

Linear Model / Estimation / Least square Estimation

Model 1 :

$$Y_{ij} = \mu + \mu_i + \varepsilon_{ij} \quad \begin{matrix} 1 \leq i \leq k \\ 1 \leq j \leq n_i \end{matrix}$$

• $\mu, \mu_i \in \mathbb{R}, 1 \leq i \leq k$

• $\varepsilon_{ij} \sim \text{Normal}(0, \sigma^2)$ & independent
 $1 \leq i \leq k, 1 \leq j \leq n_i$

Population-1, Population 2, ..., Population k

Y_{11}

Y_{21}

Y_{k1}

Y_{12}

$\vdots$

Y_{k2}

$\vdots$

$\vdots$

Y_{1n_1}

Y_{2n_2}

$\vdots$

$Y_{kn_{1k}}$

$n = n_1 + n_2 + n_3 + \dots + n_k$ - is total
Sample size.

Questions :-

(i) Can we estimate μ ? If yes then
what is the estimate of μ ?

(ii) Can we estimate $\mu + \mu_i$? If yes then
what is the estimate of $\mu + \mu_i$?

(iii) Can we estimate $\mu_i - \mu_j$? If yes then what is the estimate of $\mu_i - \mu_j$?
 $1 \leq i \neq j \leq k$

Within group Variation

- Variance of group i can be estimated by

$$\left. \begin{array}{l} i\text{th Sample} \\ \text{group variance} \end{array} \right\} := \frac{1}{n_i - 1} \sum_{j=1}^{n_i} (Y_{ij} - \bar{Y}_{i.})^2 \quad \left. \begin{array}{l} \text{unbiased} \\ \text{estimate} \end{array} \right\}$$

where $\bar{Y}_{i.} = \frac{1}{n_i} \sum_{j=1}^{n_i} Y_{ij}$

- Sum of squares within:

$$SSW := \sum_{i=1}^k \sum_{j=1}^{n_i} (Y_{ij} - \bar{Y}_{i.})^2$$

Between group Variation

- Sum of squares between

$$SSB = \sum_{i=1}^k n_i (\bar{Y}_{i.} - \bar{Y}_{..})^2$$

where $\bar{Y}_{..} = \frac{1}{n} \sum_{i=1}^k \sum_{j=1}^{n_i} Y_{ij}$

Observations:

$$\begin{aligned} \text{(i)} \quad SSW + SSB &= \sum_{i=1}^k \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_{i.})^2 \\ &\quad + \sum_{i=1}^k n_i (\bar{y}_{i.} - \bar{y}_{..})^2 \\ &= \sum_{i=1}^k \left[\sum_{j=1}^{n_i} (y_{ij} - \bar{y}_{i.})^2 + (\bar{y}_{i.} - \bar{y}_{..})^2 \right] \\ &= \sum_{i=1}^k \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_{..})^2 \\ &= \text{Total sum of squares} \end{aligned}$$

(2) (Ex.) SSW is independent of SSB.

(3) (Ex.)

$$SSW \stackrel{d}{=} \chi_{k-1}^2 \sigma^2 \quad \begin{array}{l} \text{- Degrees of freedom - } k-1 \\ \text{for SSW} \end{array}$$

$$SSB \stackrel{d}{=} \chi_{n-k}^2 \sigma^2 \quad \begin{array}{l} \text{- Degrees of freedom - } n-k \\ \text{for SSB} \end{array}$$

$$TSS \stackrel{d}{=} \chi_{n-1}^2 \sigma^2 \quad \begin{array}{l} \text{- Degrees of freedom - } n-1 \\ \text{for total} \end{array}$$

Parameter vector:

$$\beta = (\mu; \mu_1, \mu_2, \dots, \mu_k) \in \mathbb{R}^{k+1}$$

Design Matrix :-

$$X_{n \times k}; \quad E[\underline{Y}] = X\beta \quad \text{where}$$

$\underline{Y}_{n \times 1}$ - data vector

$$= (Y_{11}, \dots, Y_{1n_1}; Y_{21}, \dots, Y_{2n_2}; \dots; Y_{k1}, \dots, Y_{kn_k})$$

$$\Rightarrow X = \begin{pmatrix} 1 & 1 & 0 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & 0 & 1 & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \end{pmatrix} \begin{matrix} \left. \begin{matrix} 0 \\ \vdots \\ 0 \end{matrix} \right\} n_1 \\ \left. \begin{matrix} 0 \\ \vdots \\ 0 \end{matrix} \right\} n_2 \\ \vdots \\ \left. \begin{matrix} 0 \\ \vdots \\ 0 \end{matrix} \right\} n_k \end{matrix}$$

So

Model

$$\underline{Y} = X\beta + \underline{\varepsilon}$$

$$Y_{ij} = \mu + \mu_i + \varepsilon_{ij}$$

$$1 \leq i \leq k$$

$$1 \leq j \leq n_i$$

where

$$\underline{\varepsilon} = (\varepsilon_{11}, \varepsilon_{12}, \dots, \varepsilon_{1n_1}; \varepsilon_{21}, \dots, \varepsilon_{2n_2}; \dots; \varepsilon_{k1}, \dots, \varepsilon_{kn_k})$$

Fact:-

What is the rank of X ?

$$\text{Ex: } \text{rank}(X) = k.$$

How many parameters are there?

$$\text{A: } k+1 \text{ parameters}$$

Question: Can we estimate all of them?

Answer: No!

Model 1: $Y_{ij} = \mu + \mu_i + \varepsilon_{ij} \quad 1 \leq i \leq k$
 $1 \leq j \leq n_i$

let. $\tilde{Y}_{ij} = \tilde{\mu} + \tilde{\mu}_i + \varepsilon_{ij}$

Model 2: $\tilde{\mu} = \mu + c \quad 1 \leq i \leq k$
 $\tilde{\mu}_i = \mu_i - c \quad 1 \leq j \leq n_i$

Alas: $\tilde{Y}_{ij} = Y_{ij} \quad \therefore$ Models are same!

Problem: The parameters are not identifiable.

Towards solving / understanding the Problem:

(1) Is there a parameter in the model that is identifiable?

A: Yes! E.g.:- $Y_{12} = \mu_1 - \mu_2$

Ans: $\therefore Y_{11} - Y_{21}$ is an unbiased estimate.

(as $E[Y_{11} - Y_{21}] = \mu_1 - \mu_2$)

Is it optimal? What is an optimal Estimate?

(2) Suppose we assume $\sum_{i=1}^k \mu_i = 0$

$\Rightarrow$ Yes! then all parameters are identifiable.

What is an optimal Estimate?

Note: In either (1) / (2) there are k -parameters (or their combinations) are identifiable.

Why k ? : Because of $\text{rank}(X) = k$?

Theory of linear Models:

- Suppose we have n observations $\{y_1, y_2, \dots, y_n\}$ which are realizations of random variables.
- let $(\beta_1, \beta_2, \dots, \beta_m)$ be the parameters on which the distribution of the data depends.

Consider the following model:

$$y_i = x_{i1} \beta_1 + x_{i2} \beta_2 + \dots + x_{im} \beta_m + \varepsilon_i \quad 1 \leq i \leq n$$

where $(\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n)$ are uncorrelated error random variables such that $E(\varepsilon_i) = 0 \quad 1 \leq i \leq n$
 $\text{Var}(\varepsilon_i) = \sigma^2$

$\& (x_{ij})$ are known constants.

Such a model is called a **Linear Model**.

Remark

(i) The model also can be described

$$E[y_i] = \sum_{j=1}^n x_{ij} \beta_j \quad 1 \leq i \leq n$$

$\& \{y_i\}_{1 \leq i \leq n}$ are uncorrelated random

variables with variance $= \sigma^2$.

(ii) It is called **linear** as the expected value is linear in the parameters

Example:

• One way classification

$$y_{ij} = \mu + \mu_i + \epsilon_{ij} \quad \begin{matrix} 1 \leq i \leq k \\ 1 \leq j \leq n_i \end{matrix}$$

• Nested classification

$$y_{ijk} = \mu + \mu_i + \delta_{ij} + \epsilon_{ijk} \quad \begin{matrix} 1 \leq k \leq n_{ij} \\ 1 \leq j \leq n_i \\ 1 \leq i \leq k \end{matrix}$$

• Three way classification Model

$$y_{ijk} = \mu + \alpha_i + \beta_j + \gamma_{ij} + \varepsilon_{ijk}$$

$$1 \leq k \leq n_{ij}$$

$$1 \leq j \leq J$$

$$1 \leq i \leq I$$

All models:

$$\underline{y} = (y_1, y_2, \dots, y_n)^T$$

$$\underline{\beta} = (\beta_1, \beta_2, \dots, \beta_m)^T$$

$$X = (x_{ij})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq m}}$$

$$\underline{\varepsilon} = (\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n)^T$$

Then

$$\underbrace{\underline{y}}_{n \times 1} = \underbrace{X}_{n \times m} \underbrace{\underline{\beta}}_{m \times 1} + \underbrace{\underline{\varepsilon}}_{n \times 1}$$

Σ - variance-covariance matrix $\sigma^2 I_{n \times n}$.

Model: $(\underline{y}, X\underline{\beta}, \sigma^2 I_{n \times n})$

Observation

$E[\underline{y}]$

Variance-covariance matrix of the data.

The Problem of Estimation of the Parameters

Suppose $(Y, X\beta, \sigma^2 I)$ be a linear model.

The parameters are $\beta_{n \times 1}$ and σ^2

i.e. (H) is $\mathbb{R}^m \times \mathbb{R}_+$

linear functions of the parameters

For any $p \in \mathbb{R}^m$, a linear combination

$$p_1\beta_1 + p_2\beta_2 + \dots + p_m\beta_m = p^T \beta$$

of the parameter vector $\underline{\beta}$ is called a linear parametric function.

Note: In a Linear model we will be interested only in linear parametric functions.

Definition: Suppose $(Y, X\beta, \sigma^2 I)$ be a linear model and $p \in \mathbb{R}^m$.

The linear parametric function $p^T \beta$ is

said to be estimable if

$\exists \underline{l} \in \mathbb{R}^n$ s.t.

$$E[\underline{l}^T \underline{y}] = \underline{p}^T \underline{\beta}.$$

(i.e. in other words $\exists \underline{l} \in \mathbb{R}^n$ s.t.
 $\underline{l}^T \underline{y}$ is unbiased estimate of $\underline{p}^T \underline{\beta}$)

Theorem: let $(\underline{y}, \underline{X}\underline{\beta}, \sigma^2 \underline{I}_n)$ be a linear model & $\underline{p} \in \mathbb{R}^m$.

Then $\underline{p}^T \underline{\beta}$ is estimable $\Leftrightarrow \underline{p} \in \mathcal{C}(\underline{X}^T)$

(where $\mathcal{C}(\underline{X}^T) = \{ \underline{z} : \exists \underline{y} \quad \underline{X}^T \underline{y} = \underline{z} \}$)

Proof:

$\underline{p}^T \underline{\beta}$ is estimable

$$\Leftrightarrow \exists \underline{l} \in \mathbb{R}^n \text{ s.t. } E[\underline{l}^T \underline{y}] = \underline{p}^T \underline{\beta} \quad \forall \underline{\beta} \in \mathbb{R}^m$$

$$\Leftrightarrow \exists \underline{l} \in \mathbb{R}^n \text{ s.t. } \underline{l}^T \underline{X} \underline{\beta} = \underline{p}^T \underline{\beta} \quad \forall \underline{\beta} \in \mathbb{R}^m$$

(linearity of expectation)

$$\Leftrightarrow \exists \underline{l} \in \mathbb{R}^n \quad \underline{X}^T \underline{l} = \underline{p}$$

(Ex: linear Algebra)

$$\Leftrightarrow \underline{p} \in \mathcal{C}(\underline{X}^T) \quad \square$$