

Supplemental Notes

EE503 Week 13

Dr. Franzke

HW: (none)

① Read: Gubner, Ch. 12

Leon Garcia, Ch. 11

② Daily Derivations (Final Exam prep)

Topics: Markov Chains and Estimation.

① $m/m/1$ Queue "Birth Death" Process

"Global Balance Equations"

$$\begin{cases} \lambda p_0 = \mu p_1, & \begin{array}{cc} \text{"birth"} & \text{"death"} \\ \swarrow & \searrow \end{array} \\ (\lambda + \mu) p_n = \lambda p_{n-1} + \mu p_{n+1} & \text{if } n \geq 1 \end{cases}$$
$$\therefore p_n = (1-p) p^n \quad \text{Geometric pdf}$$

② (DTFS) Markov Chain: $\pi_e P = \pi_e$
 $\uparrow$ equilibrium pdf.

$$P = \begin{pmatrix} 1-a & a \\ b & 1-b \end{pmatrix} \Rightarrow \pi_e = \left(\frac{b}{a+b}, \frac{a}{a+b} \right)$$

$\uparrow$ the "google eigenvector"

M/M/1 Queue

- Queue with 1 server & infinite buffer

$$\therefore \text{State space} = \{0, 1, 2, 3, \dots\}$$

- First "M": Poisson arrivals (customers)

with rate $\lambda > 0$ (Poisson process)

- i.i.d. exponential inter-arrival times

(renewal process)

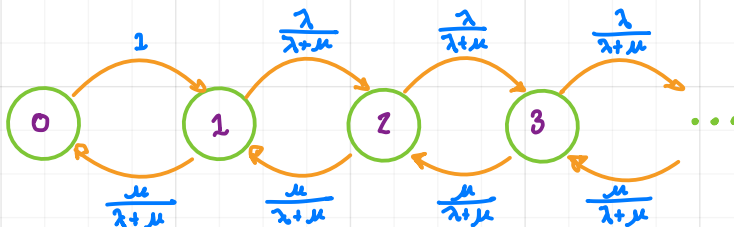
- Second "M": Exponential Service times

with rate $\mu > 0$

$$\therefore \text{pdf } f(t) = \mu e^{-\mu t}$$

$\therefore$ Continuous Markov Chain: with "jumps"
(Birth-Death process)

jumps $\longrightarrow$ EMBEDDED CHAIN.



$\therefore \lambda$ = average "birth" rate
 μ = average "death" rate

$P_n(t) = \text{Prob}(\text{system is in state } \textcircled{n} \text{ at time } t)$

$\therefore$ at equilibrium: $p_n =$ proportion of time system is in state $\textcircled{n}$

$$p_n = P_n(\infty) = \lim_{t \rightarrow \infty} p_n(t) \quad (\text{if exists})$$

Assume: $\lambda < \mu$ (else unstable)

Define: $\rho \triangleq \frac{\lambda}{\mu} \quad \therefore 0 < \rho < 1.$

★ Global Balance Equations (memorize)

$$\begin{cases} \lambda p_0 = \mu p_1 & \begin{array}{l} \text{"birth"} \\ \downarrow \end{array} \\ (\lambda + \mu) p_n = \lambda p_{n-1} + \mu p_{n+1} & \begin{array}{l} \text{"death"} \\ \downarrow \end{array} \end{cases} \quad \text{At equilibrium if } n \geq 1$$

Will show: $p_n = \left(\frac{\lambda}{\mu}\right)^n p_0 = \rho^n p_0 \quad (\text{pdf: } \sum_{n=0}^{\infty} p_n = 1)$

$$\therefore p_n = (1-\rho) \rho^n$$

Geometric pdf

$$\therefore \sum_{n=0}^{\infty} p_0 \cdot \rho^n = \frac{p_0}{1-\rho}$$

$$\therefore p_0 = 1 - \rho.$$

$$\therefore P(N \geq n) = \rho^n$$

$$\therefore E[N] = \frac{\rho}{1-\rho} = \frac{\lambda}{\mu - \lambda}$$

← expected number of customers in the system

$$V[N] = \frac{\rho}{(1-\rho)^2}$$

Derivation

Assume: - $P(1 \text{ BIRTH in } (t, t+\Delta t)) \approx \lambda \cdot \Delta t + O(\Delta t)$

- $P(\geq 2 \text{ BIRTHS in } (t+\Delta t)) \approx O(\Delta t)$

- $P(1 \text{ DEATH in } (t+\Delta t)) \approx \mu \cdot \Delta t + O(\Delta t)$

service

$$\text{for } \lim_{\Delta t \rightarrow 0} \frac{O(\Delta t)}{\Delta t} = 0$$

$$\therefore O(\Delta t) = o(\Delta t)$$

$$O(\Delta t) + O(\Delta t) = O(\Delta t)$$

Put $p_n(t) = P(\text{population} = n \text{ at time } t)$ for $n \geq 0$

$$\left(\begin{array}{l} \lambda_n = \lambda \quad \forall n, \text{ but} \\ \mu_n = 0 \quad \& \quad \mu_n = \mu \text{ for } n \geq 1 \end{array} \right)$$

Case 1: $n \geq 1$

4 birth-death (D-B) cases plus $O(\Delta t)$

$$\begin{aligned} \therefore p_n(t + \Delta t) &= p_n(t) \left(\overset{B=0 \& D=0}{1 - \lambda \Delta t} \right) \left(\overset{B=0 \& D=0}{1 - \mu \Delta t} \right) + p_n(t) \cdot \left(\overset{B=1 \& D=1}{\lambda \Delta t} \right) \left(\overset{B=0 \& D=1}{\mu \Delta t} \right) \\ &+ p_{n-1}(t) \left(\overset{B=1 \& D=0}{\lambda \Delta t} \right) \left(\overset{B=0 \& D=0}{1 - \mu \Delta t} \right) + p_{n+1}(t) \cdot \left(\overset{B=0 \& D=1}{1 - \lambda \Delta t} \right) \left(\overset{B=1 \& D=1}{\mu \Delta t} \right) + O(\Delta t) \\ &= p_n(t) \left[1 - \lambda \Delta t - \mu \Delta t + \lambda \mu (\Delta t)^2 \right] + p_n(t) \lambda \mu (\Delta t)^2 \\ &+ p_{n-1}(t) \left(\lambda \Delta t - \lambda \mu (\Delta t)^2 \right) + p_{n+1}(t) \left(\mu \Delta t - \lambda \mu (\Delta t)^2 \right) + O(\Delta t) \\ &= p_n(t) \left(1 - \lambda \Delta t - \mu \Delta t \right) + p_{n-1}(t) \lambda \Delta t + p_{n+1}(t) \mu \Delta t \\ &+ \underbrace{\left[2 p_n(t) \lambda \mu - p_{n-1}(t) \lambda \mu - p_{n+1}(t) \lambda \mu \right] (\Delta t)^2}_{= O(\Delta t) \text{ since } \lim_{\Delta t \rightarrow 0} \frac{(\Delta t)^2}{(\Delta t)} = \lim_{\Delta t \rightarrow 0} \Delta t = 0} + O(\Delta t) \end{aligned}$$

$$= p_n(t) (1 - \lambda \Delta t - \mu \Delta t) + p_{n-1}(t) \lambda \Delta t + p_{n+1}(t) \mu \Delta t + o(\Delta t)$$

since $o(\Delta t) + o(\Delta t) = o(\Delta t)$

$$= p_n(t) - p_n(t) (\lambda + \mu) \Delta t + p_{n-1}(t) \lambda \Delta t + p_{n+1}(t) \mu \Delta t + o(\Delta t)$$

$$\therefore \dot{p}_n(t) = \frac{dp_n(t)}{dt} = \lim_{\Delta t \rightarrow 0} \frac{p_n(t + \Delta t) - p_n(t)}{\Delta t}$$

$$= \lim_{\Delta t \rightarrow 0} \frac{-p_n(t) (\lambda + \mu) \cancel{\Delta t} + p_{n-1}(t) \lambda \cancel{\Delta t} + p_{n+1}(t) \mu \cancel{\Delta t}}{\cancel{\Delta t}} + \lim_{\Delta t \rightarrow 0} \underbrace{\frac{o(\Delta t)}{\Delta t}}_{=0}$$

$$= -p_n(t) (\lambda + \mu) + p_{n-1}(t) \lambda + p_{n+1}(t) \mu$$

$$\therefore \text{At equilibrium: } \dot{p}_n(t) = 0$$

$$\longleftrightarrow \underline{(\lambda + \mu) p_n = \lambda p_{n-1} + \mu p_{n+1}}$$

Case 2: $n = 0$

(i.e. $p_{-1}(t) = 0$)

3 B-D cases to (Δt) since can't have $B=0$ & $D=1$ $\therefore \mu_0 = 0$

$$\therefore p_0(t + \Delta t) = p_0(t) \overset{B=0}{\left(1 - \lambda \Delta t\right)} \overset{D=0}{\left(1 - \mu_0 \Delta t\right)} + p_0(t) \overset{B=1 \text{ \& } D=1}{\left(\lambda \Delta t\right) \left(\mu \Delta t\right)} + p_1(t) \overset{B=0 \text{ \& } D=1}{\left(1 - \lambda \Delta t\right) \left(\mu \Delta t\right)} + o(\Delta t)$$

$$= p_0(t) \left(1 - \lambda \Delta t - \mu_0 \Delta t + \lambda \mu_0 (\Delta t)^2\right) + p_0(t) \lambda \mu (\Delta t)^2$$

$$+ p_1(t) \mu \Delta t - p_1(t) \lambda \mu (\Delta t)^2 + o(\Delta t)$$

$$= p_0(t) - p_1(t) (\lambda + \mu_0) \Delta t + \underbrace{p_1(t) \mu \Delta t}_{O(\Delta t)} + \underbrace{\lambda \mu (p_0(t) - p_1(t)) (\Delta t)^2}_{+ O(\Delta t)} + O(\Delta t)$$

$$= p_0(t) - (\lambda + \overset{=0}{\mu_0}) p_0(t) \Delta t + \mu p_1(t) \Delta t + O(\Delta t)$$

$$= p_0(t) - \lambda p_0(t) \Delta t + \mu p_1(t) \Delta t + O(\Delta t)$$

$$\therefore 0 = \dot{p}_0(t) = \lim_{\Delta t \rightarrow 0} \frac{p_0(t + \Delta t) - p_0(t)}{\Delta t}$$

$$= \lim_{\Delta t \rightarrow 0} \frac{-p_0(t) \cancel{\lambda \Delta t} + p_1(t) \cancel{\mu \Delta t}}{\cancel{\Delta t}} + \lim_{\Delta t \rightarrow 0} \frac{O(\Delta t)}{\Delta t}$$

$$= -\lambda p_0(t) + \mu p_1(t)$$

$$\therefore \dot{p}_0(t) = 0 \iff \lambda p_0 = \mu p_1$$

QED - Global Balance Equations

Thm: (for Global Balance Equations)

$$p_n = \left(\frac{\lambda}{\mu}\right)^n p_0$$

Prf: (by induction on n)

Basis step: $n=1$ (& $n=2$)

$$\text{For } n=1: p_0 \lambda = p_1 \mu \iff \underline{p_1 = \left(\frac{\lambda}{\mu}\right)^1 p_0}$$

$$\text{For } n=2: p_1 (\lambda + \mu) = p_0 \lambda + p_2 \mu$$

$$\begin{aligned}\therefore p_2 &= \frac{1}{\mu} [p_1(\lambda + \mu) - p_0 \lambda] \stackrel{n=1}{=} \frac{1}{\mu} \cdot \left[\left(\frac{\lambda}{\mu} \right) p_0 (\lambda + \mu) - p_0 \lambda \right] \\ &= \frac{1}{\mu} \left(\frac{\lambda^2}{\mu} p_0 + \cancel{\lambda p_0} - \cancel{\lambda p_0} \right) = \left(\frac{\lambda}{\mu} \right)^2 \cdot p_0\end{aligned}$$

QED: Basis

Induction step:

Assume Induction Hypothesis (IH): $p_n = \left(\frac{\lambda}{\mu} \right)^n p_0$

Show: $p_{n+1} = \left(\frac{\lambda}{\mu} \right)^{n+1} p_0$

$$p_{n+1} \cdot \mu = p_n (\lambda + \mu) - p_{n-1} \lambda \quad \text{from G.B. } (\lambda + \mu) p_n = \lambda p_{n-1} + \mu p_{n+1}$$

$$\therefore p_{n+1} = \left(\frac{\lambda + \mu}{\mu} \right) p_n - \left(\frac{\lambda}{\mu} \right) p_{n-1}$$

$$\stackrel{\text{I.H.}}{=} \left(\frac{\lambda + \mu}{\mu} \right) \left(\frac{\lambda}{\mu} \right)^n p_0 - \left(\frac{\lambda}{\mu} \right) p_{n-1}$$

$$\stackrel{\text{I.H.}}{=} \left(\frac{\lambda + \mu}{\mu} \right) \left(\frac{\lambda}{\mu} \right)^n p_0 - \left(\frac{\lambda}{\mu} \right) \left(\frac{\lambda}{\mu} \right)^{n-1} p_0$$

$$= \left(\frac{\lambda + \mu}{\mu} \right) \left(\frac{\lambda}{\mu} \right)^n p_0 - \left(\frac{\lambda}{\mu} \right)^n p_0$$

$$= \left[\frac{\lambda + \mu}{\mu} - 1 \right] \left(\frac{\lambda}{\mu} \right)^n p_0$$

$$= \left[\frac{\lambda}{\mu} + \cancel{1} - \cancel{1} \right] \left(\frac{\lambda}{\mu} \right)^n p_0 = \left(\frac{\lambda}{\mu} \right)^{n+1} p_0$$

QED

$$\therefore p_n = p_0^n p_0 \quad \text{if } p_0 = \frac{\lambda}{\mu}$$

$$\therefore 1 = \sum_{n=0}^{\infty} p_n \stackrel{\text{pdf}}{=} \sum_{n=0}^{\infty} p^n p_0 \stackrel{\text{thm}}{=} p_0 \cdot \sum_{n=0}^{\infty} p^n$$

$$\stackrel{\text{G.S.}}{=} p_0 \cdot \frac{1}{1-p}$$

$$\therefore p_0 = 1-p$$

$$\therefore p_n = p^n (1-p) \quad \star$$

geometric structure

$$\therefore P(N \geq n) = \sum_{k=n}^{\infty} p_k \stackrel{\text{thm}}{=} \sum_{k=n}^{\infty} (1-p) p^k = (1-p) \sum_{k=n}^{\infty} p^k$$

$$\stackrel{\text{G.S.}}{=} (1-p) \frac{p^n}{1-p} = p^n$$

$$\therefore P(N \geq n) = p^n$$

Ex: In-n-Out Burger drive-thru (M/M/1)

Say 3 customers arrive per minute (late night): $\lambda = 3$

Find service time μ so that 95% of time queue has fewer than 10 customers

$$\therefore P(N \geq 10) = p^{10} = \left(\frac{\lambda}{\mu}\right)^{10} = \left(\frac{3}{\mu}\right)^{10}$$

$$\therefore 0.95 = P(N < 10) = 1 - P(N \geq 10) = 1 - \left(\frac{3}{\mu}\right)^{10} = \frac{19}{20}$$

$$\therefore \left(\frac{3}{\mu}\right)^{10} = \frac{1}{20} \quad \longleftrightarrow \quad \mu = 3 \sqrt[10]{20} \approx 4.048 \text{ customers per minute.}$$

Recall: $|a| < 1 \rightarrow \sum_{k=0}^{\infty} a^k \stackrel{\text{G.S.}}{=} \frac{a^0}{1-a}$

$$\therefore \sum_{n=0}^{\infty} n a^n = \frac{a}{(1-a)^2} \quad \left(= \sum_{n=1}^{\infty} n a^n \right)$$

$$\therefore \sum_{n=0}^{\infty} n^2 a^n = \frac{a+a^2}{(1-a)^3} \quad \left(= \sum_{n=1}^{\infty} n^2 a^n \right)$$

$$\therefore E_N[N] = \sum_{n=0}^{\infty} n \cdot p_n = \sum_{n=0}^{\infty} n (1-p) p^n = (1-p) \sum_{n=0}^{\infty} n p^n$$

$$\stackrel{\text{G.S.}}{=} (1-p) \cdot \frac{p}{(1-p)^2} = \boxed{\frac{p}{1-p}} = \frac{\lambda/\mu}{1-\lambda/\mu} = \frac{\lambda}{\mu-\lambda}$$

$$E_N[N^2] = \sum_{n=0}^{\infty} n^2 \cdot p_n = \sum_{n=0}^{\infty} n^2 \cdot p^n (1-p) = (1-p) \sum_{n=0}^{\infty} n^2 p^n$$

$$\stackrel{\text{G.S.}}{=} (1-p) \cdot \frac{p+p^2}{(1-p)^3} = \frac{p+p^2}{(1-p)^2}$$

$$\therefore V_N[N] = E_N[N^2] - E_N^2[N] = \frac{\cancel{p}+p^2}{(1-p)^2} - \frac{\cancel{p}^2}{(1-p)^2} = \boxed{\frac{p}{(1-p)^2}}$$

$$\therefore E_N[N] = \frac{p}{1-p} \quad \text{and} \quad V_N[N] = \frac{p}{(1-p)^2}$$

Ex: from above: $\lambda = 3$ and $\mu \approx 4 \quad \therefore p = 3/4$

$$\therefore E_N[N] = \frac{p}{1-p} = \frac{3/4}{1-3/4} = \frac{3/4}{1/4} = 3 \text{ customers}$$

$$V_N[N] = \frac{p}{(1-p)^2} = \frac{3/4}{(1-3/4)^2} = \frac{3/4}{1/16} = \frac{3 \cdot 16}{4} = 12$$

$$\therefore \sigma_N = \sqrt{V_N[N]} = \sqrt{12} = 2\sqrt{3} \text{ customers.}$$

Markov chains

Markov idea: Future = $f(\text{Present})$

↑ does not depend on the Past.

$$f(x_{n+1} | x_1, \dots, x_n) = f(x_{n+1} | x_n)$$

Multiplication Theorem:

↙ always holds

$$P\left(\bigcap_{k=1}^n A_k\right) = P(A_1) \cdot P(A_2 | A_1) \cdots P(A_n | A_1, \dots, A_{n-1})$$

Independence: $P\left(\bigcap_{k=1}^n A_k\right) = \prod_{k=1}^n P(A_k)$

Markovity lies in between

$$P\left(\bigcap_{k=1}^n A_k\right) = P(A_1) \cdot P(A_2 | A_1) \cdot P(A_3 | A_2) \cdots P(A_n | A_{n-1})$$

DTFS: Discrete Time and Finite State

State: $s_1, s_2, \dots, s_m$

↳ ("random walk")

Ex: - Frog jumping from lily pad to lily pad

- Chess-board configuration

- Google "Page rank" algorithm

(Larry page)

↳ US Patent # 6,285,999 B1.

page rank $\longleftrightarrow$ main eigenvector from
Markov-chain matrix.

Defn: $X_1, X_2, \dots$ is a Markov Chain iff

$$P(X_{k+1} = x_{k+1} \mid X_1 = x_1, \dots, X_k = x_k) = P(X_{k+1} = x_{k+1} \mid X_k = x_k) \quad \forall k$$

Ex: iid sum $S_n = \sum_{k=1}^n X_k$ is a Markov Chain
(infinite state space)

Prf: (SIT technique)

$$P(S_{n+1} = s_{n+1} \mid S_1 = s_1, \dots, S_n = s_n)$$

$$= P(X_{n+1} + S_n = s_{n+1} \mid S_1 = s_1, \dots, S_n = s_n)$$

$$\stackrel{\text{Sub}}{=} P(X_{n+1} + S_n = s_{n+1} \mid S_1 = s_1, \dots, S_n = s_n)$$

$$= P(X_{n+1} = s_{n+1} - S_n \mid S_1 = s_1, \dots, S_n = s_n)$$

$$\stackrel{\text{ind}}{=} P(X_{n+1} = s_{n+1} - S_n) \quad \text{by defn of } S_n$$

$$= P(X_{n+1} + S_n = s_{n+1})$$

$$\stackrel{\text{Sub}}{=} P(X_{n+1} + S_n = s_{n+1} \mid S_n = s_n)$$

$$= P(S_{n+1} = s_{n+1} \mid S_n = s_n)$$

$\therefore \{S_n\}$ is a Markov Chain

QED.

History: Andrei Markov - 1906 (St. Petersburg)

- Markov wanted to show that independence was not needed for the LLN and CLT

He succeeded: there are Markov LLN and CLT.

- 1913: Analyzed 20,000 letters of Pushkin's poem

found: $p(\text{vowel}) = 0.432$ ← ^{Eugeny Onegin} stationary value

$$p(\text{vowel} | \text{vowel just occurred}) = 0.128 = P(v \rightarrow v)$$

$$p(\text{vowel} | \text{consonant just occurred}) = 0.663 = P(c \rightarrow v)$$

Simplist case: DTFS

n-states: $\textcircled{1}, \textcircled{2}, \dots, \textcircled{n}$ (like pads or web pages)

$$p_{ij} = P(\textcircled{i} \rightarrow \textcircled{j}) = P(X_{k+1} = j | X_k = i)$$

↑
"transition probability"

Time-Homogeneous: $P(X_{k+1} = j | X_k = i) = p_{ij}$ for $\forall k$.
↑ $p_{ij} \geq 0$

Transition Matrix P :

$$P = \begin{pmatrix} P_{11} & P_{12} & \dots & P_{1n} \\ P_{21} & P_{22} & \dots & P_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ P_{n1} & P_{n2} & \dots & P_{nn} \end{pmatrix}$$

$n \times n$

← each row is a pdf
∴ stochastic matrix.

infinite matrix for infinite state Markov chain

So that each row sums to 1:

$$\sum_{k=1}^n P_{jk} = 1 \quad \text{for } \forall j.$$

∴ $P: S^n \rightarrow S^n$ maps probability vector π to probability vector

$\uparrow$
n-simplex

∴ P has a fixed point pdf π :

$$\pi P = \pi$$

(from Brouwer's fixed point theorem since P is linear and continuous and since S^n is convex and compact)
closed and bounded

∴ π is an eigenvector with eigenvalue $\lambda = 1$

Fact: All other eigenvalues λ' : $|\lambda'| < 1$.

↳ from Perron-Frobenius Theorem for non-negative matrices.

For any pdf u :

$$u(k) = u(0) P^k$$

for time-evolution $u(0), u(1), \dots, u(k)$

↳ from Chapman-Kolmogorov equations

Defn: P is irreducible (regular) if $\exists k > 0$: P^k is a positive

matrix: $p_{ij}^{(k)} > 0 \quad \forall i, j$

↑ cannot happen if P has a 1 on a main diagonal unless $p_{ii} = 1$. $\therefore p_{kk} = 1 \rightarrow \sim \text{irreducible}$

Defn: P is aperiodic iff the number of steps between any two states ① and ② is an integer multiple (not forced into fixed cycle).

Note: P is a periodic if no eigenvalue λ' : $\lambda' = -1$

★

Thm: ("AI") If the transition matrix P is (1) aperiodic

and (2) irreducible then P has a stationary pdf π_e :
unique since finite n

$$(1) \pi_e P = \pi_e$$

$$(2) \lim_{k \rightarrow \infty} P^k = \begin{pmatrix} -\pi_e - \\ -\pi_e - \\ \vdots \\ -\pi_e - \end{pmatrix} \quad \text{"in the long run"}$$

$$(3) \lim_{k \rightarrow \infty} u P^k = \pi_e \quad \text{for } \forall u \in S^n.$$

Basis of MCMC approximation (Gibbs sampler)

↳ identity problem with solution = π_e .

Ex: Special case: 2×2 P

$$P = \begin{matrix} & \begin{matrix} 1 & 2 \end{matrix} \\ \begin{matrix} 1 \\ 2 \end{matrix} & \begin{pmatrix} 1-a & a \\ b & 1-b \end{pmatrix} \end{matrix} \quad \text{for } a, b \in [0, 1].$$



∴ Proposition: $\pi_e = \left(\frac{b}{a+b}, \frac{a}{a+b} \right)$ if P is 2×2

Prf: $\pi_e P = \left(\frac{b}{a+b}, \frac{a}{a+b} \right) \begin{pmatrix} 1-a & a \\ b & 1-b \end{pmatrix}$

$$= \left(\frac{b - ab + ab}{a+b}, \frac{ab + a - ab}{a+b} \right)$$

$$= \left(\frac{b}{a+b}, \frac{a}{a+b} \right) = \pi_e \quad \therefore \text{fixed point.}$$

QED

Ex: Suppose each day John walks (W) or bikes (B) to

class. If John bikes one day then he never bikes the next day. If John walks one day then the next day he is equally likely to bike or walk.

Q: In the long run... how often does John bike to class?

A: $P = \begin{matrix} & \begin{matrix} B & W \end{matrix} \\ \begin{matrix} B \\ W \end{matrix} & \begin{pmatrix} 0 & 1 \\ \frac{1}{2} & \frac{1}{2} \end{pmatrix} \end{matrix} \quad \text{from above facts}$

$$= \begin{pmatrix} 1-a & a \\ b & 1-b \end{pmatrix}$$



Note: $\lambda_1 = 1$ and $\lambda_2 = 1 - a - b = -\frac{1}{2} \neq -1$ ∴ P is aperiodic (A)

$$P^2 = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ \frac{1}{4} & \frac{3}{4} \end{pmatrix} \quad \therefore \text{all positive entries}$$

$\therefore P$ is irreducible (I)

$\therefore$ Theorem applies for unique equilibrium pdf: π_e .

$$\pi_e = \left(\overset{B}{\frac{b}{a+b}}, \overset{W}{\frac{a}{a+b}} \right) = \left(\frac{\frac{1}{2}}{\frac{3}{2}}, \frac{\frac{1}{2}}{\frac{3}{2}} \right) = \left(\overset{B}{\frac{1}{3}}, \overset{W}{\frac{2}{3}} \right)$$

$\therefore$ Bike $\frac{1}{3}$ the time.

$$\begin{aligned} \text{Check: } \pi_e P &= \left(\frac{1}{3}, \frac{2}{3} \right) \begin{pmatrix} 0 & 1 \\ \frac{1}{2} & \frac{1}{2} \end{pmatrix} = \left(\frac{2}{3} \cdot \frac{1}{2}, \frac{1}{3} + \frac{2}{3} \cdot \frac{1}{2} \right) \\ &= \left(\frac{1}{3}, \frac{2}{3} \right) \\ &= \pi_e \end{aligned}$$

Extension: Hidden Markov Model (HMM)

- train with E-M algorithm.

Ex: Google Page Rank Algorithm "random surfer on N web pages"

$r(A)$ = rank of page A

$$= \frac{\alpha}{N} + (1-\alpha) \cdot \left(\frac{r(B_1)}{|B_1|} + \dots + \frac{r(B_n)}{|B_n|} \right) \quad \therefore \sum r(A) = 1$$

where $B_1, \dots, B_n$ are the "backlink" pages of A with ranks $r(B_1), \dots, r(B_n)$

- $\alpha \in [0, 1]$
- $|B_k|$ = # "forward links" of B_k
- N = total # of pages on the web

$\alpha \approx P(\text{surfer jumps randomly to any web page})$

$\alpha =$ "damping factor" $\leftarrow$ keeps algorithm from getting stuck in loops.

≈ 0.15

Put $P = \frac{\alpha}{N} \underset{N \times N}{\mathbf{1}} + (1-\alpha) \underset{N \times N}{B}$ where

$\swarrow$ matrix of all 1's

- $b_{ij} = \begin{cases} \frac{1}{n_i} & \text{if node } i \text{ points to node } j \\ 0 & \text{else} \end{cases}$

- $n_i =$ total # of forward links from i

$\therefore \exists$ unique $\underset{1 \times N}{\pi_e} : \pi_e P = \pi_e$ since P is aperiodic and irreducible.

$\therefore \pi_e = p_\infty$ is Google ranking of all N webpages

Markov chains $\{X_k\}$ obey forms of LLNs and CLTs if they are suitably "ergodic" (aperiodic and irreducible and positive recurrent) - even though the chain $X_1, X_2, \dots$ is NOT independent.

Markov Law of Large Numbers (MCLLN)

- If
- π is the stationary pdf
 - Markov chain $X_1, X_2, \dots$ is (Harris) ergodic
 - $f: X \rightarrow \mathbb{R}$ is measurable
 - $E_\pi[|f|] < \infty$

Then with probability one ("o"):

$$\bar{f}_n(x) = \frac{1}{n} \cdot \sum_{k=1}^n f(x_k) \xrightarrow{o} E_\pi[f(x)]$$

With these and further technical conditions:

MC CLT:

$$\sqrt{n}(\bar{f}_n(x) - E_\pi[f]) \xrightarrow{d} W \sim N(0, \sigma_f^2)$$

if $\sigma_f^2 < \infty$ but now σ_f^2 has the form:

$$\sigma_f^2 = V_\pi[f(x_1)] + \text{Cov}_\pi[f(x_1), f(x_k)].$$

$\therefore$ Markov was right: independence is not necessary for LLN or CLT.

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 $\hat{\Theta}_n = g(\mathcal{X}_n)$
 $= g(X_1, \dots, X_n)$ $\therefore \hat{\Theta}_n$ is a statistic.

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

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

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

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

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

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

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

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

total exp.

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

Linear

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

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

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 thrm.

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

linear

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

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

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

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

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

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

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

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

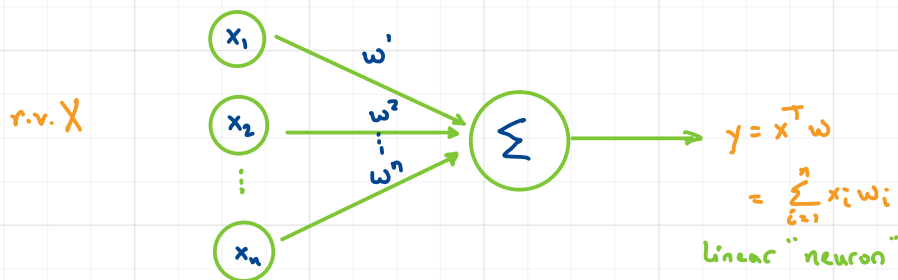
$$\therefore \text{Consistent: } E[\theta|X] \xrightarrow{\text{plim}} p \quad \text{by } m \rightarrow p$$

Adaptive MS Estimation (LMS)

EE583

"Linear combiner" $\begin{cases} \rightarrow \text{adaptive noise suppression} \\ \rightarrow \text{adaptive equalizer and modems} \end{cases}$

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:
- μ : constant learning rate
- $\nabla_w J$: LMS estimate
- $-2\varepsilon_k x_k$: (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)

$$\begin{aligned} y_k &= f(x_k) \\ &= a + b x_k + \epsilon_k \\ &\triangleq \beta_0 + \beta_1 x_k + \epsilon_k. \end{aligned}$$

constant $\downarrow$ noise $\swarrow$ modelling error $\nwarrow$

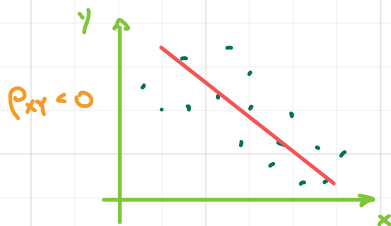
y_k is a (normal) r.v. because iid $\epsilon_k \sim N(0, \sigma^2)$

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

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

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

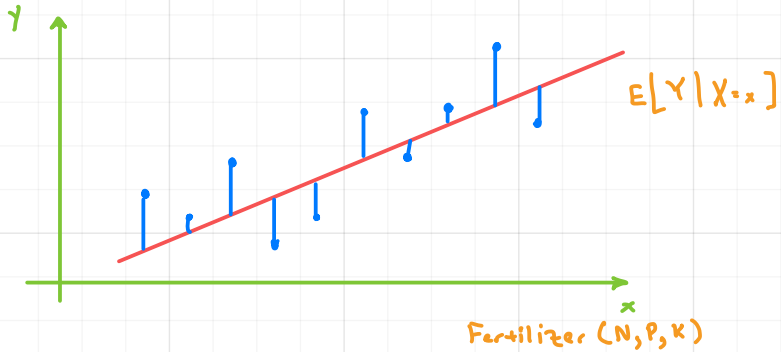
$$\begin{aligned} \star b_n^{LS} &= \frac{s_{xy}}{s_x^2} = R_{xy} \cdot \frac{s_y}{s_x} \\ &= \hat{\beta}_1 \quad \text{unbiased and consistent.} \end{aligned}$$



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

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

else look at WLS.

OLS Normal Equations (Ordinary Least Squares)

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

$(k+1) \times 1$

$$Y = X \beta + \epsilon$$

$\swarrow$ Observation matrix
 $n \times 1 \quad n \times (k+1) \quad (k+1) \times 1 \quad n \times 1$
 $\quad \quad \quad \underbrace{\quad \quad \quad}_{n \times 1}$

$$y_k = \beta_0 + \beta_1 x_1 + \dots + \beta_k x_k + \epsilon_k$$

$$\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} = \begin{pmatrix} 1 & x_{11} & x_{12} & \dots & x_{1k} \\ 1 & x_{21} & x_{22} & \dots & x_{2k} \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & x_{n1} & x_{n2} & \dots & x_{nk} \end{pmatrix} \begin{pmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_k \end{pmatrix} + \begin{pmatrix} \epsilon_1 \\ \epsilon_2 \\ \vdots \\ \epsilon_n \end{pmatrix}$$

$$Y = X \beta + \epsilon$$

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

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

Recall:

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

$\swarrow$ linear form

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

$\swarrow$ 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. $(X^T X)^T = X^T (X^T)^T = X^T X.$

$\therefore$ for $\forall w \neq 0$

$$\begin{aligned} \underbrace{w^T (X^T X) w}_{1 \times 1} &\stackrel{\text{assoc.}}{=} (w^T X^T) (X w) \\ &\stackrel{T}{=} (X w)^T (X w) \\ &= z^T \cdot z \quad \text{if } z = X w \\ &\geq 0 \quad \text{since } z \text{ real vector} \end{aligned}$$

$\therefore X^T X \geq 0$ 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 $\varepsilon \sim N(0, \sigma^2 I)$

$$\therefore \hat{\beta} \sim N \quad \text{QED (1)}$$

(2) Unbiasedness

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

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

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

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

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

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

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

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

$$= \beta$$

$\therefore \hat{\beta}$ 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

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

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

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.

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

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

