

Introduction to Mathematical Statistics Summary

Introduction

Chapter 1 - what this module is about

- ① why bother: Stats is distinguishing signal from noise. This module introduces the basic grammar you need. Data $\xrightarrow{\text{models (parameters)}}$ inference (about params)

Chapter 2 - Basic Concepts

② Probability Space: $(\Omega, \mathcal{F}, \mathbb{P})$

- Ω : $\omega \in \Omega$ are the points / elementary events
- $\mathcal{F}$: subsets of Ω forming a σ -algebra [later measure theory!]
- $\mathbb{P}$: probability measure, $\mathbb{P}(E)$ assigned to each event E

③ Random variables: A function $X: \Omega \rightarrow \mathbb{R}$ measurement, $X(\omega)$ where ω is the outcome that actually happens. $\{X \in B\} = \{\omega \in \Omega: X(\omega) \in B\}$, the law is the function $B \mapsto \mathbb{P}(X \in B)$ - its the distribution.

- Discrete: pmf, $p_X(x) = \mathbb{P}(X=x)$
- Continuous: pdf, $f_X: \mathbb{R} \rightarrow [0, \infty)$ $\mathbb{P}(X \in B) = \int_B f_X(x) dx$
- CDF captures all the info about the distribution, $F_X(x) = \mathbb{P}(X \leq x)$

You can calculate expectation, variance, moments $\mathbb{E}[X^n]$ $\leftarrow$ n th moment

- $X_1, \dots, X_n$ independent $\Rightarrow \mathbb{P}(X_1 \leq a_1, \dots, X_n \leq a_n) = \mathbb{P}(X_1 \leq a_1) \dots \mathbb{P}(X_n \leq a_n)$
- If A, B events, $\mathbb{P}(B) > 0$, then the conditional prob of A given B is

$$\mathbb{P}(A|B) = \frac{\mathbb{P}(A \cap B)}{\mathbb{P}(B)} = M_X(u)$$



"whole set is B , what is size of set A in B ?"

④ Moment Generating Functions: If $\mathbb{E}[\exp(uX)] < \infty \quad \forall u \in (-c, c)$, $c > 0$ MGF $\checkmark$ when the MGF exists, it characterises the distribution

Thm 2.1: (Uniqueness thm for MGFs). X, Y RVs w/ $M_X(u) = M_Y(u) \quad \forall u \in (-c, c)$ for some $c > 0$. Then X and Y identically distributed

⑤ Conditional Expectations: Conditional pmf of Y given X

- If $X > 0 \Rightarrow \mathbb{E}[X|Y] \geq 0$
- $\mathbb{E}[aX + bZ|Y] = a\mathbb{E}[X|Y] + b\mathbb{E}[Z|Y]$
- $\mathbb{E}[\mathbb{E}[Y|X]] = \mathbb{E}[Y]$
- $\mathbb{E}[f(Y)X|Y] = f(Y)\mathbb{E}[X|Y]$

$p_{Y|X=x}(y) = \mathbb{P}(Y=y|X=x)$
[positivity]
[linearity]
[Tower Property]
[Taking out what's known]

Statistical Models

(*) Calculations to practice (*)

Chapter 3 - A Glossary of Statistical Models

A statistical model is a parametrized family of distributions $(\mathbb{P}_\theta: \theta \in \Theta)$

⑥ Bernoulli(p): $p \in (0, 1)$

failure $\swarrow$ success
 $X = \{0, 1\}$

- pmf: $\mathbb{P}(X=1) = p = 1 - \mathbb{P}(X=0)$
- models: wins/losses at roulette, success/fail
- $\mathbb{E}[X] = p$, $\text{var}(X) = p(1-p)$, $M_X(u) = pe^u + 1-p$ ($\mathbb{E}[e^{uX}] = e^u \cdot p + e^0(1-p)$)

⑦ Binomial(n, p): $n = \#$ trials, $p \in (0, 1)$ prob. success

- pmf: $\mathbb{P}(X=k) = \binom{n}{k} p^k (1-p)^{n-k} \quad k = 0, 1, 2, \dots$

- $E[X] = np$, $Var(X) = np(1-p)$, $M_X(u) = (1 + p(e^u - 1))^n$ $\underline{pt:}$ $X = \sum_{i=1}^n X_i$ Bernoulli trials n independent & use above...
- Models: n iid coin flips

- ⑧ Geometric(p): $N \sim \text{Geometric}(p)$ is counting time till the first success
- pmf: $P(N=k) = p(1-p)^{k-1}$ for $k=1, 2, 3, \dots$
 - $P(N>n) = (1-p)^n$, $E[N] = \frac{1}{p}$, $Var(X) = \frac{1-p}{p^2}$, $M_X(u) = \frac{pe^u}{1-(1-p)e^u}$
 - Memoryless! $P(N=n+k | N>n) = \frac{p(1-p)^{n+k-1}}{(1-p)^n} = p(1-p)^{k-1} = P(X=k)$ [No effect of previous event]
 - Use: $E[X] = \sum_{k=1}^{\infty} k P(X=k) = \sum_{k=0}^{\infty} P(X>k)$

- ⑨ Neg Binomial(r, p): n iid Bernoulli, r th index $N_r - 1$ occurs.
- pmf: $P(N=k) = \binom{k-1}{r-1} p^r (1-p)^{k-r}$ $k=r, r+1, \dots$
 - $E[N_r] = \frac{r}{p}$, $Var(N_r) = r \frac{1-p}{p^2}$, $M_X(u) = \left(\frac{pe^u}{1-(1-p)e^u} \right)^r$ [independence + ⑧]
- Binomial(n, p) w/ $\lambda = np$ & large n
- Sum of r geometric
↳ # sequences length n w/
 r ones & ends w/ a 1

- ⑩ Poisson(λ): counts rare incidents, e.g. # typographical errors in long text
- pmf: $P(X=k) = \frac{\lambda^k}{k!} e^{-\lambda}$
 - $E[X] = \lambda$, $Var(X) = \lambda$, $M_X(u) = e^{\lambda(e^u - 1)}$
- $E[X] = \sum_{k=0}^{\infty} k \cdot \frac{\lambda^k}{k!} e^{-\lambda} = \dots = \lambda$
 $E[e^{uX}] = \sum_{k=0}^{\infty} e^{uk} \cdot \frac{\lambda^k}{k!} e^{-\lambda} = e^{-\lambda} \sum_{k=0}^{\infty} \frac{(\lambda e^u)^k}{k!} = e^{-\lambda} e^{\lambda e^u} = e^{\lambda(e^u - 1)}$

- ⑪ Categorical(p): Models questionnaire responses... $\mathbb{R}^k$ dimensional
- pmf: $P(X=i) = p_i$ $i \in \{1, \dots, k\}$ $\sum p_i = 1$
 - easy to fit to data, but many parameters. Occam's razor!!! Simplest explanation max, $(0, 0, \dots, 1, \dots, 0) = e_x$
- prob. simplex

- ⑫ Multinomial(n, p): lots of categorical variables together. view $X \sim \text{categorical}(p)$ as
- pmf: $P(M = \binom{n}{m_1, \dots, m_k}) = \binom{n}{m_1, \dots, m_k} \prod_{i=1}^k p_i^{m_i}$ ($\sum m_i = n$) so $X = e_i$ if $X=i$
 - see hard stuff later...
- $\binom{n}{m_1, \dots, m_k} = \frac{n!}{m_1! \dots m_k!}$

- ⑬ Uniform(a, b): model location of random objects on an interval
- $E[u] = \frac{1}{2}(a+b)$, $Var(u) = \frac{1}{12}(b-a)^2$, CDF etc... all easy to calc...
 - $E[e^{pu}] = \frac{1}{b-a} \int_a^b e^{px} dx = \frac{e^{pb} - e^{pa}}{p(b-a)}$ can be $\propto$ area/volume in high dim...
- models radioactive decays

- ⑭ Exponential(λ): continuous version of geometric. Memoryless / constant failure rate
- pdf: $f_T(x) = \lambda e^{-\lambda x}$ for $x>0$, 0 o/w.
 - $E[T] = \frac{1}{\lambda}$, $Var(T) = \frac{1}{\lambda^2}$
 - CDF: $F_T(x) = 1 - e^{-\lambda x}$ $x>0$
 - Memoryless: $P(T>u+t | T>u) = P(T>t)$
 - MGF: $E[e^{uT}] = \frac{\lambda}{\lambda - u}$ for $u<\lambda$
- calculations will need this...

Def: (The Gamma Function)

$$\Gamma(\nu) = \int_0^{\infty} x^{\nu-1} e^{-x} dx, \quad \nu > 0$$

- $\Gamma(\nu+1) = \nu \Gamma(\nu)$
- $\Gamma(n) = (n-1)!$

Think substitution to get this form

- ⑮ Gamma(ν, λ): waiting times when NOT memoryless
- pdf: $f_X(x) = \frac{\lambda^\nu}{\Gamma(\nu)} x^{\nu-1} e^{-\lambda x}$, $x>0$, 0 o/w.
 - $E[T] = \dots = \frac{\nu}{\lambda}$, $Var(T) = \frac{\nu}{\lambda^2}$
 - Gamma(1, λ) = exponential(λ) $\Rightarrow \sum_{i=1}^n$ iid exp(λ) is gamma(n, λ)
 - Gamma(ν_1, λ) + Gamma(ν_2, λ) = Gamma($\nu_1 + \nu_2, \lambda$)
 - Mode: compute $\frac{d}{dx} \log f_X(x)$ & assume why max w/ limits

will discuss more distributions when they turn up

Chapter 4 - Transformations

- ⑯ Univariate Transformation Theorem: $g: \mathbb{R} \rightarrow \mathbb{R}$ dsly diff, strictly increasing, X obs its P.V.
- CDF F_X , pdf f_X . Then $Y = g(X)$ obs its
- $$F_Y(y) = F_X(g^{-1}(y)), \quad f_Y(y) = \frac{1}{|g'(g^{-1}(y))|} f_X(g^{-1}(y))$$
- be careful of course of function

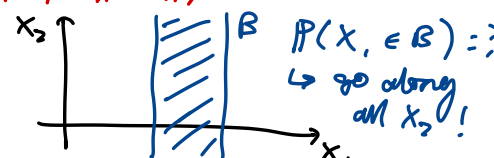
Pf: $f_Y(y) = F'_Y(y) = \frac{\partial}{\partial y} (F_X(g^{-1}(y)))$, chain rule & $\frac{\partial}{\partial y} (g(g^{-1}(y))) = 1 \dots$

- (17) Standard Normal Distribution: $Z \sim N(0,1)$, $f_Z(z) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}z^2} \quad \forall z \in \mathbb{R}$
 recall double integrals & polar coordinates trick...
 $E[Z] = 0$, $\text{var}(Z) = 1$. $F_Z(x) = \Phi(x)$

- (18) General Normal Distribution: $X \sim N(\mu, \sigma^2)$, $f_X(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{1}{2\sigma^2}(x-\mu)^2} \quad \forall x \in \mathbb{R}$
 $-$ Note $X = \mu + \sigma Z$
 $-$ CDF $F_X(x) = \Phi(\frac{x-\mu}{\sigma})$
 $-$ Linearity of expectation etc: $E[X] = \mu$, $\text{var}(X) = \sigma^2$
 $-$ MGF $M_X(t) = e^{\mu t + \frac{1}{2}\sigma^2 t^2}$ [complete the square]
 $-$ Independence: $X \sim N(\mu_1, \sigma_1^2)$, $Y \sim N(\mu_2, \sigma_2^2)$, $X \perp Y \Rightarrow X+Y \sim N(\mu_1 + \mu_2, \sigma_1^2 + \sigma_2^2)$
Pf: just use MGF, independence & uniqueness theorem for MGFs $\blacksquare$
asset prices/exchange rates/amount material

- (19) Lognormal Distribution: To model non-negative data. $Y \sim \text{LogN}(\mu, \sigma^2)$ if
 $Y = \exp(X)$ w/ $X \sim N(\mu, \sigma^2)$
 $-$ use transformation trick to get pdf/cdf.
 $-$ MGF does NOT exist! explodes... see notes...
 $- E[Y^n] = E[e^{nX}] = M_X(n) = \exp(\mu n + \frac{1}{2}\sigma^2 n^2)$

- (20) χ^2 Distribution: Central when assessing statistical procedure. $Y \sim \chi_d^2$ $Y = Z_1^2 + \dots + Z_d^2$
 where $Z_i \sim N(0,1)$ iid. $d = \#$ degrees of freedom of chi-squared distribution
 $-$ If $d=1$, $Y = Z^2$, $F_Y(y) = \Phi(\sqrt{y}) - \Phi(-\sqrt{y})$, $\checkmark$ d independent gamma
 $-$ $\chi_1^2 = \text{Gamma}(\frac{1}{2}, \frac{1}{2}) \Rightarrow \chi_d^2 = \text{Gamma}(\frac{d}{2}, \frac{1}{2})$ $\text{Gamma}(\frac{1}{2}, \frac{1}{2})$
Pf: $f_Y(y) = F'_Y(y) = \frac{1}{2\sqrt{y}} \Phi'(\sqrt{y}) + \frac{1}{2\sqrt{y}} \Phi'(-\sqrt{y}) = \frac{1}{\sqrt{y}} \Phi'(\sqrt{y}) = \frac{1}{\sqrt{2\pi}y} \exp(-\frac{y}{2})$

- (21) Joint distributions: when R.V.s interact. $X \in \mathbb{R}^n$, then joint dist is $P(X \in A)$
 $-$ Multivariate distribution function $F_{X_1, \dots, X_n}(x_1, \dots, x_n) = P(X_1 \leq x_1, \dots, X_n \leq x_n)$
 $-$ $f_X(x) = \frac{\partial^n}{\partial x_1 \dots \partial x_n} F_X(x_1, \dots, x_n)$
 $-$ Multivariate MGF: $M(a) = E[\exp(a^T X)]$
 $-$ To get marginal distribution, just integrate out all else


- (22) Covariance / Covariance matrices: How much do R.V.s depend on each other?
 $- \text{Cov}(X, Y) = E[(X - E[X])(Y - E[Y])] = E[XY] - E[X]E[Y]$. Note $X \perp Y \Rightarrow \text{Cov}(X, Y) = 0$
 $- \text{Var}(aX + bY) = a^2 \text{Var}(X) + 2ab \text{Cov}(X, Y) + b^2 \text{Var}(Y)$
Pf: $b = E[(aX + bY - E[aX + bY])^2] = \dots$
 $- -1 \leq \text{corr}(X, Y) \leq 1$ $\leftarrow$ quadratic wrt $a \Rightarrow b^2 - 4ac \leq 0$
Pf: $\text{Var}(aX + bY) \geq 0 \Rightarrow (2b \text{Cov}(X, Y))^2 - 4b^2 \text{Var}(X) \text{Var}(Y) \leq 0$
 $- X_1, \dots, X_n$ independent w/ finite 2nd moments $\Rightarrow \text{Var}(\sum a_i X_i) = \sum a_i^2 \text{Var}(X_i)$
 $-$ If $X \in \mathbb{R}^n$, $\text{Cov}(X)_{ij} = \text{Cov}(X_i, X_j)$ [diagonals are variance]

Example: Covariance matrix of $X \in \mathbb{R}^k$ multinomial (n, p) !!! EXAM HINT ???
"see example, worth being there"

$E[X] = (np_1, \dots, np_k)$, $\text{Var}(X) = np_i(1-p_i) \because X_i \sim \text{Binomial}(n, p_i)$

$\text{Cov}(X_i, X_j)$ @ $i=j \Rightarrow \text{cov} = \text{Var}(X_i) = p_i(1-p_i)$

@ $i \neq j \Rightarrow \text{cov} = E[X_i X_j] - E[X_i]E[X_j] = -p_i p_j$

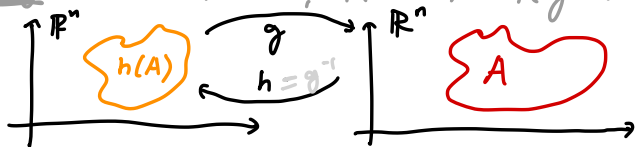
so

$$\text{Cov}(X) = \begin{pmatrix} np_1(1-p_1) & -np_1 p_2 & \dots & -np_1 p_k \\ -np_2 p_1 & np_2(1-p_2) & & \vdots \\ \vdots & & \ddots & \\ -np_k p_1 & \dots & & np_k(1-p_k) \end{pmatrix}$$

$X =$ sum of n iid categorical (p)

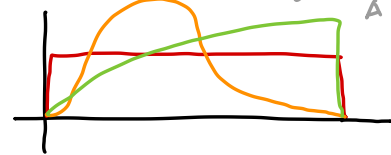
(23) Multivariate Transformation Theorem: $X \in \mathbb{R}^n$ abs dts with density $f_X(x)$, $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ is 1:1 & onto & chky diff with chky diff inverse h , then $Y = G(X)$ is abs dts with density $f_Y(y) = J_h(y) f_X(h(y))$, J_h is Jacobian $J_h = \left| \det \left(\frac{\partial h}{\partial y} \right) \right| = \left| \det \begin{pmatrix} \frac{\partial h_1}{\partial y_1} & \dots & \frac{\partial h_1}{\partial y_n} \\ \vdots & \ddots & \vdots \\ \frac{\partial h_n}{\partial y_1} & \dots & \frac{\partial h_n}{\partial y_n} \end{pmatrix} \right|$

Proof: Take $A \subset \mathbb{R}^n$, $P(Y \in A) = P(g(X) \in A) = P(X \in h(A))$



$$= \int_{h(A)} f_X(x) dx = \int_A f_X(h(y)) J_h(y) dy = \int_A f_Y(y) dy$$

change of variable
formula for
multiple integrals



(24) Beta(α, β): Very flexible, can model whatever!

- pdf: $f_X(x) = \frac{\Gamma(\alpha + \beta)}{\Gamma(\alpha)\Gamma(\beta)} x^{\alpha-1} (1-x)^{\beta-1}$, $x \in (0,1)$, $\alpha, \beta > 0$

just norm constant

- If $X \sim \text{Gamma}(\alpha, 1)$, $Y \sim \text{Gamma}(\beta, 1)$, $X \perp Y \Rightarrow \frac{X}{X+Y} \sim \text{Beta}(\alpha, \beta)$

Pf: Define $Z = X+Y$, $B = \frac{X}{X+Y}$, $(B, Z) = g(X, Y) = (\frac{X}{X+Y}, X+Y)$, use (23), find marginal

(25) Quick Linear Algebra Review:

a) Gram-Schmidt orthogonalization: Given a basis $\{x_1, \dots, x_n\}$ of $\mathbb{R}^n$, orthogonalize by

① $u_1 = \frac{x_1}{\|x_1\|_2}$ ② $u_2 = x_2 - (x_2 \cdot u_1)u_1$ ③ normalise u_2 ④ repeat.

b) $V \in \mathbb{R}^{n \times n}$ symmetric $\Rightarrow$ real eigenvalues & orthonormal basis of eigenvectors. $Vu_i = \lambda_i u_i$

c) If $U = \begin{pmatrix} u_1 & \dots & u_n \end{pmatrix}$, then $U^T U = I_n$

$$A^T A = (U L^{\frac{1}{2}} U^T)^T (U L^{\frac{1}{2}} U^T)$$

$$= U L^{\frac{1}{2}} U^T U L^{\frac{1}{2}} U^T = U L U^T = V$$

d) $V = U \text{diag}(\lambda_1, \dots, \lambda_n) U^T$

e) V positive-semidefinite $\Leftrightarrow x^T V x \geq 0 \quad \forall x \in \mathbb{R}^n$ [strict if full rank]

f) V has a sqrt $A = U L^{\frac{1}{2}} U^T$, $L = \text{diag}(\lambda_1^{\frac{1}{2}}, \dots, \lambda_n^{\frac{1}{2}})$ [avoids complex $\lambda_i^{\frac{1}{2}}$ w/ $\lambda_i < 0$]

g) If V symmetric positive semi-definite, then so is A .

(26) Multivariate Normal: $X \sim \text{MVN}(\mu, V)$ $\Leftrightarrow \exists \mu \in \mathbb{R}^n$, $A \in \mathbb{R}^{n \times m}$ s.t. $X = \mu + A Z$

$Z \in \mathbb{R}^m$ with $z_i \sim N(0,1)$ and $V = A A^T$

- $E[X_i] = E[\mu_i + \sum_j A_{ij} z_j] = \mu_i$

- $\text{Cov}(X_i, X_j) = E[(X_i - E[X_i])(X_j - E[X_j])] = E[(A z)_i (A z)_j] = E[(A z)(A z)^T]_{ij} = A A^T_{ij}$

- density: If V has full rank,

$$f_X(x) = \frac{1}{\sqrt{(2\pi)^n \det(V)}} \exp\left(-\frac{1}{2}(x-\mu)^T V^{-1}(x-\mu)\right)$$

Pf: $X = g(Z) = \mu + A Z \Rightarrow Z = h(X) = A^{-1}(X - \mu) \Rightarrow f_X(x) = J_h(x) f_Z(h(x))$

$\frac{\partial h}{\partial x} = A^{-1} \Rightarrow J_h(x) = |\det(A^{-1})| = \frac{1}{\sqrt{\det(V)}} \quad \therefore V = A A^T$, know $f_Z(z) = \frac{1}{\sqrt{(2\pi)^n}} e^{-\frac{1}{2}\|z\|_2^2}$

so $f_Z(h(x)) = \frac{1}{\sqrt{(2\pi)^n}} \exp\left(-\frac{1}{2}\|h(x)\|_2^2\right)$

$$= \frac{1}{\sqrt{(2\pi)^n}} \exp\left(-\frac{1}{2}(A^{-1}(x-\mu))^T (A^{-1}(x-\mu))\right)$$

$$(A^{-1})^T A^{-1} = (A A^T)^{-1} = V^{-1}$$

$$= \frac{1}{\sqrt{(2\pi)^n}} \exp\left(-\frac{1}{2}(x-\mu)^T V^{-1}(x-\mu)\right)$$

- MGF: $M_X(u) = E[\exp(u^T X)] = \dots = \exp(u^T \mu + \frac{1}{2} u^T V u)$

(27) Fisher's Theorem: $X_1, \dots, X_n$ independent $N(0, \sigma^2)$, $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$, $S_{XX}^2 = \sum_{i=1}^n (X_i - \bar{X}_n)^2$

Then, $\bar{X}_n$ and S_{XX}^2 are independent & $\frac{\bar{X}_n}{\sigma/\sqrt{n}} \sim N(0,1)$, $\frac{1}{\sigma^2} S_{XX}^2 \sim \chi_{n-1}^2$ independent of $Y_2, \dots, Y_n$

(wlog take $\sigma=1$)

Pf: Take $u_i = (\frac{1}{\sqrt{n}}, \dots, \frac{1}{\sqrt{n}})^T$, Gram Schmidt $\Rightarrow$ orthonormal basis $u_1, \dots, u_n$ of $\mathbb{R}^n$

Set $Y = U^T X$ with $U^T U = I_n$, $Y_i \sim \text{iid } N(0,1)$, $Y_1 = \sum_{i=1}^n \frac{1}{\sqrt{n}} X_i = \frac{\sqrt{n}}{n} \sum_{i=1}^n X_i = \sqrt{n} \bar{X}_n$

Note: $\sum_{i=1}^n Y_i^2 = \|U^T X\|^2 = X^T U U^T X = X^T X = \sum_{i=1}^n X_i^2$

so $S_{XX}^2 = \sum_{i=1}^n (X_i - \bar{X}_n)^2 = \sum_{i=1}^n X_i^2 - n \bar{X}_n^2 = \sum_{i=1}^n X_i^2 - Y_1^2 = \sum_{i=2}^n Y_i^2 \sim \chi_{n-1}^2$

independent of Y_1 $\therefore$ not Y_1 formula!

Corollary: $X_1, \dots, X_n \sim \text{iid } N(\mu, \sigma^2)$, Then $\frac{\bar{X}_n - \mu}{\sigma/\sqrt{n}} \sim N(0, 1)$, $\Rightarrow \bar{X}_n \sim N(\mu, \frac{\sigma^2}{n})$
Pf: easy! Just write out...

(28) The F-distribution: $Y \sim F_{m,n}$ (m, n are degrees of freedom) if
 $Y = \frac{X_m/m}{X_n/n}$ where $X_m \sim \chi_m^2$ and $X_n \sim \chi_n^2$

- lemma: $f_Y(y) = \frac{\Gamma(\frac{m+n}{2})}{\Gamma(\frac{m}{2})\Gamma(\frac{n}{2})} (\frac{m}{n})^{\frac{m}{2}} y^{\frac{m}{2}-1} (1 + \frac{m}{n}y)^{-\frac{1}{2}(n+m)}$

used in statistical comparison of variance

Pf: Stretched in notes...

Chapter 5 - Approximation Theorems

(29) Useful inequalities:

(a) Markov: $X \geq 0$, $P(X \geq a) \leq \frac{E[X]}{a}$

Pf: $Y = \begin{cases} 1 & \text{if } X \geq a \\ 0 & \text{o/w} \end{cases} \Rightarrow 0 \leq aY \leq X$

$0 \leq aP(X \geq a) \leq E[aY] \Rightarrow$ expectations...

(b) Chebyshev: X R.V. $E[X] = \mu$, $\text{Var}(X) = \sigma^2$, then $P(|X - \mu| \geq k) \leq \frac{\sigma^2}{k^2}$ Pf: Markov on $|X - \mu|^2$

(30) Convergence in Random Variables: $X, X_1, X_2, \dots$ R.V. defined on the same prob space (quadratic mean)

• Convergence in Quadratic mean: $X_n \xrightarrow{2\text{-mean}} X$ if $E[|X_n - X|^2] \xrightarrow{n \rightarrow \infty} 0$

Pf: $P(|X_n - X| > \varepsilon) = P(|X_n - X|^2 > \varepsilon^2) \leq \frac{E[|X_n - X|^2]}{\varepsilon^2} \xrightarrow{n \rightarrow \infty} 0$ Counter: ~~Diagram showing a sequence of points diverging from a limit.~~

• Convergence in Probability: $X_n \xrightarrow{P} X$ in prob if $\forall \varepsilon > 0, P(|X_n - X| > \varepsilon) \xrightarrow{n \rightarrow \infty} 0$

Pf: Omitted but not hard

Counter: $U \sim \text{Uniform}(0, 1)$, $X_n = \begin{cases} n & \text{if } U \leq \frac{1}{n} \\ 0 & \text{o/w} \end{cases}$
 $E[X_n] = n \times \frac{1}{n} = 1$
 $X_n \xrightarrow{P} X$

• Convergence in Distribution: $X_n \xrightarrow{\text{weakly}} X$ in distribution if $E[h(X_n)] \xrightarrow{n \rightarrow \infty} E[h(X)]$ (for every bbd cts fctn h)
 $\Leftrightarrow h(X_n) \Rightarrow h(X)$

(31) Weak Law of Large Numbers: $X_1, \dots, X_n$ iid, mean μ , var σ^2 . If $X_n = \frac{1}{n} \sum_{i=1}^n X_i$ is sample mean, then $\bar{X}_n \xrightarrow{P} \mu$ in prob.

Pf: $E[\bar{X}_n] = \frac{1}{n} \sum_{i=1}^n E[X_i] = \mu$, $\text{Var}(\bar{X}_n) = \frac{1}{n^2} \text{Var}(\sum X_i) = \frac{\text{Var}(X_1)}{n} = \frac{\sigma^2}{n}$, pick $\varepsilon > 0$,
Chebyshev: $P(|\bar{X}_n - \mu| > \varepsilon) \leq \frac{\frac{\sigma^2}{n}}{\varepsilon^2} \xrightarrow{n \rightarrow \infty} 0$

Note: Convergence of MGFs can force weak convergence: $Z_1, Z_2, \dots$ R.V. w/ MGF defined on $(-c, c)$, $c > 0$.
If $M_{Z_n}(u) \rightarrow M_Z(u)$ for $-c < u < c$. THEN, $Z_n \Rightarrow Z$ [weak convergence]

(32) Central Limit Theorem: $X_1, X_2, \dots$ independent iid, mean zero, $\text{Var}(X_i) = \sigma^2 < \infty$, have a common MGF defined on $(-c, c) \ni 0$, $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$, Then $\sqrt{n} \bar{X}_n \Rightarrow N(0, \sigma^2)$

Pf: $M(0) = E[e^{0X}] = 1$, $M'(0) = 0$, $M''(0) = \sigma^2$, plug into power series: $M(u) = 1 + \frac{1}{2}\sigma^2 u^2 + O(u^3)$
 $S_n = \sum X_i$, $M_{S_n}(u) = (M(u))^n$ [iid], $Z_n = \frac{S_n}{\sigma\sqrt{n}} = \frac{1}{\sigma\sqrt{n}} S_n \Rightarrow M_{Z_n}(u) = M_{S_n}(\frac{u}{\sigma\sqrt{n}}) = (M(\frac{u}{\sigma\sqrt{n}}))^n$
 $M_{Z_n}(u) = (1 + \frac{\sigma^2 u^2}{2(\sigma\sqrt{n})^2} + o(\frac{1}{n}))^n \xrightarrow{n \rightarrow \infty} e^{-\frac{u^2}{2}} \leftarrow \text{MGF of } N(0, 1), \text{ so write } \sqrt{n} \bar{X}_n \xrightarrow{d} N(0, \sigma^2)$

"weak convergence preserved under continuous mappings"

(33) Continuous Mapping Theorem: $X_i \in \mathbb{R}^k$ R.V. $X_n \Rightarrow X$, $g: \mathbb{R}^k \rightarrow \mathbb{R}^m$ cts. Then $g(X_n) \Rightarrow g(X)$

Pf: $X_n \Rightarrow X \Rightarrow f(X_n) \Rightarrow f(X)$ $\forall$ bbd cts $f: \mathbb{R}^k \rightarrow \mathbb{R}$. Pick $h: \mathbb{R}^m \rightarrow \mathbb{R}$ bbd cts, then $h \circ g: \mathbb{R}^k \rightarrow \mathbb{R}$ bbd cts so $(X_n \Rightarrow X) \Rightarrow h(g(X_n)) \Rightarrow h(g(X))$ so $\#$

(34) Slutsky's Thm: $X_n \Rightarrow X$, $Y_n \Rightarrow c$, then $X_n + Y_n \Rightarrow X + c$, $X_n Y_n \Rightarrow cX$, $\frac{X_n}{Y_n} \Rightarrow \frac{X}{c}$ ($c \neq 0$)
Pf: omitted

Statistical Inference

♥ of module

Chapter 6 - Likelihood

Process

- 1 Get data $\{x_1, \dots, x_n\}$
- 2 Model via $\{X_1, \dots, X_n; \theta\}$ *hypothesis tests*
- 3 estimate $\hat{\theta}$, accurate? *consistent w/ data?*
confidence intervals

35 Observed Likelihood function: Data $x_1, \dots, x_n$ modelled as observed values of R.V.s w/ joint density/pmf $f_{x_1, \dots, x_n}(x_1, \dots, x_n; \theta)$, Observed likelihood function is $L(\theta) = f_{x_1, \dots, x_n}(x_1, \dots, x_n; \theta)$

- Not pmf/density so no need to integrate to 1, function of θ , data is fixed!
(MLE)

36 Maximum Likelihood Estimate: MLE $\hat{\theta}$ is the value of θ that maximises $L(\theta)$ assuming it exists & is unique: $\hat{\theta} = \arg \max_{\theta} L(\theta) = \arg \max_{\theta} \ell(\theta)$ [same $\because$ log monotonic]

MLE (i) makes the data more probable

log-likelihood

(ii) value of θ that has most support from data

$$\ell(\theta) = \log L(\theta)$$

coincidence or an event happening

warning: correlation $\nrightarrow$ causation.

Rain CAUSES ppl to stay indoors (makes sense)

Staying indoors CAUSES rain [$\neq$ makes sense]

beware of confounding variables

37 Example MLE calculations:

CORE SKILL OF THE MODULE !!!

Assumption: INDEPENDENCE !!

$\hookrightarrow$ Q: How were populations sampled?

- 1 If you sample from internet - you self-select for ready internet trade...
- 2 Phone for political survey: NOT EVERYONE PICKS UP THE PHONE $\Rightarrow$ NOT a random sample of pop. of interest...
- 3 Snowball sampling: ask friends to ask answer $\Rightarrow$ Health/wealth Qs $\Rightarrow$ very bad!

Getting responses > good quality data
dangerous !!!

a MLE for sequence of Bernoulli R.V.s

only records Σx_i & n so good for privacy!

- Likelihood: $L(\theta) = \theta^{\Sigma x_i} (1-\theta)^{n-\Sigma x_i} \xrightarrow{\log} \ell(\theta) = \log \theta (\Sigma x_i)$

- Diff: $\frac{\partial \ell}{\partial \theta} = \frac{1}{\theta} \Sigma x_i - \frac{1}{1-\theta} (n - \Sigma x_i) + \log(1-\theta)(n - \Sigma x_i)$

- $\frac{\partial \ell}{\partial \theta} = 0 \Leftrightarrow \dots \Leftrightarrow \hat{\theta}_{MLE} = \frac{1}{n} \Sigma x_i$

- JV: $L(\theta) \rightarrow 0$ when $\theta \rightarrow 0$ or $\theta \rightarrow 1 \Rightarrow \text{MAX}$

b MLE for data modelled by binomial RV

- Model: $Y = \Sigma X_i$ w/ X_i binomial & remember n , $Y \sim \text{Binomial}(n, \theta)$

- Likelihood: $L(\theta) = P_{\theta}(Y=y) = \binom{n}{y} \theta^y (1-\theta)^{n-y} \Rightarrow \ell(\theta) = \log \binom{n}{y} + y \log \theta + (n-y) \log(1-\theta)$

- Diff: $\frac{\partial \ell}{\partial \theta} = \frac{y}{\theta} - \frac{n-y}{1-\theta}$, $\frac{\partial \ell}{\partial \theta} = 0 \Leftrightarrow \hat{\theta}_{MLE} = \frac{y}{n}$

- JV: $\theta \rightarrow 1$, $\theta \rightarrow 0$, $\ell(\theta) \rightarrow 0$, $\ell(\theta) \rightarrow 0 \Rightarrow \text{MAX}$

c MLE for data modelled by Normal RVs (I)

- Model $x_1, \dots, x_n \sim n$ iid $N(\mu, \sigma^2)$ $\theta = (\mu, \sigma) \in \mathbb{R} \times (0, \infty)$

- $L(\theta) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2\sigma^2}(x_i - \mu)^2} \Rightarrow \ell(\theta) = -\frac{n}{2} \log(2\pi\sigma^2) - \frac{\sum (x_i - \mu)^2}{2\sigma^2}$: log easier to diff

- Diff $\frac{\partial \ell}{\partial \mu} = 0$ & $\frac{\partial \ell}{\partial \sigma} = 0 \Rightarrow \frac{\partial \ell}{\partial \mu} = \frac{\sum (x_i - \mu)}{\sigma^2} \Rightarrow \mu = \frac{1}{n} \sum x_i$, $\frac{\partial \ell}{\partial \sigma} = -\frac{n}{2\sigma} - \frac{1}{2\sigma^3} \sum (x_i - \mu)^2$

- JV: Hessian has strictly negative eigenvalues / $L(\mu, \sigma^2) \rightarrow 0$ on boundary of $\mathbb{R} \times (0, \infty) \Rightarrow \text{MAX}$

- so can say that $\hat{\mu}_{MLE} = \frac{1}{n} \sum x_i = \bar{x}$, $\hat{\sigma}_{MLE}^2 = \frac{1}{n} \sum (x_i - \bar{x})^2$

$$-\frac{1}{2} \sigma^{-2} \rightarrow \frac{1}{2} \sigma^{-1}$$

d MLE for data modelled by discrete uniform RVs

no need for calculus here !!!

- Model: bag has unknown # balls n . Select k at random: data $x_1, \dots, x_k$ each #s $\wedge$

sample WITH replacement

$\hookrightarrow$ use $X_1, \dots, X_k$ as independent $\text{Uniform}(\{1, \dots, n\})$ to model this data.

- $L(n) = \prod_{i=1}^k \frac{1}{n} \mathbb{1}\{x_i \in \{1, \dots, n\}\} = \frac{1}{n^k} \mathbb{1}\{n \geq \max \{x_1, \dots, x_k\}\}$

- $L(n)$ monotone $\therefore \frac{1}{n^k} \Rightarrow L(n) \downarrow$ for $n > \max \{x_1, \dots, x_k\} \Rightarrow \hat{n}_{MLE} = \max \{x_1, \dots, x_k\}$

e MLE for simple linear regression [Normal RVs (II)]

"all models are wrong but some are useful"

calculations easier

- Model: $x_i = \text{age}$, $y_i = \text{blood pressure} \Rightarrow Y_i = \gamma + \beta X_i + \sigma \varepsilon_i$; $\varepsilon_i \sim \text{iid } N(0, 1)$

- re-parameterise: $\bar{x} = \frac{1}{n} \sum x_i$, $x_i' = x_i - \bar{x}$ [centered], $Y_i = \alpha + \beta x_i' + \sigma \varepsilon_i$ [$\alpha = \gamma + \beta \bar{x}$]

- switch: $\varepsilon_i = \frac{1}{\sigma} (y_i - \alpha - \beta x_i')$ $\Rightarrow L(\alpha, \beta, \sigma) = -\frac{n}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} \sum (y_i - (\alpha + \beta x_i'))^2$

- Diff: $\frac{\partial \ell}{\partial \alpha} = 0 \Leftrightarrow \hat{\alpha}_{MLE} = \bar{y}$, $\frac{\partial \ell}{\partial \beta} = 0$, $\frac{\partial \ell}{\partial \sigma} = 0 \Leftrightarrow \hat{\beta}_{MLE} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2}$, $\hat{\sigma}_{MLE}^2 = \frac{1}{n} \sum (y_i - (\hat{\alpha} + \hat{\beta}(x_i - \bar{x})))^2$

- JV: Can $L(\alpha, \beta, \sigma) \rightarrow \infty$ as $(\alpha, \beta, \sigma) \rightarrow \partial(\mathbb{R} \times \mathbb{R} \times (0, \infty))$?

$\hookrightarrow$ Involved, check LN. Generalises to multiple covariates. VERY important! GLMs...

f) MLE for multinomial Distribution

"Think: political opinion poll: # individuals $\rightarrow$ categories"

- Setup: Sample n individuals from pop. w/ k types. n_i of type i , want p_i , $\sum_{i=1}^k p_i = 1$
- Model: $N_1, \dots, N_k \sim \text{Multinomial}(n, p)$, $p \in \mathbb{R}^k$. $\sum p_i = 1$, $p_i \geq 0 \quad \forall i \in \{1, \dots, k\}$, $p_k = 1 - p_1 - \dots - p_{k-1}$
- Likelihood: $L(p) = \binom{n}{n_1, \dots, n_k} \prod_{i=1}^k p_i^{n_i}$, $n \in \mathbb{R}^k$, max p over parameter space (Lagrange multipliers)
- maximiser: $\ell(p) = \log(\hat{n}) + \sum_{i=1}^k n_i \log p_i$, For $i = 1, \dots, k-1$, split $\dots \Rightarrow \hat{p}_{i, \text{MLE}} = \frac{n_i}{n}$
- JV: when $n_i > 0 \quad \forall i$, $\ell(p) \rightarrow 0$ as $p_i \rightarrow 0$

major tool in modern data science

"when we raise money its AI, when we hire its machine learning, when we do the work, it's logistic regression."

blurred line between stats, ML, data science etc... often data science sounds sexy & gets funding but substance is this...

g) MLE for Logistic regression

- Scenario: $x_i = \text{BMI}$, $y_i = \begin{cases} 1 & \text{type II diabetes} \\ 0 & \text{otherwise} \end{cases}$
- Model: $Y_i \sim \text{Bernoulli}(f(x_i; \beta_0, \beta_1))$ where $f(x; \beta_0, \beta_1) = \frac{\exp(\beta_0 + \beta_1 x)}{1 + \exp(\beta_0 + \beta_1 x)}$ logistic functions!
- Attractive for interpretation: $1 - f(x; \beta_0, \beta_1) = \frac{1}{1 + \exp(\beta_0 + \beta_1 x)} \Rightarrow \log\left(\frac{P(Y_i = 1)}{P(Y_i = 0)}\right) = \beta_0 + \beta_1 x$ linear
- Likelihood: $L(\beta_0, \beta_1) = \prod_{i=1}^n f(x_i; \beta_0, \beta_1)^{y_i} (1 - f(x_i; \beta_0, \beta_1))^{1-y_i}$
- MLE: compute numerically.

Chapter 7 - Repeated Sampling Principle

How accurate are our MLE estimates?

Does NOT depend on the model parameters!!!
e.g... MLE

38) Statistic: A statistic is simply a function $t(x_1, \dots, x_n)$ of the observed data $x_1, \dots, x_n$.

39) Repeated Sampling Principle: Statistical decision procedures should be evaluated on the basis of their behaviour in hypothetical repetitions of the experiment that generated the original data $x_1, \dots, x_n$. If we observe data that $\neq$ fit model $\Rightarrow$ reject model.

note: if data $x_1, \dots, x_n$ modelled by $X_1, \dots, X_n$ w/ param θ , if $t(x_1, \dots, x_n)$ estimates model parameter θ , then $T = t(X_1, \dots, X_n)$ is an estimator for θ .

want small BUT large time & - we don't cancel...

40) Bias of an Estimator: Bias of the estimator $T = T(X_1, \dots, X_n)$ for θ is $\mathbb{E}_\theta[T] - \theta$

word both small

distance between estimator & param

41) Mean-Square-Error of an estimator: MSE of T for θ is $\mathbb{E}_\theta[(T - \theta)^2]$

scales linearly rather than quadratically

root mean squared

deviation

square so $\pm$ cancel out

42) Standard error of an estimator: S.E. $\approx \text{RMSE} \approx \sqrt{\text{MSE}}$ of T for θ is $\sqrt{\mathbb{E}_\theta[(T - \theta)^2]}$

43) Bias-variance decomposition of MSE: $\text{MSE} = \mathbb{E}_\theta[(T - \theta)^2] = \text{Var}_\theta(T) + (\mathbb{E}_\theta(T) - \theta)^2$

pf: $\mathbb{E}[(T - \theta)^2] = \mathbb{E}[(T - \mathbb{E}_\theta[T] + \mathbb{E}_\theta[T] - \theta)^2]$

$= \underbrace{\mathbb{E}[(T - \mathbb{E}_\theta[T])^2]}_{\text{variance}} + \underbrace{2\mathbb{E}_\theta[(T - \mathbb{E}_\theta[T])(\mathbb{E}_\theta[T] - \theta)]}_{\text{non-random}} + \underbrace{(\mathbb{E}_\theta[T] - \theta)^2}_{\text{bias}^2}$

44) Consistency of Estimators: A sequence of estimators $T_1, T_2, \dots$ is consistent for θ if $T_n \xrightarrow{n \rightarrow \infty} \theta$ in probability.

note: $T_n = T_n(X_1, \dots, X_n)$ estimates θ . $\text{MSE}(T_n) = \mathbb{E}_\theta[(T_n - \theta)^2] \xrightarrow{n \rightarrow \infty} 0$

convergence in 2-mean.

30) $\Rightarrow T_n \rightarrow \theta$ in probability

note: MLE could be biased but consistent.

- Data $x_1, \dots, x_n \sim \text{iid Uniform}(0, \theta)$, θ unknown

- MLE $\hat{\theta} = M_n = \max\{x_1, \dots, x_n\}$, $\ell(\theta) = \frac{1}{\theta^n}$ for $\theta > \max\{x_1, \dots, x_n\}$

- CDF of M_n : $F_{M_n}(m) = \left(\frac{m}{\theta}\right)^n$ for $0 < m < \theta$, $P(M_n \leq m) = (P(X_1 \leq m))^n$

- PDF of M_n : $f_{M_n}(m) = \frac{nm^{n-1}}{\theta^n}$ for $0 < m < \theta$ [diff]

- $\mathbb{E}_\theta[M_n] = \int_0^\theta m \cdot \frac{nm^{n-1}}{\theta^n} dm = \frac{n}{n+1} \theta \Rightarrow \text{bias} = \mathbb{E}_\theta[M_n] - \theta = -\frac{1}{n+1} \theta$

- $\mathbb{E}_\theta[M_n^2] = \dots = \frac{n\theta^2}{n+2}$, $\text{MSE} = \mathbb{E}_\theta[(M_n - \theta)^2] = \dots = \frac{2}{(n+1)(n+2)} \theta^2$

more calc shows biased but not consistent...

45) MLE General Picture: Under 'suitable regularity conditions', sequences of MLE estimators, (i) asymptotically unbiased, (ii) sequences of MLEs have MSE which are asymptotically as small as possible

Hence MLE consistent

(***)

- sequences of MLEs $\neq$ converge to boundary of param space
- # data $\rightarrow \infty$
- dim parameters θ cannot increase for sequences of MLEs

MLE must not be the 'best' estimator in terms of MSE

(46) MLE may not yield optimal MSE: idea is that intuitively natural estimators can actually be biased. Need to adjust to make them unbiased. *Saying we wanna estimate $\mu, \sigma^2 \dots$*

$MSE = \text{Variance} + \text{bias}^2$ *Pick your position...*

This is straightforward

(a) Estimation of mean using sample mean

- Take $X_1, \dots, X_n$ iid mean μ , variance σ^2 . If $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$ is sample mean
- $E[\bar{X}_n] = \frac{1}{n} \sum E[X_i] = \frac{n\mu}{n} = \mu \Rightarrow \text{Bias: } E[\bar{X}_n] - \mu = 0 \Rightarrow \text{unbiased}$
- $MSE = \text{var} + \text{bias}^2 = \text{Var}(\bar{X}_n) = \frac{1}{n^2} \sum \text{Var}(X_i) = \frac{\sigma^2}{n}$ *MSE*

MLE variance is biased sample variance is unbiased but ↑ variance

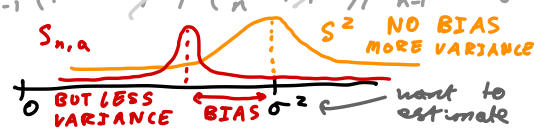
$\text{Var}(Z) = E[Z^2] - E[Z]^2 \Rightarrow E[Z^2] = \text{Var}(Z) + E[Z]^2$

(b) Estimation of variance using sample variance

- MLE is biased! $\hat{\sigma}_{MLE}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2$
- Sample variance is unbiased! $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2 = \frac{1}{n-1} \sum (X_i^2 - 2\bar{X}_n X_i + \bar{X}_n^2) = \frac{1}{n-1} (n E[X_i^2] - 2n E[\bar{X}_n^2] + n E[\bar{X}_n^2])$
- why? $E[S^2] = \frac{1}{n-1} E[\sum X_i^2 - 2\sum \bar{X}_n X_i + \sum \bar{X}_n^2] = \frac{1}{n-1} (n E[X_i^2] - 2n E[\bar{X}_n^2] + n E[\bar{X}_n^2])$
So bias = $E[S^2] - \sigma^2 = 0$
- let $S_{n,a} = \frac{1}{a} \sum_{i=1}^n (X_i - \bar{X}_n)^2$
- $\arg\min(MSE(S_{n,a})) = a^* = n+1$

$MSE(\hat{\sigma}_{MLE}^2) < MSE(S^2)$
better MSE *worse MSE*

Trade off between bias & variance



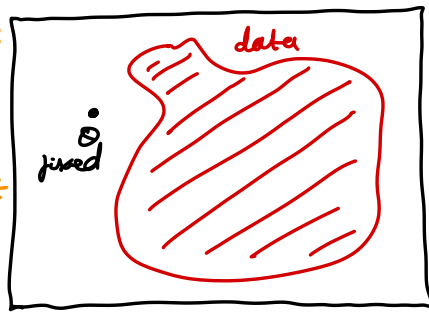
note: method of moments is an alternative to MLE. 200 years old. MLE is gold standard!

AIM: construct intervals using data $\Rightarrow$ where unknown θ live...

Chapter 8 - Confidence Intervals

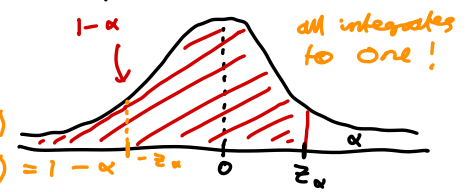
(47) Confidence Interval: Data $x_1, \dots, x_n$ modelled by $X_1, \dots, X_n$ w/ param $\theta \in \Theta$. A pair of statistics $\ell(x_1, \dots, x_n), r(x_1, \dots, x_n)$ specify an interval $[\ell(x_1, \dots, x_n), r(x_1, \dots, x_n)]$ which is called a $100(1-\alpha)\%$ CI for $\phi(\theta)$ if for any $\theta \in \Theta$, $P_\theta[\ell(X_1, \dots, X_n) \leq \phi(\theta) \leq r(X_1, \dots, X_n)] = 1 - \alpha$

"Throw data at wall: $100(1-\alpha)\%$ chance of hitting θ in your wet sponge of data"



(48) Gaussian percentiles: For $\alpha \in (0,1)$ set the $100(1-\alpha)\%$ percentile of the standard Normal distribution to be z_α s.t. if $z \sim N(0,1)$, then $P(z \leq z_\alpha) = 1 - \alpha$

note: (a) $P(z \leq z_\alpha) = 1 - \alpha \Leftrightarrow \Phi(z_\alpha) = 1 - \alpha \Leftrightarrow z_\alpha = \Phi^{-1}(1 - \alpha)$
(b) $P(z > -z_\alpha) = P(z < z_\alpha) = 1 - \alpha$



(49) Normal CI with known variance, unknown mean: data $x_1, \dots, x_n$ modelled by $X_i \stackrel{iid}{\sim} N(\mu, \sigma^2)$ unknown μ , known $\sigma^2 > 0$. choose MLE estimate for μ : $\bar{X}_n = n^{-1} \sum_{i=1}^n X_i$ $[\bar{X}_n - z_{\alpha/2} \frac{\sigma}{\sqrt{n}}, \bar{X}_n + z_{\alpha/2} \frac{\sigma}{\sqrt{n}}]$
 $\bar{X}_n \sim N(\mu, \frac{\sigma^2}{n}) \Rightarrow \frac{\bar{X}_n - \mu}{\frac{\sigma}{\sqrt{n}}} \sim N(0,1) \Rightarrow P_{\mu, \sigma^2}[\alpha \leq \frac{\bar{X}_n - \mu}{\frac{\sigma}{\sqrt{n}}} \leq b] = \Phi(b) - \Phi(a) \Rightarrow$ has prob. $1 - \alpha$ of containing μ

(50) Pivot: is a random variable whose distribution $\neq$ depend on params of underlying model

note: If we can find a pivot of simple form $\Rightarrow$ easy to build a CI!!! *n degrees of freedom*

(51) Student's t-distribution: $Z \sim N(0,1)$, $V \sim \chi_n^2$ are independent R.V.s $T = \frac{Z}{\sqrt{\frac{V}{n}}} \sim t_n$

(52) More Normal Confidence Intervals: Data $x_1, \dots, x_n$ modelled by $X_i \stackrel{iid}{\sim} N(\mu, \sigma^2)$

(a) Estimation of variance (ignore mean)

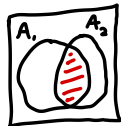
- MLE: $\hat{\mu} = \bar{x}$, $\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{n-1}{n} S^{*2}$ where $S^{*2} = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$ *pivot!*
- Fischer's thm $\Rightarrow \hat{\mu}$ independent from $\hat{\sigma}^2$ and $\frac{n-1}{\sigma^2} S^{*2} = \frac{1}{\sigma^2} \sum_{i=1}^n (x_i - \bar{x})^2 \sim \chi_{n-1}^2$
- Pivot: $\alpha \in (0,1) \Rightarrow P_{\mu, \sigma^2}(\chi_{n-1, 1-\alpha/2}^2 \leq \frac{n-1}{\sigma^2} S^{*2} \leq \chi_{n-1, \alpha/2}^2) = 1 - \alpha$

*$\chi_{n-1, \alpha}^2$ is the $100(1-\alpha)\%$ percentile of $Y \sim \chi_{n-1}^2$
 $P(Y \leq \chi_{n-1, \alpha}^2) = 1 - \alpha$*

⑥ Estimation of mean w/ variance unknown

- Setup: $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i \sim N(\mu, \frac{\sigma^2}{n}) \Rightarrow \frac{\bar{X}_n - \mu}{\frac{\sigma}{\sqrt{n}}} \sim N(0,1)$, $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2$
- Fisher's thm: $\frac{n-1}{\sigma^2} S^2 \sim \chi^2_{n-1}$ independent of $\bar{X}_n \Rightarrow T = \frac{\bar{X}_n - \mu}{\frac{S}{\sqrt{n}}} \sim t_{n-1}$
- CI: pivot $\Rightarrow T = \frac{\bar{X}_n - \mu}{\frac{S}{\sqrt{n}}} \Rightarrow P(-t_{n-1, \alpha/2} \leq T \leq t_{n-1, \alpha/2}) = 1 - \alpha$

Note: widths of CI for μ & σ^2 are not independent...



$$P(A_1 \cap A_2) \geq 1 - P(A_1^c) - P(A_2^c)$$

$\mu \in CI_\mu$, $\sigma^2 \in CI_{\sigma^2}$ w/ prob $1 - \alpha$
 $\mu, \sigma^2 \in \text{their CI}$

$$\Rightarrow P_{\mu, \sigma^2}(\bar{X}_n - t_{n-1, \alpha/2} \frac{S}{\sqrt{n}} \leq \mu \leq \bar{X}_n + t_{n-1, \alpha/2} \frac{S}{\sqrt{n}}) = 1 - \alpha$$

$$\Rightarrow [\bar{x} - t_{n-1, \alpha/2} \frac{s}{\sqrt{n}}, \bar{x} + t_{n-1, \alpha/2} \frac{s}{\sqrt{n}}] \text{ is } 100(1-\alpha)\% \text{ CI for } \mu$$

⑦ Regression case (use Fisher's thm for linear regression)...

⑧ CI for Binomial Distribution: Unknown proportion p of pop. carry genetic marker

- $\hat{p} = \frac{x}{n}$ where x carry marker in sample size of n , model as $X = \sum_{i=1}^n B_i$, $B_i \stackrel{iid}{\sim} \text{Bernoulli}(p)$
- CLT: $x = \text{sum iid R.V.} \Rightarrow X \approx N(np, np(1-p)) \Rightarrow \frac{(\frac{X}{n} - p)\sqrt{n}}{\sqrt{p(1-p)}} \approx N(0,1)$ (doesn't involve p)
- WLLN: $\frac{X}{n} \xrightarrow{p} p$
- CLT mapping thm: $\sqrt{\frac{X}{n}(1-\frac{X}{n})} \xrightarrow{d} \sqrt{p(1-p)}$
- Slutsky's thm: $\frac{\sqrt{n}(\frac{X}{n} - p)}{\sqrt{\frac{X}{n}(1-\frac{X}{n})}} \xrightarrow{d} N(0,1) \Rightarrow P_p(-z_{\alpha/2} \leq \frac{\sqrt{n}(\frac{X}{n} - p)}{\sqrt{\frac{X}{n}(1-\frac{X}{n})}} \leq z_{\alpha/2}) = 1 - \alpha$

no p in denom

$$\Rightarrow [\hat{p} - z_{\alpha/2} \sqrt{\frac{\hat{p}(1-\hat{p})}{n}}, \hat{p} + z_{\alpha/2} \sqrt{\frac{\hat{p}(1-\hat{p})}{n}}]$$

approximate CI for $p \rightarrow$ need big n

⑨ Multinomial CI: lead of one political party over another

- Setup: data $n_1, \dots, n_k$; $\sum n_i = n$, modelled by multinomial(n, p). Estimate $p_1 - p_2$
- MLE: $p_1 - p_2 \Rightarrow \hat{p}_1 - \hat{p}_2 = \frac{n_1}{n} - \frac{n_2}{n}$
- Mean: $E[\hat{p}_1 - \hat{p}_2] = \dots = p_1 - p_2$, $\text{Var}(\hat{p}_1 - \hat{p}_2) = \text{Var}(\hat{p}_1) - 2\text{Cov}[\hat{p}_1, \hat{p}_2] + \text{Var}(\hat{p}_2) = \dots$
- Approx pivot: $\frac{\frac{n_1 - n_2}{n} - (p_1 - p_2)}{\sqrt{\frac{\hat{p}_1 - \hat{p}_2 - (\hat{p}_1 - \hat{p}_2)^2}{n}}} \approx N(0,1) \Rightarrow \hat{p}_1 - \hat{p}_2 \pm \sqrt{\frac{\hat{p}_1 - \hat{p}_2 - (\hat{p}_1 - \hat{p}_2)^2}{n}} \times z_{\alpha/2}$

replace some params w/ their estimators via WLLN/CLT

approx CI for $p_1 - p_2$

⑩ Median instead of mean:

- $X_1, \dots, X_n \stackrel{iid}{\sim} \text{cauchy} \Rightarrow \frac{1}{n} \sum X_i \sim \text{cauchy} \Rightarrow$ cannot estimate θ based on $\bar{X}_n$
- Trick: compute median instead of mean! Define $Y_i = \begin{cases} 1 & \text{if } X_i \leq \theta \\ 0 & \text{or w} \end{cases} Y_i \sim \text{Bern}(\frac{1}{2})$
- $P(X_{(k)} \leq \theta \leq X_{(n-k+1)}) = P(\sum_{i=1}^k Y_i \in (k, n-k+1)) = A_k$
- Pivot: k s.t. $P(A_k) > 1 - \alpha$
- Pivot: $\sum_{i=1}^k Y_i \sim \text{Binomial}(n, \frac{1}{2})$
- Have another look... a little confusing

Chapter 9 - Statistical Tests

From estimating $\rightarrow$ Testing!

⑪ Testing Setup: Given data $x_1, \dots, x_n$ modelled by a joint distribution depending on $\theta \in \Theta$. The statistical model $\rightarrow$ likelihood $L(\theta; x_1, \dots, x_n)$. Define:

(i) null hypothesis $\theta \in \Theta_0 \subset \Theta$

(ii) alternative hypothesis $\theta \in \Theta_1 \subset \Theta$

intuition: $L(\theta) \approx$ how much does data support param value θ ?

⑫ Likelihood Ratio: $W(x_1, \dots, x_n) = W(\theta_0, \theta_1) = \frac{\sup_{\theta \in \Theta_1} \{L(\theta)\}}{\sup_{\theta \in \Theta_0} \{L(\theta)\}}$

note: Large value of $W \Rightarrow$ "evidence vs null hypth."

$\hookrightarrow$ There are $\theta \in \Theta_1$ that are much more strongly supported by the data than any $\theta \in \Theta_0$

best case scenario under alt. hypothesis

best case scenario under null hypothesis

Conventionally, these are nested hypotheses $\Theta_0 \subset \Theta_1$

possible param space

Can be simple: $\Theta_0 = \{\theta_0\}$

or compound: $\Theta_0 = (-\infty, \theta_0]$

[Likelihood Ratio Test - LRT]

$\hookrightarrow$ unlikely if null hypothesis is true $\Rightarrow$ small $p \Rightarrow$ reject θ_0

measures the maximum probability of seeing evidence against the null hypothesis when null hypothesis is actually true.

⑬ P-value of a Test Statistic: Suppose $W(x_1, \dots, x_n)$ is a statistic to test θ_0 against θ_1 , $W(x_1, \dots, x_n) \sim \text{P.V.}$

The p value for $W(x_1, \dots, x_n)$ is

$$p(x_1, \dots, x_n) = \sup_{\theta \in \Theta_0} P_{\theta}(W(x_1, \dots, x_n) \geq W(x_1, \dots, x_n))$$

$\hookrightarrow$ small p value $\Rightarrow$ strong evidence vs null.

How often will test statistic be wrong?

- (59) Decision Theoretic Interpretation:** Null hypothesis simple so $\Theta_0 = \{\theta_0\} \subset \Theta$, [nested]
- a) Pick significance level $\alpha \in (0,1)$ [Based on contextual evaluation of risk of making mistakes]
 - b) Model the scenario: obtain data $x_1, \dots, x_n$ & model w/ P.V.s $\rightarrow$ [likelihood $L(\theta; x_1, \dots, x_n)$]
 - c) Test null hypothesis vs. alt. hypothesis using LRT (or other chosen test)
 - d) Consider the resulting p-value:
 - $p(x_1, \dots, x_n) \leq \alpha \Rightarrow$ reject null hypothesis $\Theta = \theta_0$ at significance level α
 - $p(x_1, \dots, x_n) > \alpha \Rightarrow$ do not reject θ_0 at s.l. α **[~~do not~~ accept θ_0 , - could be wrong!]**

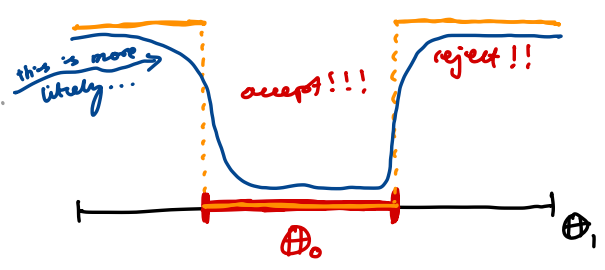
Useful lemma: X R.V. w/ cts CDF F_X ,
Then $P(F_X(X) \leq \alpha) = \alpha \quad \alpha \in [0,1]$

- (60) Uniform Distribution of p-value under a simple null hypothesis:**
Suppose $\Theta_0 = \{\theta_0\}$, $w(x)$ has cts d-set for $\theta = \theta_0$. Then, under null, P_{θ_0} , p-value $P(X) \sim \text{Uniform}(0,1)$
Pf: Let z, X be P.V. $z \perp X$, distributed P_{θ_0} , by test statistic def ^{distribution of $w(x)$ when $\theta = \theta_0$}
 $p(X) = P_{\theta_0}(w(z) \geq w(x) | X) = 1 - P_{\theta_0}(w(z) < w(x) | X) = 1 - F_w(w(x); \theta_0)$
 $P_{\theta_0}(p(X) \leq \alpha) = P_{\theta_0}(1 - F_w(w(x); \theta_0) \leq \alpha) = P_{\theta_0}(F_w(w(x); \theta_0) \geq 1 - \alpha) = 1 - P_{\theta_0}(F_w(w(x); \theta_0) \leq 1 - \alpha) = 1 - (1 - \alpha) = \alpha$
 (The t-test)

- (61) One Sample t-test:**
- Setup: data $x_1, \dots, x_n$ modeled by $X_1, \dots, X_n \stackrel{iid}{\sim} N(\mu, \sigma^2)$. $\Theta_0 = \{(\mu, \sigma): \mu = 0, \sigma > 0\}$, $\Theta_1 = \{(\mu, \sigma): \mu \in \mathbb{R}, \sigma > 0\}$
 - Likelihood Ratio: $w(\theta_0, \theta_1) = \frac{L(\bar{x}, \hat{\sigma}^2)}{L(0, \hat{\sigma}^2)} = \dots = \left(1 + \frac{n\bar{x}^2}{\sum_{i=1}^n (x_i - \bar{x})^2}\right)^{\frac{n}{2}} = \left(1 + \frac{t^2}{n-1}\right)^{\frac{n}{2}}$ where $t = \frac{\bar{x}}{\sqrt{\frac{\hat{\sigma}^2}{n}}}$, $\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$
 - Fischer's thm: under $\mu=0$, $t(x) \sim t_{n-1}$ [30]
 - $w(x)$ ↑ function of $t^2 \Rightarrow p(x) = P_{0, \sigma^2}(w(x) > w(x')) = P_{0, \sigma^2}(t(x)^2 > t(x')^2) = 2(1 - F_{t_{n-1}}(|t(x)|))$
 - Reject null hypothesis if $|t(x)| \geq t_{n-1, \alpha/2}$
- note: rejecting null hypothesis $\Leftrightarrow \theta \notin 100(1-\alpha)\%$ confidence interval for μ . **LINTH!**
- be very careful what α value you choose!

- (62) Significance & power of a test:** Θ - parameter space, $\mathcal{H}_0: \theta \in \Theta_0$, $\mathcal{H}_1: \theta \in \Theta_1$, $\Theta_0 \subset \Theta_1 \subset \Theta$
- A test has significance level α if it has prob $\leq \alpha$ of rejecting $\mathcal{H}_0$ when it is true
 - A test $\rightarrow$ test statistic $w(x)$: reject $\mathcal{H}_0$ if $w(x) \geq c$ where c chosen to achieve sig. level α .
Set $R_\alpha = \{x: w(x) \geq c\}$. Reject $\mathcal{H}_0$ if $x \in R_\alpha$ [rejection region]
 - Type I error: rejecting $\mathcal{H}_0$ when $\mathcal{H}_0$ is true.
 $\hookrightarrow$ significance/size α bounds $P(\text{Type I error})$: $\sup_{\theta \in \Theta_0} P(w(x) \geq c) = \sup_{\theta \in \Theta_0} P(x \in R_\alpha) \leq \alpha$
 - Type II error: not rejecting $\mathcal{H}_0$ when $\mathcal{H}_0$ false.
 $\hookrightarrow$ The power of a test controls $P(\text{Type II error})$: $\beta(\theta) = P_\theta(X \notin R_\alpha)$ for $\theta \in \Theta_1 \setminus \Theta_0$
 - In brief:
 - Power function of a test: $\beta(\theta) = P_\theta(\mathcal{H}_0 \text{ rejected})$
 - Significance level / size of test: $\alpha = \sup_{\theta \in \Theta_0} \beta(\theta)$
- NOTE: Ideally $\beta(\theta) \approx 1$ for $\theta \in \Theta_1 \setminus \Theta_0$
 $\alpha \approx 0$ for $\theta \in \Theta_0$
- Power = $P(\text{reject } \mathcal{H}_0)$ when $\mathcal{H}_0$ false
 Size = $\max P(\text{reject } \mathcal{H}_0)$ when $\mathcal{H}_0$ true

$\hookrightarrow$ Not possible $\because \beta(\theta)$ is continuous...



Boostr notes to recall:

- (A) MGFs Specify Distributions: If X, Y are two random variables on $\mathbb{R}$ with $M_X(t) = M_Y(t)$ on some interval $(-c, c)$ with $0 \in (-c, c)$. Then X and Y have the same distribution.
 X has c.d.f.v.
- (B) Univariate Transformation: If $Y = g(X)$ where $g: \mathbb{R} \rightarrow \mathbb{R}$ is increasing, bijective, continuously differentiable & ~~positive~~, then $F_Y(y) = F_X(g^{-1}(y))$ and $f_Y(y) = \frac{1}{g'(g^{-1}(y))} f_X(g^{-1}(y))$
- (C) $X \sim \text{MVN}$: $X \in \mathbb{R}^n$ has MVN distribution if $X = \mu + AZ$ where $\mu \in \mathbb{R}^n$ is the mean vector, $Z \in \mathbb{R}^m$ is an m -vector of m iid $N(0, 1)$ random numbers and $A \in \mathbb{R}^{n \times m}$ is a matrix such that $V = AA^T \in \mathbb{R}^{n \times n}$ is the variance-covariance matrix, so that $X \sim \text{MVN}(\mu, V)$
- (D) A distribution is memoryless if $P(X > u+t | X > t) = P(X > u)$
- (E) Fischer's theorem: If $X_1, \dots, X_n$ are a sequence of $X(\mu, \sigma^2)$ random variables, then setting $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$ and $S^{*2} = \sum_{i=1}^n (X_i - \bar{X}_n)^2$, then the following are independent:
 $\frac{\bar{X}_n - \mu}{\sigma/\sqrt{n}} \sim N(0, 1)$, $\frac{1}{\sigma^2} S^{*2} \sim \chi_{n-1}^2$
- (F) Convergence:
- $X_n \rightarrow X$ in 2-mean if $E[(X_i - X)^2] \xrightarrow{n \rightarrow \infty} 0$
 $\Downarrow$
 - $X_n \rightarrow X$ in prob if $\forall \varepsilon > 0$ $P(|X_n - \mu| > \varepsilon) \xrightarrow{n \rightarrow \infty} 0$
 $\Downarrow$
 - $X_n \rightarrow X$ in distribution if $E[h(X_n)] \xrightarrow{n \rightarrow \infty} E[h(X)]$ $\forall$ bdd $h: \mathbb{R} \rightarrow \mathbb{R}$
- (G) WLLN: $X_1, X_2, \dots$ sequence of random variables iid with finite mean μ and variance σ^2 . If $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$ is the sample mean, then $\bar{X}_n \rightarrow \mu$ in prob.
Pf: $P(|\bar{X}_n - \mu| > \varepsilon) \leq \frac{\text{var}(\bar{X}_n)}{\varepsilon^2} = \frac{\sigma^2}{n\varepsilon^2} \xrightarrow{n \rightarrow \infty} 0 \quad \forall \varepsilon > 0$
- (H) CLT: $X_1, X_2, \dots$ sequence of iid random variables w/ zero mean and finite variance, and have a common MGF defined on some interval $(-c, c)$ with $0 \in (-c, c)$. Then $\sqrt{n} \bar{X}_n \Rightarrow N(0, \sigma^2)$
- (I) c.t.s mapping Thm: If $X_n \Rightarrow X$ in $\mathbb{R}^m$ and $g: \mathbb{R}^m \rightarrow \mathbb{R}^n$ is continuous then $g(X_n) \Rightarrow g(X)$
Slutsky's & c.t.s mapping are both convergence in dist.
- (J) Slutsky's Thm: If $X_n \Rightarrow X$, $Y_n \Rightarrow c$, then (i) $X_n + Y_n \Rightarrow X + c$ (ii) $X_n Y_n \Rightarrow Xc$ and (iii) $\frac{X_n}{Y_n} \Rightarrow \frac{X}{c}$ if $c \neq 0$
- (K) An estimator is consistent if $T_n \xrightarrow{P} \theta$
- (L) p-value: $p(x_1, \dots, x_n) = \sup_{\theta \in \theta_0} P_\theta(W(X) \geq W(x))$. Max probability of seeing evidence against the null hypothesis when θ_0 is actually true.
- (M) Type I: $P(\text{reject } \theta_0 \text{ when } \theta_0 \text{ true})$
Type II: $P(\text{don't reject } \theta_0 \text{ when } \theta_0 \text{ false})$

Rejection region $R_\alpha = \{x: w(x) \geq c\}$ where c chosen to achieve s.l. α .
 reject H_0 if $x \in R_\alpha$.

$$\sup_{\theta \in \Theta_0} P_\theta(w(X) \geq c) = \sup_{\theta \in \Theta_0} P_\theta(X \in R_\alpha) \leq \alpha$$

Power $\beta(\theta) = P_\theta(H_0 \text{ rejected})$

$$\alpha = \sup_{\theta \in \Theta_0} \beta(\theta)$$

Recall checks:

- If $\mathbb{E}[e^{tx}] < \infty \wedge \forall t \in (-c, c)$ then MGF of X $m_X(t) = \mathbb{E}[e^{tx}] = 1 + t\mathbb{E}[X] + \frac{1}{2}t^2\mathbb{E}[X^2] + \dots$
 $\forall t \in (-c, c)$ where $c > 0$ some constant.
- $\int_x^\infty \lambda e^{-\lambda x} dx = \left[-e^{-\lambda x} \right]_x^\infty = e^{-\lambda x} \leftarrow \text{neat!}$

List of Sneaky Tricks

- $X_1, \dots, X_n$ w/ $X_i \sim \text{Exponential}(\lambda)$
 - $\sum_{i=1}^n X_i \sim \text{Gamma}(n, \lambda)$ [Sum of n iid exp is gamma]
 - $P(X > y) = e^{-\lambda y}$ which is easier to work w/!
 - $E\left[\frac{1}{\sum_{i=1}^n X_i}\right] = E\left[\frac{1}{Y}\right] = \int_0^\infty \frac{1}{y} \frac{1}{\Gamma(n)} \lambda^n y^{n-1} e^{-\lambda y} dy$
 $Y \sim \text{Gamma}(n, \lambda)$
 - Memoryless $P(X > u+t | X > u) = P(X > t) \quad \forall u, t > 0$
 - $\text{Gamma}(n, \frac{1}{2}) \sim \chi_n^2$
- Uniform distribution on a weird area: $\frac{1}{\text{Area}(A)} \int_{\mathbb{R}^2} \mathbb{I}\{(x,y) \in A\} dA$
 - ↳ use symmetry $\text{Cov}(X, Y) = \text{Cov}(X, -Y)$
- Transformations then requires STRICTLY $\uparrow$ function
- State WHEN you use a named Theorem.
 - ↳ Sum iid Random Variables $\Rightarrow$ CLT applies
 - ↳ exp cts $\Rightarrow$ apply CMT
- Fisher's Thm for sequence $X_1, \dots, X_n$ but μ, σ^2 unknown.
 - ↳ $\frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X})^2 \sim \chi_{n-1}^2$
 - $n-1$ $Z \sim N(0,1)$ R.V.
so var > 1
 - so $E\left[\sum_{i=1}^n (X_i - \bar{X})^2\right] = \sigma^2 E\left[Z_1^2 + \dots + Z_{n-1}^2\right] = (n-1)\sigma^2$
- Be VERY careful!!! NOT $X_1 \sim \text{Exp}(\lambda_1), X_2 \sim \text{Exp}(\lambda_2)$
- Justify MLE makes.
- NOTE: SHIFTED DISTRIBUTIONS $\Rightarrow X = Z + c$ (not odd factors)
- DEFINE ALL YOUR VARIABLES. (leave nothing unexplained).
- CLT
 - iid
 - common m.g.f
 WLLN
 - iid
 - finite μ, σ^2 .
- Say test statistic $w(X) \downarrow$, defined on $X > 0$, then w will be biggest on $[0, c]$ so that is the rejection region.
 - ↳ use $\uparrow/\downarrow$ of test statistic to determine shape of R_α

- $X \sim N(0,1) \quad Y = X^2$



$$\hookrightarrow F_Y(y) = P(Y \leq y) = P(X^2 \leq y) = \int_{-\infty}^{\infty} \mathbb{1}_{\{x^2 \leq y\}} dx = \int_{-\sqrt{y}}^{\sqrt{y}} f_X(x) dx$$

$$\hookrightarrow f_Y(y) = F'_Y(y)$$

USE $\mathbb{1}$ functions
to identify ranges

- $Z \in \mathbb{R}^n \quad E[Z] = \mu$

$$\hookrightarrow V_{ij} = E[(Z_i - E[Z_i])(Z_j - E[Z_j])]$$

$$\begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix} = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix} = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix}$$

$$= E[Z_i Z_j] - \mu_i \mu_j \Rightarrow V = Z Z^T + \mu \mu^T$$

- conditional expectations:

- Take out what's known

- Tower property (can sneak an extra $E[\cdot | X]$ in!)

- Sum n iid R.v. $\Rightarrow$ apply CLT $\Rightarrow \log N$.

- Several Normals $\rightarrow$ MVN(μ, V) where $V \in \mathbb{R}^{n \times n}$

$$\hookrightarrow$$
 Specify rows/columns of matrix carefully. $AA^T = V$

$$\hookrightarrow$$
 watch out for X_{n-2}^2 at any point = ALL RELATED!

- $Y_i = \gamma + \beta x_i + \sigma \varepsilon_i \Rightarrow Y_i \sim N(\gamma + \beta x_i, \sigma^2) \rightarrow$ PDF!

- Lognormal multipliers...

$$T = \frac{\bar{Z}}{\sqrt{\frac{U}{n}}} \sim t_n$$

$Z \sim N(0,1)$
 $U \sim \chi_n^2$

- For cI

$$\hookrightarrow \mu, \sigma^2$$
 unknown $\Rightarrow$ Fishers, t_{n-1} create it!

$$\hookrightarrow \sigma^2$$
 known $\Rightarrow$ simpler χ_{n-1}^2

- Likelihood ratio statistic $\Rightarrow$ PICT values of parameters in diff scenarios. E.g.
$$\mu = \frac{n\bar{x} + m\bar{y}}{n+m}$$

- NLCN proof: $H_0: \theta = 0$

$$P(|\bar{X}_n - \mu| > \varepsilon) = P((\bar{X}_n - \mu)^2 > \varepsilon^2) \leq \frac{\text{Var}(\bar{X}_n - \mu)}{\varepsilon^2} = \frac{n\sigma^2}{n^2\varepsilon^2} \xrightarrow{n \rightarrow \infty} 0$$

- Choose questions CAREFULLY!

$$\hookrightarrow$$
 make sure you can do all parts/ have ideas for how to attack.

- If you can't JV a logical step, LEAVE IT & come back.

$$\hookrightarrow$$
 all double, just can't go wrong.