \quad z = x/4 ; \quad \frac{dz}{dx} = \frac{1}{4}.$$

$$\text{So, } p(x) = p(z) \left| \frac{1}{4} \right| = \left(\frac{1}{2} \right) \left(\frac{1}{4} \right) = \frac{1}{8}$$

$$\text{or } 0 \leq x \leq 8.$$

Example:

$$z \sim N(0, 1) ; \quad -\infty < z < \infty$$

$$x = \sigma z + \mu$$

$$\text{So, } z = (x - \mu) / \sigma ; \quad \frac{\partial z}{\partial x} = \frac{1}{\sigma}$$

$$p(x) = p(z) \left| \frac{1}{\sigma} \right|$$

$$= \left(\frac{1}{\sigma} \right) \left| \frac{1}{\sqrt{2\pi}} \right| e^{-z^2/2}$$

So, if $\sigma > 0$:

$$p(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2} \left(\frac{x - \mu}{\sigma} \right)^2}$$

$$= \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2\sigma^2} (x - \mu)^2} ; \quad -\infty < x < \infty$$

$$\sim N[\mu, \sigma^2].$$

10.

Recall that $\mathcal{L}(\Theta | \mathbf{x})$ is the joint data density, viewed as a fctn. of Θ . If the x_i 's are independent, then we have $\mathcal{L}(\Theta | \mathbf{x}) = \prod_i p(x_i | \Theta)$. If not, we need to know the joint data density. For example

Definition:

Let $\tilde{\mathbf{x}}$ be an $(n \times 1)$ random vector. Then $\tilde{\mathbf{x}}$ is Multivariate Normally distributed, $N(\mu, V)$, if:

$$p(\tilde{\mathbf{x}} | \mu, V) = (2\pi)^{-n/2} |V|^{-1/2} \exp \left[-\frac{1}{2} (\tilde{\mathbf{x}} - \mu)' V^{-1} (\tilde{\mathbf{x}} - \mu) \right]$$

(Clearly V must be p.d.s. Also, if $n=1$ then $V = \sigma^2$ & we get usual normal density.)

Example:

$$\mathbf{y} = \mathbf{X}\beta + \mathbf{E}$$

$$\mathbf{E} \sim MVN(0, \sigma^2 \mathbf{\Omega}) ; \quad \mathbf{\Omega} \text{ known}$$

$$\text{So, } p(\mathbf{E}) = (2\pi)^{-n/2} |\sigma^2 \mathbf{\Omega}|^{-1/2} \exp \left[-\frac{1}{2} \mathbf{E}' (\sigma^2 \mathbf{\Omega})^{-1} \mathbf{E} \right]$$

Consider the Jacobian for $\mathbf{E} \rightarrow \mathbf{y}$:

$$\partial \mathbf{E}_i / \partial \mathbf{y}_i = 1 ; \quad \partial \mathbf{E}_i / \partial \mathbf{y}_j = 0 \quad (i \neq j)$$

$$\text{So } \mathbf{J} = \text{abs. } |\mathbf{I}_n| = 1.$$

So, $p(y|\beta, \sigma^2, \Omega) = (2\pi)^{-n/2} |\sigma^2 \Omega|^{-1/2}$

$$= (2\pi\sigma^2)^{-n/2} |\Omega|^{-1/2} \exp\left[-\frac{1}{2} (y - X\beta)' (\sigma^2 \Omega)^{-1} (y - X\beta)\right]$$

$$= L(\beta, \sigma^2 | y);$$

$$\log L = -\frac{n}{2} \log 2\pi - \frac{n}{2} \log(\sigma^2) - \frac{1}{2} \log |\Omega|$$

$$- \frac{1}{2\sigma^2} (y - X\beta)' \Omega^{-1} (y - X\beta)$$

$$(i) \left(\frac{\partial \log L}{\partial \beta} \right) = -\frac{1}{2\sigma^2} \frac{\partial}{\partial \beta} \left[y' \Omega^{-1} y + \beta' X' \Omega^{-1} X \beta \right. \\ \left. - 2y' \Omega^{-1} X \beta \right]$$

$$= -\frac{1}{2\sigma^2} [2X' \Omega^{-1} X \beta - 2X' \Omega^{-1} y]$$

$$\Rightarrow \boxed{\hat{\beta} = (X' \Omega^{-1} X)^{-1} X' \Omega^{-1} y}$$

$$(ii) \left(\frac{\partial \log L}{\partial \sigma^2} \right) = -\frac{n}{2} \left(\frac{1}{\sigma^2} \right) - \left(\frac{1}{2} \right) (-1) (\sigma^2)^{-2} \cdot (y - X\beta)' \Omega^{-1} (y - X\beta)$$

= 0

$$\Rightarrow \left(\frac{1}{2\sigma^2} \right) = \left(\frac{1}{2\sigma^2} \right) (y - X\hat{\beta})' \Omega^{-1} (y - X\hat{\beta})$$

$$\Rightarrow \boxed{\hat{\sigma}^2 = \frac{1}{n} (y - X\hat{\beta})' \Omega^{-1} (y - X\hat{\beta})}$$

(Confirm 2nd-order condition.)

Note: (i) $\hat{\beta} = GLS$ estimator of β .

(ii) $\hat{\sigma}^2$ is a biased estimator.

(iii) If $\Omega = I$, we just have OLS

This is for the case where Ω is known.

Usually, Ω is unknown, but $\Omega = \Omega(\theta)$

where θ has small dimension. Then we

$$\text{solve : } \left. \begin{aligned} \left(\frac{\partial \log L}{\partial \beta} \right) &= 0 \\ \left(\frac{\partial \log L}{\partial \sigma^2} \right) &= 0 \\ \left(\frac{\partial \log L}{\partial \theta} \right) &= 0 \end{aligned} \right\}$$

Simultaneously, for the MLE's of all of the parameters.

e.g. : AR(1) errors -

$$\Omega = \Omega(\rho) = \frac{1}{1-\rho^2}$$

$$\begin{bmatrix} 1 & \rho & \rho^2 & \dots & \rho^{n-1} \\ \rho & 1 & \rho & \dots & \rho^{n-2} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho^{n-1} & \rho^{n-2} & \dots & \dots & 1 \end{bmatrix}$$

13

Sometimes we can determine the MLE by inspection - avoid lots of calculus!

Example: $y = X\beta + \varepsilon$; $\varepsilon \sim N(0, \sigma^2 I_n)$

Estimate β , s.t. $R\beta = q$.

$$L(\beta, \sigma^2) = \left(\frac{1}{\sigma\sqrt{2\pi}}\right)^n \exp\left[-\frac{1}{2\sigma^2}(y-X\beta)'(y-X\beta)\right]$$

$$\log L = -\frac{n}{2} \log(\sigma^2) - \frac{1}{2\sigma^2} (y-X\beta)'(y-X\beta)$$

we want to maximize $\log L$, s.t. $R\beta = q$.

That is, minimize $(y-X\beta)'(y-X\beta)$, s.t. $R\beta = q$.

$\Rightarrow$ Restricted Least Squares estimator of β .

$$(\hat{\sigma}^2 = (y-X\hat{\beta})'(y-X\hat{\beta})/n ; \hat{\beta} = \text{MLE} = \text{RLS})$$

Example: $y = X\beta + \varepsilon$

Errors follow multivariate Student-t dist'n:

$$p(\varepsilon | v) = \frac{C}{\sigma} [v + \varepsilon'\varepsilon/\sigma^2]^{-(n+v)/2}$$

$$q \in (\varepsilon) = 0 ; V(\varepsilon) = \left(\frac{\sigma^2 v}{v-2}\right) I_n$$

(if $v > 2$)

14

Once again, $J_+ = 1$, &

$$L(\beta, \sigma^2) = p(y | \beta, \sigma^2) = p(\varepsilon | v)$$

$$= \frac{C}{\sigma} [v + (y-X\beta)'(y-X\beta)/\sigma^2]^{-(n+v)/2}$$

$$\log L = \log(C) - \left(\frac{n+v}{2}\right) \log [v + (y-X\beta)'(y-X\beta)/\sigma^2]$$

Immediately, to maximize $\log L$, we see we

need to minimize $(y-X\beta)'(y-X\beta)$ w.r.t. β ,

so $\hat{\beta} = \text{OLS}$ in this case (just like Normal

(Of course the MLE of σ^2 still needs to be considered.)

What can we observe about properties of MLE so far?

(i) May be biased/unbiased in finite samples

(ii) May be efficient/inefficient in finite samples.

(iii) If MLE is unique, then it must be a function of a sufficient statistic.
(Geometric: $\hat{\beta} = 1/(1+B)$; $\sum y_i$ & hence $\bar{y}$ are sufficient.)

Also: MLE's are invariant to continuous transformations. So, if $g(\cdot)$ is any continuous fcn., & if $\tilde{\theta}$ is MLE for θ , then $g(\tilde{\theta})$ is MLE for $g(\theta)$.

e.g. If $\tilde{\sigma}^2$ is MLE for σ^2 , then $\sqrt{\tilde{\sigma}^2}$ is MLE for σ .

If $\tilde{\beta}$ is MLE for β (kx1), then $(\tilde{\beta}_1/\tilde{\beta}_2)$ is MLE for (β_1/β_2) .

Clearly, this result can save a lot of work!

(Note: Some texts say that $g(\cdot)$ must be

a 1-1 mapping, this is not needed for invariance.)

Geometric - $\tilde{\rho} = \frac{1}{1+\tilde{y}} = \text{MLE for } \rho$.

$E(Y_i) = (1-\rho)/\rho$, so MLE for $E(Y_i)$ is $(1-\tilde{\rho})/\tilde{\rho} = \tilde{y}$. More than just a "natural" estimator for $E(Y_i)$.

Now, let's consider a more complex example that illustrates the need to take care with respect to the Jacobian, & illustrates the MUV error case when Ω is unknown.

Example:

$$y_t = x_t' \beta + u_t$$

$$u_t = \rho u_{t-1} + \varepsilon_t \quad ; t=1, \dots, n$$

$$\varepsilon_t \sim \text{i.i.d. } N(0, \sigma_\varepsilon^2) \quad ; |\rho| < 1$$

$$\text{So: } E(y) = 0 \quad ; \quad V(y) = \sigma_\varepsilon^2 \Omega(\rho)$$

where

$$\Omega(\rho) = \frac{1}{1-\rho^2} \begin{bmatrix} 1 & \rho & \rho^2 & \dots & \rho^{n-1} \\ \rho & 1 & \rho & \dots & \rho^{n-2} \\ \rho^2 & \rho & 1 & \dots & \rho^{n-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho^{n-1} & \rho^{n-2} & \rho^{n-3} & \dots & 1 \end{bmatrix}$$

What do we usually do to estimate this model

Note that

$$\Omega^{-1/2} = \begin{bmatrix} \sqrt{1-\rho^2} & 0 & \dots & \dots & 0 \\ -\rho & 1 & \dots & \dots & 0 \\ 0 & -\rho & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & \dots & \dots & 0 \end{bmatrix}$$

& transform both sides of model by $\Omega^{-1/2}$.

17

$$y_t^* = \begin{bmatrix} \sqrt{1-\rho^2} y_1 \\ y_2 - \rho y_1 \\ \vdots \\ y_n - \rho y_{n-1} \end{bmatrix}; \text{ etc}$$

Then fit the model:

$$y_t^* = x_t^* \beta + \varepsilon_t$$

$$\text{or, } y_t = \rho y_{t-1} + x_t' \beta - x_{t-1}' \rho \beta + \varepsilon_t$$

$$\sqrt{1-\rho^2} y_1 = x_1' \sqrt{1-\rho^2} \beta + \varepsilon_1 \quad (t=2, \dots, n)$$

(This is effectively a NLS problem.)

Now, is this the same as MLE? Actually,

not quite!

Note: $y_t | y_{t-1} = \rho y_{t-1} + x_t' \beta - x_{t-1}' \rho \beta + \varepsilon_t$

$$\left\{ \begin{array}{l} y_1 = x_1' \beta + \varepsilon_1 / \sqrt{1-\rho^2} \\ (t=2, \dots, n) \end{array} \right.$$

18

We need to construct the Likelihood Function
We don't have the joint data density, &
the observations are not independent (they
are serially correlated). What can we do?

Recall: $p(x_2 | x_1) = p(x_2, x_1) / p(x_1)$

or: $p(x_2, x_1) = p(x_2 | x_1) p(x_1)$

Generalizing —

$$p(x_0, x_{n-1}, \dots, x_1) = p(x_n | x_{n-1}) \cdot p(x_{n-1} | x_{n-2}) \cdot \dots \cdot p(x_2 | x_1) p(x_1)$$

$$\text{Now, } p(y_t | y_{t-1}) = \frac{1}{\sigma_{\varepsilon} \sqrt{2\pi}} \exp \left[-\frac{1}{2\sigma_{\varepsilon}^2} \{ (y_t - \rho y_{t-1}) - (x_t' - \rho x_{t-1}') \beta \}^2 \right]$$

$$\sigma_{\varepsilon}^2 = (y_1 - x_1' \beta) \sqrt{1-\rho^2} \quad (t=2, \dots, n)$$

So $(\partial \varepsilon_1 / \partial y_1) = \sqrt{1-\rho^2}$

$$\sigma p(y_1) = p(\varepsilon_1) \sqrt{1-\rho^2}$$

$$\begin{aligned} \text{or, } p(y_1) &= (1-\rho^2)^{1/2} \frac{1}{\sigma_\varepsilon \sqrt{2\pi}} \exp\left[-\frac{1}{2\sigma_\varepsilon^2} \varepsilon_1^2\right] \\ &= (1-\rho^2)^{1/2} \frac{1}{\sigma_\varepsilon \sqrt{2\pi}} \exp\left[-\frac{(1-\rho^2)}{2\sigma_\varepsilon^2} (y_1 - x_1'\beta)^2\right]. \end{aligned}$$

Now,

$$\begin{aligned} L(\beta, \rho, \sigma_\varepsilon^2) &= p(y_0, \dots, y_1 | \beta, \rho, \sigma_\varepsilon^2) \\ &= p(y_n | y_{n-1}) p(y_{n-1} | y_{n-2}) \dots p(y_2 | y_1) p(y_1) \\ &= (2\pi)^{-(n-1)/2} (\sigma_\varepsilon^2)^{-(n-1)/2} \\ &\quad \cdot \exp\left\{-\frac{1}{2\sigma_\varepsilon^2} \sum_{t=2}^n [(y_t - \rho y_{t-1}) - (x_t' - \rho x_{t-1}')\beta]^2\right\} \\ &\quad \cdot (2\pi)^{-1/2} (\sigma_\varepsilon^2)^{-1/2} (1-\rho^2)^{1/2} \\ &\quad \cdot \exp\left[-\frac{(1-\rho^2)}{2\sigma_\varepsilon^2} (y_1 - x_1'\beta)^2\right] \end{aligned}$$

We need to take logs, & maximize w.r.t.

$$\beta, \sigma_\varepsilon^2 \text{ \& } \rho. \quad (\text{Beach \& MacKinnon}) \quad (1978)$$

Note:

(i) If we use NCLS (Cochrane-Orcutt)

we are omitting Jacobian term - $\log L$ is mis-specified & all parameter estimates are affected.

(ii) The first-order conditions (the likelihood equations) are highly non-linear in the parameters & cannot be solved explicitly. We need to solve them numerically. (This situation arises often)

Example: Non-linear Model

$$y_t = \beta_1 + \beta_2 x_t^{\beta_1} + \varepsilon_t \quad ; \quad \varepsilon_t \sim \text{iid } N(0, \sigma^2)$$

$$T+ = 1 \text{ \& so}$$

$$L(\beta_1, \beta_2, \sigma^2) = (2\pi\sigma^2)^{-n/2} \exp\left[-\frac{1}{2\sigma^2} \sum_t (y_t - \beta_1 - \beta_2 x_t^{\beta_1})^2\right]$$

Using $\log L$, the likelihood equations are:

$$(i) \quad \frac{1}{T} \sum_t [(y_t - \beta_1 - \beta_2 x_t^{\beta_1}) (1 + \beta_2 x_t^{\beta_1} \log x_t)] = 0$$

$$(ii) \quad \frac{1}{T} \sum_t [(y_t - \beta_1 - \beta_2 x_t^{\beta_1}) x_t^{\beta_1}] = 0$$

$$(iii) \quad -\frac{n}{2\sigma^2} + \frac{1}{2\sigma^4} \sum_t (y_t - \beta_1 - \beta_2 x_t^{\beta_1}) = 0$$

Again, no 'closed-form' solution for $\hat{\beta}_1, \hat{\beta}_2$ & $\hat{\sigma}^2$ - use numerical methods.