

Supplemental Notes

EE503 Week 14

Dr. Franzke

HW #13

① Handout - Least squares

② Leon-Garcia 6.71-6.72

Topics: Markov Chains and Estimation.

① Mean-Square Estimation (MSE)

$$\hat{\theta}_n^{ms} = E[\theta | \mathcal{X}_n]$$

for any $g(\mathcal{X}_n)$

Orthogonality Principle: $E_{\theta, \mathcal{X}}[(\theta - \hat{\theta}_n^{ms}) g^T(\mathcal{X}_n)] = 0$

② Simple Least Squares: $\hat{y} = \hat{a}_n + \hat{b}_n \cdot x$

$$\therefore \hat{a}^{LS} = \bar{y}_n - \hat{b}^{LS} \bar{x}_n$$

$$\hat{b}^{LS} = s_{xy} / s_{xx}^2$$

③ Ordinary Least Squares (OLS): $Y = X\beta + \varepsilon$.

$$\therefore \hat{\beta}^{LS} = (X^T X)^{-1} X^T Y \sim N(\beta, \sigma^2 (X^T X)^{-1})$$

Optimal Estimators

Recall: ① Maximum Likelihood (ML)

$$\hat{\theta}_n^{ML} = \arg \max_{\theta} g(x|\theta)$$

- need closed form pdf g
- Consistency: $\hat{\theta}_n^{ML} \xrightarrow{P} \theta$ (smoothness condition)
- Asymptotic normality: $\sqrt{n}(\hat{\theta}_n^{ML} - \theta) \xrightarrow{d} W \sim N(0, \frac{1}{I(\theta)})$
for Fisher Information I (Cramer-Rao)
- Invariance principle: $\widehat{g(\theta)}^{ML} = g(\hat{\theta}^{ML})$
- Applications:
 - EM algorithm for missing data
 - Viterbi algorithm
(dynamic programming on "hidden" Markov chain)

② Maximum A Posteriori (MAP)

$$\hat{\theta}_n^{MAP} = \arg \max_{\theta} f(\theta|x) \quad \text{if } \theta \sim h(\theta) \text{ and } f(\theta|x) \propto h(\theta) \cdot g(x|\theta)$$

- MAP = MLE if $h(\theta) \sim \text{uniform}$. prior likelihood
- Bayes minimization of squared error loss
 $\rightarrow E[\theta|x]$ is optimal estimator, (MSE result)
- Gibbs Sampler (MCMC)
- Subjectivity: who gives prior $h(\theta)$?

③ Mean square estimation (MSE)

$$\hat{\theta}_n^{MS} = E[\theta|x] \quad \text{for random } \theta \text{ (in general)}$$

focus of these notes.

Data set $\mathcal{X}_n = \{X_1, X_2, \dots, X_n\}$ for random vectors X_k

Estimator $\hat{\Theta}_n = g(\mathcal{X}_n)$
 $= g(X_1, \dots, X_n) \quad \therefore \hat{\Theta}_n \text{ is a statistic.}$
 $k \times 1$

Random Parameter $\Theta \sim f(\Theta_1, \dots, \Theta_k)$
 $k \times 1$
generalizes deterministic case

Assume: Data $\mathcal{X}_n$ depends on Θ

$\therefore$ there exists a joint pdf: $f(\Theta, \mathcal{X}_n) = f(\Theta, X_1, \dots, X_n)$

★ Thrm: $\hat{\Theta}_n^{ms} = E[\Theta | \mathcal{X}_n]$ if $\mathcal{X}_n = \{X_1, \dots, X_n\}$ and Θ

obey joint pdf $f(\Theta, \mathcal{X}_n)$

Prf: $0 \leq \text{MSE}[\hat{\Theta}_n^{ms}]$
 1×1

$$= E_{\Theta \mathcal{X}} [(\Theta - \hat{\Theta}_n^{ms})^T (\Theta - \hat{\Theta}_n^{ms})]$$

$1 \times k$ $k \times 1$

where $\Theta \mathcal{X} \longleftrightarrow f(\Theta, \mathcal{X}_n) = f(\Theta_1, \dots, \Theta_n, X_1, \dots, X_n)$

total exp.

$$= E_{\mathcal{X}} [E[(\Theta - \hat{\Theta}_n^{ms})^T (\Theta - \hat{\Theta}_n^{ms}) | \mathcal{X}_n]]$$

Linear

$$= E_{\mathcal{X}} [E[\Theta^T \Theta | \mathcal{X}_n] + E[\hat{\Theta}_n^{ms T} \hat{\Theta}_n^{ms} | \mathcal{X}_n] \\ - E[\Theta^T \hat{\Theta}_n^{ms} | \mathcal{X}_n] - E[\hat{\Theta}_n^{ms} \Theta | \mathcal{X}_n]]$$

$$= E_x \left[E[\theta^T \theta | x_n] + \hat{\theta}_n^{msT} \hat{\theta}_n^{ms} - E[\theta^T | x_n] \cdot \hat{\theta}_n^{ms} - \hat{\theta}_n^{msT} E[\theta | x_n] \right].$$

since $\hat{\theta}_n^{ms} = g(x_n)$ and $E[g(x)z | x] = g(x) \cdot E[z | x]$

comp. the square = $E_x \left[E[\theta^T \theta | x_n] - E[\theta^T | x_n] \cdot E[\theta | x_n] + (\hat{\theta}_n^{ms} - E[\theta | x_n])^T (\hat{\theta}_n^{ms} - E[\theta | x_n]) \right]$

since $(\hat{\theta}_n^{ms} - E[\theta | x_n])^T (\hat{\theta}_n^{ms} - E[\theta | x_n])$
 $= \hat{\theta}_n^{msT} \hat{\theta}_n^{ms} - \hat{\theta}_n^{msT} E[\theta | x_n] - E[\theta^T | x_n] \hat{\theta}_n^{ms} + E[\theta^T | x_n] \cdot E[\theta | x_n]$

linear
 $= E_x \left[E[\theta^T \theta | x_n] - E[\theta^T | x_n] E[\theta | x_n] \right] + E_x \left[(\hat{\theta}_n^{ms} - E[\theta | x_n])^T (\hat{\theta}_n^{ms} - E[\theta | x_n]) \right]$
 $\uparrow$ only this term depends on $\hat{\theta}_n^{ms}$

$\therefore \hat{\theta}_n^{ms}$ minimizes $MSE[\hat{\theta}_n^{ms}] \geq 0$ if it

minimizes $\underbrace{E_x \left[(\hat{\theta}_n^{ms} - E[\theta | x_n])^T (\hat{\theta}_n^{ms} - E[\theta | x_n]) \right]}_{\geq 0}.$

Latter occurs iff $\hat{\theta}_n^{ms} = E[\theta | x_n].$

QED

Same result if minimize $E_{\theta x} \left[(\theta - \hat{\theta}_n^{ms})^T W (\theta - \hat{\theta}_n^{ms}) \right].$

weight matrix $W = W^T$ is positive semi-definite

Unbiased: $E_x[\hat{\theta}_n^{ms}] \stackrel{\text{thm}}{=} E_x[E[\theta|\chi_n]] \stackrel{\text{total exp.}}{=} E_\theta[\theta] \quad \forall n.$

★ Thrm: (Orthogonality Principle) MSE error is orthogonal to data (functions of data): $\varepsilon^{ms} \perp f(x).$

$$E_{\theta x}[(\theta - \hat{\theta}_n^{ms}) g^T(\chi_n)] = 0$$

for ANY function $g(\chi_n) = g(x_1, \dots, x_n)$

- applications of orthogonal projections.

- Wiener-filter (Wiener-Hopf equation) EE562
- Yule-Walker equations EE583
- Kalman-filter (linear Gauss-Markov model)

Prf:

$$E_{\theta x}[(\theta - \hat{\theta}_n^{ms}) g^T(\chi_n)]$$

$$\stackrel{\text{total exp.}}{=} E_x[E[(\theta - \hat{\theta}_n^{ms}) g^T(\chi_n) | \chi_n]]$$

$$= E_x[E[\theta - \hat{\theta}_n^{ms} | \chi_n] \cdot g^T(\chi_n)].$$

$$\text{since } E[g(x)Y|x] = g(x) \cdot E[Y|x]$$

MSE thm.

$$= E_x[E[\theta - E[\theta|\chi_n] | \chi_n] g^T(\chi_n)]$$

linear

$$= E_x[(E[\theta|\chi_n] - E[\theta|\chi_n] \cdot E[1|\chi_n]) g^T(\chi_n)].$$

$$\text{since } E[\theta|\chi_n] = \phi(\chi_n) \text{ and}$$

$$E[\phi(\chi_n)|\chi_n] = \phi(\chi_n) \cdot E[1|\chi_n]$$

$$\begin{aligned}
 &= E_x \left[\left(\underbrace{E[\theta | x_n] - E[\theta | x_n]}_{=0} \right) g^T(x_n) \right] \\
 &= E_x \left[0 \cdot g^T(x_n) \right] \\
 &= 0
 \end{aligned}$$

QED

- Hard to find closed form $E[\theta | x_n]$ in general
- $\hat{\theta}_n^{ms}$ is not robust to outliers or impulsive data
(unlike conditional median)

Special (but common) case: Jointly-Gaussian X and Y

$$E[X|Y] = \underbrace{\mu_x}_{n \times 1} + \underbrace{K_{xy}}_{n \times p} \underbrace{K_{yy}^{-1}}_{p \times p} (\underbrace{Y - \mu_y}_{p \times 1})$$

$$\underbrace{K_{x|y}}_{n \times n} = \underbrace{K_{xx}}_{n \times n} - \underbrace{K_{xy}}_{n \times p} \underbrace{K_{yy}^{-1}}_{p \times p} \underbrace{K_{yx}}_{p \times n}$$

Warning: Must TEST for normality (EE517)

Ex: Bayesian conjugacy for finding $E[\theta | x]$.
(proved in Week 5: "beta is conjugate to binomial")

Thrm: If $\theta \sim h(\theta) = \text{Beta}(\alpha, \beta)$ ↙ Bernoulli trials
 $g(x|\theta) = \text{binomial}(n, x, \theta)$

then $f(\theta|x) = \text{Beta}(\alpha+x, \beta+n-x)$

where θ estimates "success probability" p after n -trials and x successes
($\therefore n-x$ failures)

Recall: Also proved $E[\Theta] = \frac{\alpha}{\alpha+\beta}$ if $\Theta \sim \text{Beta}(\alpha, \beta)$

$\therefore$ For squared-error loss and $f(\theta|x)$

$$\hat{\Theta}_n^{ms} = E[\Theta | X=x]$$

conjugacy $\frac{\alpha'}{\alpha'+\beta'}$

since $f(\theta|x) = \text{Beta}(\alpha', \beta')$

$$= \frac{\alpha+x}{\alpha+x+\beta+n-x}$$

since $f(\theta|x) = \text{Beta}(\alpha+x, \beta+n-x)$

$$= \frac{\alpha+x}{\alpha+\beta+n}$$

← closed form answer.

$$= \frac{\alpha}{\alpha+\beta+n} + \frac{x}{\alpha+\beta+n}$$

$$= \frac{\alpha}{\alpha+\beta+n} \cdot \frac{\alpha+\beta}{\alpha+\beta} + \frac{x}{\alpha+\beta+n} \cdot \frac{n}{n}$$

$$= \underbrace{\frac{\alpha+\beta}{\alpha+\beta+n}}_{\lambda_n} \cdot \frac{\alpha}{\alpha+\beta} + \underbrace{\frac{n}{\alpha+\beta+n}}_{1-\lambda_n} \cdot \frac{x}{n}$$

$$= \lambda_n \cdot \underbrace{\frac{\alpha}{\alpha+\beta}}_{E[\Theta]} + (1-\lambda_n) \underbrace{\frac{x}{n}}_{\bar{X}_n = \hat{\Theta}_n^{ml}}$$

for convex coefficients
 $0 \leq \lambda_n \leq 1$

$$= \lambda_n \cdot E[\Theta] + (1-\lambda_n) \bar{X}_n$$

$$\longrightarrow \bar{X}_n \text{ as } n \longrightarrow \infty$$

because $\lambda_n = \frac{\alpha+\beta}{\alpha+\beta+n} \downarrow 0$
 $\therefore 1-\lambda_n \longrightarrow 1$.

$\therefore$ Effect $E[\Theta]$ of subjective prior "washes out" for large n .

$\therefore \hat{\Theta}_n^{ms}$ is asymptotically unbiased for Θ . since

$$E[\bar{X}_n] = \frac{E[X]}{n} = \frac{n p}{n} = p \text{ for binomial } n$$

$$\begin{aligned}
 - \text{Var}[E(\theta|X)] &= \text{Var}\left[\frac{\alpha + X}{\alpha + \beta + n}\right] \\
 &= \frac{1}{(\alpha + \beta + n)^2} \text{Var}[X] \\
 &= \frac{npq}{(\alpha + \beta + n)^2} \quad \text{since } X \sim \text{binomial} \\
 &\downarrow 0 \quad \text{as } n \uparrow \infty
 \end{aligned}$$

$$\therefore E[\theta|X] \xrightarrow{m} p$$

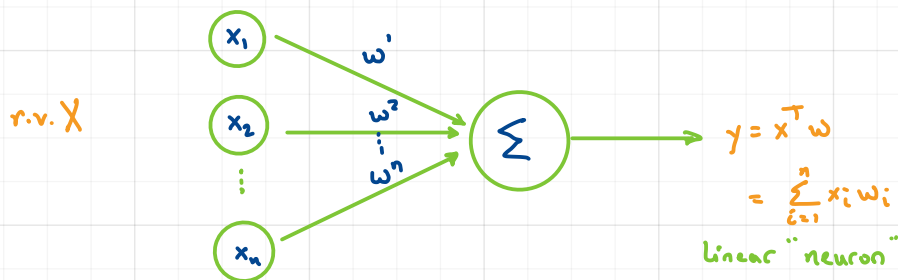
$$\therefore \text{Consistent: } E[\theta|X] \xrightarrow{\text{plim}} p \quad \text{by } m \rightarrow p$$

Adaptive MS Estimation (LMS)

EE583

"Linear combiner" $\rightarrow$ adaptive noise suppression
 $\rightarrow$ adaptive equalizer and modems

Weights: $w \in \mathbb{R}^n$



desired output $d_j \in \mathbb{R}$

Learn MS optimal w^* from $(x_1, d_1), \dots, (x_m, d_m)$

r.v. $\mathcal{E}_k = d_k - y_k = d_k - x_k^T w$

$$\therefore \text{MSE} = J_k(w) = E[\mathcal{E}_k^2]$$

$\nwarrow$ unknown pdf $p(x)$

$$\therefore \nabla_w J = 0: w^* = R^{-1}P$$

$$= (E[X_k X_k^T])^{-1} E[d_k x_k^T]$$

unknown pdfs



LMS idea: $E[\mathcal{E}_k^2] \approx \mathcal{E}_k^2$ (Widrow-Hopf, late 1950's)
 $\nearrow$
single realization of r.v.

∴ LMS algorithm

$$w_{k+1} = w_k - \mu \nabla_w J$$

Annotations:
- $\hat{\nabla_w J}$: LMS estimate
- μ : constant learning rate
- $-2\varepsilon_k x_k$: (derivative below)

general GRADIENT DESCENT algorithm

$$= w_k + 2\mu \varepsilon_k x_k \quad \star$$

↑ arguably most widely used EE algorithm.
if WSS process
(wide-sense Stationary) EE562

Not robust to outliers or impulsive (S_αS-type) noise

Robust LMS: $\varepsilon_k \mapsto \text{sign}(\varepsilon_k)$
"signed LMS"

S_αS noise: $\begin{cases} \alpha=2 & \text{Gaussian} \\ \alpha=1 & \text{Cauchy} \end{cases}$

$\alpha \downarrow \Rightarrow \text{tail-thickness} \uparrow$
($\alpha < 2 \Rightarrow \text{infinite variance}$)

∴ LMAD: "least mean absolute deviation"

- LMS similar to CENTROID learning in clustering models (pattern recognition)
- LMS: constant learning rate μ (unlike most neural network algorithms)
 ↪ still converges if WSS signal

Fact: $w_k \xrightarrow{\text{LMS}} w^*$ if

$$0 < \mu < \frac{1}{\lambda_{\max}}$$

eigenvalues of R : $0 \leq \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_{\max}$

$$\begin{aligned}\hat{\nabla}_w J &= \nabla_w \mathcal{E}_k^2 && \text{since LMS idea } E[\mathcal{E}_k^2] = \mathcal{E}_k^2 \\ &= \left(\frac{\partial \mathcal{E}_k^2}{\partial w^1}, \dots, \frac{\partial \mathcal{E}_k^2}{\partial w^n} \right) \\ &= \left(\frac{\partial}{\partial w^1} (d_k - y_k)^2, \dots, \frac{\partial}{\partial w^n} (d_k - y_k)^2 \right) \\ &= \left(-2 \underbrace{(d_k - y_k)}_{\mathcal{E}_k} \frac{\partial y_k}{\partial w^1}, \dots, -2 \underbrace{(d_k - y_k)}_{\mathcal{E}_k} \frac{\partial y_k}{\partial w^n} \right) \\ &= -2 \mathcal{E}_k \left(\frac{\partial y_k}{\partial w^1}, \dots, \frac{\partial y_k}{\partial w^n} \right) \\ &= -2 \mathcal{E}_k (x_k^1, \dots, x_k^n) && \text{since } y_k = x_k^T w = \sum_{i=1}^n x_k^i w_i \\ &= -2 \mathcal{E}_k x_k\end{aligned}$$

QED.

$$\therefore w_{k+1} = w_k + 2\mu \mathcal{E}_k x_k \quad (\text{LMS})$$

Simple Linear Regression (OLS)

$$y_k = f(x_k)$$

$$= a + b x_k + \epsilon_k$$

constant noise
 modelling error

$$\triangleq \beta_0 + \beta_1 x_k + \epsilon_k.$$

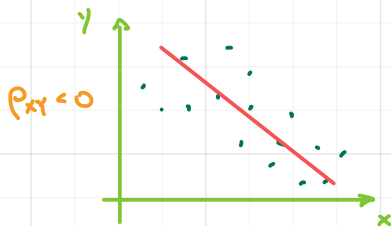
y_k is a (normal) r.v. because iid $\epsilon_k \sim N(0, \sigma^2)$

"TRUE" model: $y = a + bx$

Result: ★ $a_n^{LS} = \bar{y}_n - \hat{b}^{LS} \bar{x}_n$

$$= \hat{\beta}_0 \quad \text{unbiased and consistent.}$$

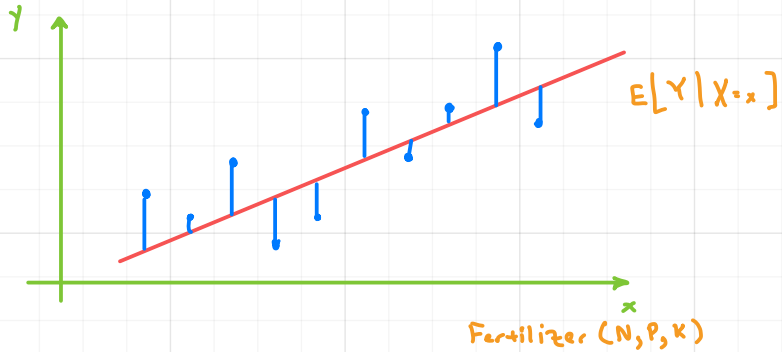
$$\star b_n^{LS} = \frac{s_{xy}}{s_x^2} = R_{xy} \cdot \frac{s_y}{s_x}$$
$$= \hat{\beta}_1 \quad \text{unbiased and consistent.}$$



Least Squares

Model: $y_k = \beta_0 + \beta_1 x_k + \varepsilon_k$

Data: $(x_1, y_1), \dots, (x_n, y_n)$; iid $\varepsilon_k \sim N(0, \sigma^2)$



Check: iid $\varepsilon_k \sim N(0, \sigma^2)$

Independence All ε_k independent

- "serial correlation" $\longrightarrow$ time series (AR(1))
Durbin-Watson test.

Homoscedasticity Constant σ_k^2

- pattern in residuals?

Normality $y_k - \hat{y}_k \sim N$

$$K_{\varepsilon\varepsilon} = \sigma^2 I = \begin{pmatrix} \sigma^4 & & 0 \\ & \ddots & \\ 0 & & \sigma^4 \end{pmatrix}$$

else look at WLS.

Find: $\hat{\Theta}_n^{ms} = \hat{\beta} = \begin{pmatrix} \hat{\beta}_0 \\ \hat{\beta}_1 \end{pmatrix}$
 2×1

Minimize J :

$$J = \text{SSE} = \sum_{k=1}^n \epsilon_k^2 = \sum_{k=1}^n (\gamma_k - \hat{\gamma}_k)^2$$

$$\stackrel{\text{LS}}{=} \sum_{k=1}^n (\gamma_k - (\hat{\beta}_0 + \hat{\beta}_1 x_k))^2$$

$$\therefore \odot = \nabla_{\beta} J = \begin{pmatrix} \partial J / \partial \beta_0 \\ \partial J / \partial \beta_1 \end{pmatrix}$$

$$= \begin{pmatrix} \sum_{k=1}^n 2(\gamma_k - (\hat{\beta}_0 + \hat{\beta}_1 x_k))(-1) \\ \sum_{k=1}^n 2(\gamma_k - (\hat{\beta}_0 + \hat{\beta}_1 x_k))(-x_k) \end{pmatrix}$$

Solve for $\hat{\beta}_0$ and $\hat{\beta}_1$ (positive definite Hessian matrix)

OLS Normal Equations (Ordinary Least Squares)

$$\hat{\beta}^{LS} = (X^T X)^{-1} X^T Y \quad \star$$

$(k+1) \times 1$

$$Y = X \beta + \varepsilon$$

$n \times 1$ $n \times (k+1)$ $(k+1) \times 1$ $n \times 1$
↙ Observation matrix
 $n \times 1$

$$y_k = \beta_0 + \beta_1 x_{1k} + \dots + \beta_k x_{kk} + \varepsilon_k$$

$$\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} = \begin{pmatrix} 1 & x_{11} & x_{12} & \dots & x_{1k} \\ 1 & x_{21} & x_{22} & \dots & x_{2k} \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & x_{n1} & x_{n2} & \dots & x_{nk} \end{pmatrix} \begin{pmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_k \end{pmatrix} + \begin{pmatrix} \varepsilon_1 \\ \varepsilon_2 \\ \vdots \\ \varepsilon_n \end{pmatrix}$$

$$Y = X \beta + \varepsilon$$

$$- \varepsilon \sim N(0, \sigma^2 I)$$

$\therefore$ White and normal and homoscedastic
(additive noise)

Recall:

$$(1) \frac{d}{dm} \left(\underbrace{b^T m}_{\substack{1 \times n \quad n \times 1 \\ 1 \times 1}} \right) = \underbrace{b}_{n \times 1} \quad \left(\frac{d}{dx} cx = c \right)$$

↙ linear form

$$(2) \text{ If } A = A^T: \frac{d}{dm} \left(\underbrace{m^T A m}_{\substack{1 \times n \quad n \times n \quad n \times 1 \\ 1 \times 1}} \right) = \underbrace{2 A m}_{\substack{n \times n \quad n \times 1 \\ n \times 1}} \quad \left(\frac{d}{dx} x^2 = 2x \right)$$

↙ quadratic form

Minimize

$$J \triangleq \text{SSE } \hat{Y} = (Y - X\beta)^T (Y - X\beta) \quad \leftarrow \text{scalar.}$$

$(Y - \hat{Y})^T (Y - \hat{Y})$

$$= Y^T Y - 2Y^T X\beta + \beta^T X^T X \beta.$$

$$\therefore \nabla_{\beta} J = 0 = -2X^T Y + 2X^T X \hat{\beta}.$$

↔ "Normal Equations"

$$(X^T X) \hat{\beta} = X^T Y \quad \star$$

$$\therefore \hat{\beta}^{LS} = (X^T X)^{-1} X^T Y = C \cdot Y \quad \leftarrow \text{linear in } Y.$$

if $(X^T X)^{-1}$ exists, iff $\det(X^T X) \neq 0$

iff $\forall$ eigenvalues $\lambda > 0$

$$X^T X = \begin{pmatrix} n & \sum_{i=1}^n x_{i1} & \cdots & \sum_{i=1}^n x_{ik} \\ \sum_{i=1}^n x_{i1} & \sum_{i=1}^n x_{i1}^2 & \cdots & \sum_{i=1}^n x_{i1} x_{ik} \\ \sum_{i=1}^n x_{i2} & \sum_{i=1}^n x_{i2} x_{i1} & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ \sum_{i=1}^n x_{ik} & \sum_{i=1}^n x_{ik} x_{i1} & \cdots & \sum_{i=1}^n x_{ik}^2 \end{pmatrix}$$

$\therefore X^T X$ is real and symmetric. $(X^T X)^T = X^T (X^T)^T = X^T X.$

$\therefore \text{for } \forall w \neq 0$

$$\begin{aligned} \underbrace{w^T (X^T X) w}_{|x|} &\stackrel{\text{assoc.}}{=} (w^T X^T) (X w) \\ &\stackrel{T}{=} (X w)^T (X w) \\ &= z^T \cdot z \quad \text{if } z = X w \\ &\geq 0 \quad \text{since } z \text{ real vector} \end{aligned}$$

$\therefore X^T X \geq 0$ Positive Semi-definite

$\longleftrightarrow$ all eigenvalues $\lambda \in \mathbb{R}^+$: $\lambda_i \geq 0$

$$\therefore \det(X^T X) = \prod_{i=1}^{n+1} \lambda_i \geq 0.$$

- Problems with numerical stability if λ_i "small" (multicollinearity)
- Pseudo-inverse always exists and is unique

Thm: $Y = X\beta + \varepsilon$ and $\varepsilon \sim N(0, \sigma^2 I)$ and

$$\hat{\beta}^{LS} = (X^T X)^{-1} X^T Y \text{ then}$$

$$\boxed{\hat{\beta}^{LS} \sim N(\beta, \sigma^2 (X^T X)^{-1})} \quad \star$$

Prf:

(1) Normality

$$\hat{\beta} = (X^T X)^{-1} X^T Y = C \cdot Y \quad - \text{linear function of } Y$$

But $Y = X\beta + \varepsilon \sim N$ because $\varepsilon \sim N(0, \sigma^2 I)$

$$\therefore \hat{\beta} \sim N$$

QED (1)

(2) Unbiasedness

$$E[\hat{\beta}] \stackrel{LS}{=} E[(X^T X)^{-1} X^T Y]$$

$$\stackrel{\text{def } Y}{=} E[(X^T X)^{-1} X^T (X\beta + \varepsilon)]$$

$$= E[(X^T X)^{-1} X^T X \beta + (X^T X)^{-1} X^T \varepsilon]$$

$$\stackrel{\text{linear}}{=} E[\underbrace{(X^T X)^{-1} (X^T X)}_{= I} \beta] + E[(X^T X)^{-1} X^T \varepsilon]$$

$$= E[\beta] + E[(X^T X)^{-1} X^T \varepsilon]$$

$$= \beta + E[(X^T X)^{-1} X^T \varepsilon] \quad \text{since } \beta \text{ not random}$$

$$= \beta + (X^T X)^{-1} X^T E[\varepsilon] \quad \text{since } X \text{ not random}$$

$$= \beta + 0 \quad \text{since } \varepsilon \sim N(0, \sigma^2 I)$$

$$= \beta$$

$\therefore \hat{\beta}$ is an unbiased estimator for β .

QED (2)

(3) Covariance (Consistency w/ Slutsky's Theorem)

$$K_{\hat{\beta}\hat{\beta}} = E[(\hat{\beta} - E[\hat{\beta}])(\hat{\beta} - E[\hat{\beta}])^T].$$

First, find K_{YY} :

$$K_{YY} = E[(Y - E[Y])(Y - E[Y])^T].$$

$$= E[(Y - X\beta)(Y - X\beta)^T]$$

since $E[Y] = E[X\beta + \varepsilon]$
 $= X\beta + E[\varepsilon] = X\beta.$

$$= E[\varepsilon\varepsilon^T]$$

since $\varepsilon = Y - X\beta.$

$$= E[(\varepsilon - 0)(\varepsilon - 0)^T].$$

by Hypo.

$$= E[(\varepsilon - E[\varepsilon])(\varepsilon - E[\varepsilon])^T]$$

since $\varepsilon \sim N(0, \sigma^2 I)$

def
 $= K_{\varepsilon\varepsilon}$

hypo.
 $= \sigma^2 I$

since $\varepsilon \sim N(0, \underline{\sigma^2 I})$

$$\therefore K_{YY} = \sigma^2 I$$

$$K_{\hat{\beta}\hat{\beta}} = E[(\hat{\beta} - E[\hat{\beta}])(\hat{\beta} - E[\hat{\beta}])^T]$$

(2)
 $= E[(\hat{\beta} - \beta)(\hat{\beta} - \beta)^T]$

since $E[\hat{\beta}] = \beta.$
 unbiased

4
 $= E[(X^T X)^{-1} X^T Y - \beta)((X^T X)^{-1} X^T Y - \beta)^T].$

"trick"
 $= E[(X^T X)^{-1} X^T Y - \underbrace{(X^T X)^{-1} X^T X \beta}_{= I})((X^T X)^{-1} X^T Y - \underbrace{(X^T X)^{-1} X^T X \beta}_{= I})^T]$

$$\begin{aligned}
 & \stackrel{\text{dist}}{=} E[(X^T X)^{-1} X^T (Y - \beta)) (X^T X)^{-1} X^T (Y - \beta))^T] \\
 &= (X^T X)^{-1} X^T \cdot E[(Y - X\beta)(Y - X\beta)^T] (X^T X)^{-1} X^T \\
 & \stackrel{T}{=} (X^T X) X^T E[(Y - X\beta)(Y - X\beta)^T] X (X^T X)^{-1} \quad \text{since } X \text{ is not random}
 \end{aligned}$$

$$\begin{aligned}
 \text{since: } (X^T X)^{-1} X^T)^T &= (X^T)^T (X^T X)^{-1})^T \\
 &= X \cdot ((X^T X)^T)^{-1} \\
 &= X \cdot (X^T (X^T)^T)^{-1} \\
 &= X \cdot (X^T X)^{-1}
 \end{aligned}$$

$$\begin{aligned}
 &= (X^T X)^{-1} X^T E[(Y - E[Y])(Y - E[Y])^T] \cdot X (X^T X)^{-1} \\
 &= (X^T X)^{-1} X^T K_{YY} X (X^T X)^{-1} \quad \text{since } E[Y] = E[X\beta + \varepsilon] = X\beta.
 \end{aligned}$$

$$\begin{aligned}
 & \stackrel{\text{lemma}}{=} (X^T X)^{-1} X^T (\sigma^2 I) X (X^T X)^{-1} \\
 &= \sigma^2 \cdot (X^T X)^{-1} \underbrace{X^T X}_{= I} (X^T X)^{-1} \\
 &= \sigma^2 \cdot (X^T X)^{-1}
 \end{aligned}$$

QED

$\therefore \hat{\beta} \xrightarrow{P} \beta$ consistent by "m-p-d" (Slutsky)

Note: asymptotic normality (CLT) if only iid ε_k and $E[\varepsilon_k] = 0$

Thrm: $\hat{\beta}^{LS} = \hat{\beta}^{ML}$

$\therefore$ invariance principle holds.

$$Y \sim N(X\beta, \sigma^2 I)$$

Prf:

$$\begin{aligned} L(\beta, \sigma^2) &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left[-\frac{1}{2} (Y-X\beta)^T \overset{\overset{I}{\cancel{X^T X}}}{K_N^{-1}} (Y-X\beta)\right] \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left[-\frac{1}{2\sigma^2} (Y-X\beta)^T I (Y-X\beta)\right] \\ &= \frac{1}{(2\pi\sigma^2)^{n/2}} \exp\left[-\frac{1}{2\sigma^2} (Y-X\beta)^T (Y-X\beta)\right]. \end{aligned}$$

$\therefore$ Maximize $L \longleftrightarrow$ minimize exponent.

$\longleftrightarrow$ minimize J

where $J = E^T E = (Y-X\beta)^T (Y-X\beta)$

$\therefore$ gives $\hat{\beta}^{LS}$

QED.

Similarly,

Thrm: $\hat{\sigma}^2_{ML} = \frac{(Y-X\hat{\beta})^T (Y-X\hat{\beta})}{n}$

Weighted Least Squares (WLS)

Positive definite $W \Rightarrow \odot$ often diagonal $W = \begin{pmatrix} w_{11} & & 0 \\ & \ddots & \\ 0 & & w_{nn} \end{pmatrix}$

$\therefore W$ has a "square root" Q : $W = Q \cdot Q^T$
- find with Cholesky Decomposition, $O(n^3)$
- Q lower triangular

W can correct over-weighted errors if errors are uncorrelated (white) but Heteroscedastic

$$\hat{\beta}^{WLS} = (X^T W X)^{-1} X^T W Y \quad (W=I)$$

$$= (X^T K_{\epsilon\epsilon}^{-1} X)^{-1} X^T K_{\epsilon\epsilon}^{-1} Y \quad \text{if } W = K_{\epsilon\epsilon}^{-1}$$

where $K_{\epsilon\epsilon}^{-1} = \begin{pmatrix} 1/\sigma_1^2 & & 0 \\ & 1/\sigma_2^2 & \\ 0 & & \ddots \\ & & & 1/\sigma_n^2 \end{pmatrix}$

since $K_{\epsilon\epsilon}$ diagonal
(uncorrelated errors)

Generalized Least Squares (GLS)

- Apply if errors heteroscedastic (non-constant variance) or correlated ($K_{\epsilon\epsilon} \neq \text{diagonal}$)

Use WLS and positive-definite $K_{\epsilon\epsilon}^{-1}$ or estimate

$$W = K_{\epsilon\epsilon}^{-1} = (QQ^T)^{-1} \quad \text{since } K_{\epsilon\epsilon} = QQ^T$$

$\therefore$ Scale with Q^{-1} to decouple

$$\text{Set } \epsilon^* = Q^{-1} \epsilon$$

$$\therefore K_{\epsilon^*\epsilon^*} = Q^{-1} K_{\epsilon\epsilon} Q = Q^{-1} (QQ^T) (Q^{-1})^T = I$$

$\therefore \epsilon^*$ white ($\therefore$ Gauss-Markov holds.)
BLUE.

$$\begin{aligned} \therefore J^{\text{GLS}} &= (Y - X\beta)^T K_{\epsilon\epsilon}^{-1} (Y - X\beta) \\ &= \left(\underbrace{Q^{-1}Y}_{Y^*} - \underbrace{(Q^{-1}X)\beta}_{X^*} \right)^T \left(Q^{-1}Y - (Q^{-1}X)\beta \right) \\ &= (Y^* - X^*\beta)^T (Y^* - X^*\beta) \end{aligned}$$

$\nwarrow$ Mahalanobis distance

$\therefore$ OLS again with scaling $Y^* = Q^{-1}Y$ and $X^* = Q^{-1}X$.