

Properties of Random Variables

We have defined Random Variables thoroughly in the Probability Theory notes, and now we will discuss some commonly referenced properties of Random Variables. The following do not necessarily exist for all random variables, but if they do exist they are defined as follows:

The expected value, also called the mean of a random variable (denoted $E[X]$ or sometimes μ):

$$E[X] = \int_{-\infty}^{\infty} x f(x) dx \text{ if } X \text{ has a density function}$$

$$E[X] = \sum_{i=1}^{\infty} x_i P(X=x_i) \text{ if } X \text{ has no density function}$$

if X has a density function, we often call it continuous, and if not, we often call it discrete.

The Variance of X , often denoted $\text{Var}(X)$:
$$\text{Var}(X) \equiv E[(X - E[X])^2] = E[X^2] - E[X]^2$$

Properties of Expectation and Variance

Linearity: $E[X+Y] = E[X] + E[Y]$

iff X and Y are independent: $\text{Var}(X+Y) = \text{Var}(X) + \text{Var}(Y)$

Multiplication by a scalar: $E[aX] = a E[X]$
 $\text{Var}(aX) = a^2 \text{Var}(X)$, e.g. $\text{Var}\left(\frac{\sum X}{\sqrt{n}}\right) = \frac{1}{n} \text{Var}(\sum X)$

Adam's Rule: $E[X] = E[E[X|Y]]$

Eve's Rule:

$$\text{Var}(X) = E[\text{Var}(X|Y)] + \text{Var}(E[X|Y])$$

Jensen's Inequality: $E[g(X)] \geq g(E[X])$
if g is convex

$$E[g(X)] = \int_{-\infty}^{\infty} g(x) f(x) dx \quad \text{or}$$

$$\sum_{i=1}^{\infty} g(x) P(X=x) \quad \text{if } X \text{ is discrete}$$

Cauchy-Schwarz Inequality: $\text{Cov}(X, Y)^2 \leq \text{Var}(X) \text{Var}(Y)$

Linearity of Expectation Example

For $X_1, \dots, X_n$ iid, prove $E\left[\frac{\bar{X}_{n-1}}{\bar{X}_n}\right] = 1$

$$\begin{aligned} E\left[\frac{\bar{X}_{n-1}}{\bar{X}_n}\right] &= E\left[\frac{\frac{X_1 + X_2 + \dots + X_{n-1}}{n-1}}{\frac{X_1 + X_2 + \dots + X_n}{n}}\right] = \frac{n}{n-1} E\left[\frac{X_1 + \dots + X_{n-1}}{X_1 + \dots + X_n}\right] \\ &= \frac{n}{n-1} \left(E\left[\frac{X_1}{X_1 + \dots + X_n}\right] + E\left[\frac{X_2}{X_1 + \dots + X_n}\right] + \dots + E\left[\frac{X_{n-1}}{X_1 + \dots + X_n}\right] \right) \end{aligned}$$

Since $X_1, \dots, X_{n-1}$ are iid, this is

$$\begin{aligned} &= \frac{n}{n-1} \left(E\left[\frac{X_1}{X_1 + \dots + X_n}\right] + E\left[\frac{X_1}{X_1 + \dots + X_n}\right], \dots + E\left[\frac{X_1}{X_1 + \dots + X_n}\right] \right) \\ &= \frac{n}{n-1} \cdot (n-1) \cdot E\left[\frac{X_1}{X_1 + \dots + X_n}\right] = n E\left[\frac{X_1}{X_1 + \dots + X_n}\right] = E\left[\frac{n X_1}{X_1 + \dots + X_n}\right] \end{aligned}$$

Again, since $X_1, \dots, X_n$ are iid, this is

$$\begin{aligned} &= E\left[\frac{X_1 + X_1 + \dots + X_1}{X_1 + X_2 + \dots + X_n}\right] = \left(E\left[\frac{X_1}{X_1 + X_2 + \dots + X_n}\right] + \dots + E\left[\frac{X_1}{X_1 + \dots + X_n}\right] \right) \\ &= \left(E\left[\frac{X_1}{X_1 + \dots + X_n}\right] + E\left[\frac{X_2}{X_1 + \dots + X_n}\right], \dots + E\left[\frac{X_n}{X_1 + \dots + X_n}\right] \right) \\ &= E\left[\frac{X_1 + X_2 + \dots + X_n}{X_1 + X_2 + \dots + X_n}\right] = E[1] = 1 \end{aligned}$$

Parameters

Often, for random variables with PDF's and/or PMF's or CDF's of the same functional form, we use parameters as placeholders for numbers. For example, the Normal distribution has PDF

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

Where μ and σ^2 are the parameters.

Most of Parametric Inference (inference with parameters) is focused on assuming we have some realizations of a random variable with known functional form and unknown parameters and using the realizations to learn about the parameters.

For the normal distribution, $E[X] = \mu$ and $\text{Var}(X) = \sigma^2$

Common Distributions

Normal Distribution:

$$f(x|\mu, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(x-\mu)^2}{\sigma^2}}$$

$$E[X] = \mu \quad \text{Var}(X) = \sigma^2$$

Uniform Distribution:

$$f(x|a, b) = \frac{1}{b-a}$$

$$E[X] = \frac{b+a}{2} \quad \text{Var}(X) = \frac{(b-a)^2}{12}$$

Binomial Distribution:

$$P(X=x | n, p) = \binom{n}{x} p^x (1-p)^{n-x}$$

↓ distribution of
the number of heads

$$E[X] = np \quad \text{Var}(X) = np(1-p) \quad \text{of } n \text{ coin flips}$$

Poisson Dist:

$$P(X=x | \lambda) = \frac{e^{-\lambda} \lambda^x}{x!}$$

$$E[X] = \lambda \quad \text{Var}(X) = \lambda$$

Exponential Dist:

$$f(x|\beta) = \frac{1}{\beta} e^{-\frac{x}{\beta}}$$

$$E[X] = \beta \quad \text{Var}(X) = \beta^2$$

Can also be parameterized equivalently with $\lambda = \frac{1}{\beta}$

Moment Generating Functions

The moment generating function generates the moments! The MGF of a r.v. X is defined as a function of t :

$$M_X(t) = E[e^{tX}]$$

The MGF, like the CDF, uniquely identifies the distribution of a r.v.. The n th moment of X is defined as $E[X^n]$, and it can be generated by taking the n th derivative of $M_X(t)$ w.r.t. t and setting $t=0$:

$$M_X^{(n)}(0) = E[X^n]$$

The moment generating function of the sum of two independent random variables is the product of their moment generating functions. This fact is often used to derive the distribution of the sum of two r.v.s

$$M_{X_1+X_2}(t) = M_{X_1}(t) M_{X_2}(t)$$

MGF Example

The normal random variable can be shown to have MGF $M_X(t) = E[e^{tx}] = \exp(\mu_X t + \frac{\sigma_X^2 t^2}{2})$ where X has parameters μ_X and σ_X^2

If Y is an independent normal with mean μ_Y and variance σ_Y^2 , then Y also has MGF $M_Y(t) = E[e^{ty}] = \exp(\mu_Y t + \frac{\sigma_Y^2 t^2}{2})$

Then the MGF of $X+Y$ is

$$\begin{aligned} M_{X+Y}(t) &= M_X(t) M_Y(t) = \exp\left(\mu_X t + \frac{\sigma_X^2 t^2}{2}\right) \exp\left(\mu_Y t + \frac{\sigma_Y^2 t^2}{2}\right) \\ &= \exp\left(\mu_X t + \frac{\sigma_X^2 t^2}{2} + \mu_Y t + \frac{\sigma_Y^2 t^2}{2}\right) \\ &= \exp\left((\mu_X + \mu_Y)t + \frac{(\sigma_X^2 + \sigma_Y^2)t^2}{2}\right) \end{aligned}$$

Thus $X+Y$ has normal distribution with mean $\mu_X + \mu_Y$ and variance $\sigma_X^2 + \sigma_Y^2$

This fact, that sums of iid normals are normal with summed variances and means, is recurrent

Covariance & Correlation

The covariance of two random variables X and Y is defined as

$$\begin{aligned}\text{Cov}(X, Y) &= E[(X - E[X])(Y - E[Y])] \\ &= E[XY] - E[X]E[Y]\end{aligned}$$

For an n -dimensional random vector $\vec{X}$, the variance-covariance matrix Σ has $\text{Var}(X_i) = \Sigma_{ii}$ and has $\text{Cov}(X_i, X_j) = \Sigma_{ij} \quad \forall i, j \leq n$, Σ is thus $n \times n$

The correlation between X and Y is defined as

$$\text{Corr}(X, Y) = \frac{\text{Cov}(X, Y)}{\sqrt{\text{Var}(X)} \sqrt{\text{Var}(Y)}}$$

The $\sqrt{\text{Var}(X)}$ is often called the standard deviation of X

Transformation of Variables

Often, we want to re-express a pair of random variables in a different way, for example, express X, Y as U, V where $U = XY$ and $V = \frac{X}{Y}$. This is done through transformation of variables when X and Y have a pdf.

First, solve for X and Y in terms of U and V . e.g. $X = \frac{U}{V}$, $Y = \frac{U/V}{V} = \frac{U}{V^2}$, $Y^2 = \frac{U}{V}$, $X = \frac{U}{\sqrt{U/V}}$

Then, the joint PDF of U and V is

$$f_{U,V}(u,v) = f_{X,Y}\left(x\left(\begin{smallmatrix} \text{in } v \text{ and} \\ \text{terms} \end{smallmatrix}\right), y\left(\begin{smallmatrix} \text{in } u \text{ and} \\ \text{terms} \end{smallmatrix}\right)\right) |J|$$

Where J is the Jacobian:

$$|J| = \left| \det \begin{bmatrix} \frac{\partial x}{\partial u} & \frac{\partial x}{\partial v} \\ \frac{\partial y}{\partial u} & \frac{\partial y}{\partial v} \end{bmatrix} \right|$$

T Transformation of Variables

Example:

$$f_{X,Y} = \begin{cases} 4xy & 0 < x < 1, 0 < y < 1 \\ 0 & \text{o.w.} \end{cases}$$

$$u = X + Y \quad v = X - Y$$

$$\text{Solve for } X \text{ and } Y: \quad X = \frac{u+v}{2} \quad Y = \frac{u-v}{2}$$

$$|J| = \left| \det \begin{bmatrix} \frac{\partial x}{\partial u} & \frac{\partial x}{\partial v} \\ \frac{\partial y}{\partial u} & \frac{\partial y}{\partial v} \end{bmatrix} \right| = \left| \begin{bmatrix} \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & -\frac{1}{2} \end{bmatrix} \right| = \left| -\frac{1}{4} - \frac{1}{4} \right| = \left| -\frac{1}{2} \right|$$

$$= \frac{1}{2}, \quad f_{u,v}(u,v) = 4 \left(\frac{u+v}{2} \right) \left(\frac{u-v}{2} \right) \cdot \frac{1}{2}$$

$$= \frac{u^2 - v^2}{2}. \quad \text{Since } 0 < x < 1 \text{ and } 0 < y < 1,$$

$$0 < \frac{u+v}{2} < 1 \Rightarrow -u < v < 2-u \quad 0 < u < 2$$

$$0 < \frac{u-v}{2} < 1 \Rightarrow -u < -v < 2-u \quad \text{for } 0 < u \leq 1$$

$$-u < v < u$$

$$\text{for } 1 < u < 2$$

$$-(2-u) < v < 2-u$$

IID Random Sample

The random variables $X_1, \dots, X_n$ are called a random sample of size n from the population $F(x)$ if $X_1, \dots, X_n$ are mutually independent random variables and the marginal pdf or pmf of each X_i is the same function $F(x)$. This is equivalent to the statement that $X_1, \dots, X_n$ are i.i.d., independent identically distributed. Often, instead of stating the pdf, the functional form of the distribution the X 's are sampled from is stated with an unknown parameter.

For example: let $X_1, \dots, X_n$ be an i.i.d. random sample from a $\text{Normal}(\mu, \sigma^2)$

An i.i.d random sample from pdf $f(x)$ or pmf $P(X=x)$ has joint pdf or pmf

$$f(x_1, \dots, x_n) = \prod_{i=1}^n f(x_i) \quad \text{or} \quad P(X_1=x_1, \dots, X_n=x_n) = \prod_{i=1}^n P(X_i=x_i)$$

In general, if the pdf/pmf can be factored into $f(x_1)f(x_2)\dots$ then $X_1, X_2, \dots$ are independent r.v.s

Statistics & Estimators

A Statistic is a function that maps the result of a random sample to the real line.

A statistic of a sample $\vec{X}$ is often notated $T(\vec{X})$

Examples of Statistics:

Sample mean: $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$

Sample variance: $\hat{\sigma}_x^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$

An estimator is a statistic used to estimate some quantity of interest called an estimand

Often, the estimand is a parameter of the distribution we are sampling from

For normal: $\bar{X}$ is an estimator of μ
 $\hat{\sigma}_x^2$ is an estimator of σ_x^2

Order Statistics

The order statistics of a random sample $X_1, \dots, X_n$ are the sampled values placed in ascending order. These are denoted $X_{(1)}, \dots, X_{(n)}$

The order statistics are random variables that satisfy $X_{(1)} \leq \dots \leq X_{(n)}$. $X_{(1)}$ is the minimum and $X_{(n)}$ is the maximum

If n is odd, the sample median is $X_{(\frac{n+1}{2})}$, the middle order statistic

If n is even, the sample median is $\frac{X_{(\frac{n}{2})} + X_{(\frac{n}{2}+1)}}{2}$, the average of the middle two order statistics

The range is $X_{(n)} - X_{(1)}$

The p th quantile is the minimum order statistic such that $\frac{p}{100}$ proportion of the sampled values are less than that order statistic

Order Statistics

Let $X_1, \dots, X_n$ be an i.i.d r.v. from a discrete distribution with pmf $f_X(x_i) = p_i$, where $x_1 < x_2 < \dots$ are the possible values of X in ascending order (e.g. for $\text{bin}(2, 0.5)$, this is $x_1=0, x_2=1, x_3=2$). Let $X_{(1)}, \dots, X_{(n)}$ denote the order statistics from the sample. Then

$$P(X_{(j)} \leq x_i) = \sum_{k=j}^n \binom{n}{k} p_i^k (1-p_i)^{n-k} \quad (1)$$

Where $P_0 = 0, P_1 = P_0 + p_1, P_2 = P_1 + p_2, P_3 = P_2 + p_3, \dots$

Also,

$$P(X_{(j)} = x_i) = \sum_{k=j}^n \binom{n}{k} [p_i^k (1-p_i)^{n-k} - p_{i-1}^k (1-p_{i-1})^{n-k}]$$

PF of (1): For each x_i the probability that any given observation is below it is p_i , so Y , the number of observations less than or equal to x_i is binomially distributed with $n=n$ and $p=p_i$. As a result, the probability $X_{(j)} \leq x_i$, the probability of n are $\leq x_i$, is $P(Y \leq x_i)$

Order Statistics Continuous

When $X_1, \dots, X_n$ is an iid random sample from a continuous distribution with pdf $f_X(x)$ and cdf $F_X(x)$, the pdf of the j th order statistic $X_{(j)}$ is

$$f_{X_{(j)}}(x) = \frac{n!}{(j-1)!(n-j)!} f_X(x) [F_X(x)]^{j-1} [1-F_X(x)]^{n-j}$$

This comes from $F_{X_{(j)}}(x)$, which is, utilizing the same exact logic as in the discrete case:

$$F_{X_{(j)}}(x) = P(Y \geq j) = \sum_{k=j}^n \binom{n}{k} [F_X(x)]^k [1-F_X(x)]^{n-k}$$

The pdf is obtained by simply differentiating the rhs of the above and doing some algebra

Order Statistics Continuous

$$\begin{aligned}
 & \frac{d}{dx} \sum_{k=j}^n \binom{n}{k} [F_x(x)]^k [1-F_x(x)]^{n-k} \\
 &= \sum_{k=j}^n \binom{n}{k} (k F_x(x)^{k-1} f_x(x) [1-F_x(x)]^{n-k} - [F_x(x)]^k [1-F_x(x)]^{n-k-1} f_x(x) (n-k)) \\
 &= \sum_{k=j}^n \binom{n}{k} k F_x(x)^{k-1} f_x(x) [1-F_x(x)]^{n-k} - \sum_{k=j}^n [F_x(x)]^k [1-F_x(x)]^{n-k-1} f_x(x) (n-k) \binom{n}{k} \\
 &= \binom{n}{j} j F_x(x)^{j-1} f_x(x) [1-F_x(x)]^{n-j} + \sum_{k=j+1}^n \binom{n}{k} k F_x(x)^{k-1} f_x(x) [1-F_x(x)]^{n-k} \\
 &\quad - \binom{n}{j} (n-j) F_x(x)^j [1-F_x(x)]^{n-j-1} f_x(x) - \sum_{k=j+1}^n [F_x(x)]^k [1-F_x(x)]^{n-k-1} f_x(x) (n-k) \binom{n}{k} \\
 &= \binom{n}{j} j F_x(x)^{j-1} f_x(x) [1-F_x(x)]^{n-j} + \left(\sum_{k=j+1}^n \binom{n}{k} k F_x(x)^{k-1} f_x(x) [1-F_x(x)]^{n-k} - \sum_{k=j+1}^n \binom{n}{k} (n-k) F_x(x)^k [1-F_x(x)]^{n-k-1} f_x(x) \right)
 \end{aligned}$$

Now, $\forall j, j+1$, $\frac{n!}{j!(n-j)!} (n-j) F_x(x)^j f_x(x) [1-F_x(x)]^{n-j-1} = \frac{n!}{(j+1)!(n-(j+1))!} j F_x(x)^{j+1-1} f_x(x) [1-F_x(x)]^{n-(j+1)}$

So the two sums cancel and since $\frac{n!}{j!(n-j)!} j = \frac{n!}{(j-1)!(n-j)!}$, we have the result:

$$= \frac{n!}{(j-1)!(n-j)!} F_x(x) [F_x(x)]^{j-1} [1-F_x(x)]^{n-j}$$

SUFFICIENCY

A statistic $T(\vec{X})$ of an i.i.d. sample $X_1, \dots, X_n$ assumed to come from a distribution with unknown parameter θ is said to be SUFFICIENT for θ if any inference about θ depends on $\vec{X}$ only through $T(\vec{X})$.

That is, if two sample points X and Y are observed such that $T(X) = T(Y)$ then inference about θ is the same whether $X=x$ or $Y=y$ is observed.

The conditional distribution of the sample $\vec{X}$ given the value of $T(\vec{X})$ does not depend on θ .

If you can factor the likelihood function of the data, the PDF or PMF considered as a function of θ , $L(\theta | \vec{X})$, into a function of the chosen $T(\vec{X})$ statistic and θ and a function of only the data, $g(\theta | T(\vec{X}))h(\vec{X})$, then $T(\vec{X})$ is sufficient.

If the likelihood divided by the pdf of $T(\vec{X})$ does not contain θ , $T(\vec{X})$ is sufficient.

If θ is multi dimensional, parameters $\theta_1, \dots, \theta_k$ are all unknown, and $T(\vec{X})$ is a k -dimensional sufficient statistic, $T(\vec{X})$ is said to be Minimal Sufficient

Completeness

A statistic of a random sample $\vec{X}$ assumed to be i.i.d. coming from a distribution with unknown parameter θ is said to be complete (or the family of distributions is said to be complete) if $\forall g \ E[g(T(\vec{X}))] = 0$ implies $g(x) = 0$, the only way to make a function of the sufficient statistic that always has expectation zero is to "cheat" by just setting it equal to zero.

A distribution is said to be part of the "exponential family" if its PDF/PMF for a single observation can be manipulated into the following form, and if it is/are, then a result about completeness follows:

$$f(x|\theta) = h(x)c(\theta) \exp\left(\sum_{j=1}^K w_j(\theta) t_j(x)\right)$$

Where $\theta = (\theta_1, \dots, \theta_K)$, if θ is K -dimensional. Then $T(\vec{X}) = \left(\sum_{i=1}^n t_1(X_i), \sum_{i=1}^n t_2(X_i), \dots, \sum_{i=1}^n t_K(X_i)\right)$ is complete unless $w_1(\theta), \dots, w_K(\theta): \theta \in \Theta$ contains no open sets in $\mathbb{R}^K$.

(Completeness) can equivalently be defined as $E_\theta[g(T)] = 0$ for all θ implies $P_\theta(g(T) = 0) = 1 \ \forall \theta$

Basu's Theorem

An Ancillary Statistic is a statistic of an i.i.d. random sample $\vec{X}$, $S(\vec{X})$, that has a distribution that does not depend on θ . For example, if we are sampling from a uniform $(\theta-1, \theta+1)$, the statistic $X_{(n)} - X_{(1)}$ does not depend on θ so $S(\vec{X}) = X_{(n)} - X_{(1)}$ is ancillary.

Basu's Theorem says if $T(\vec{X})$ is a complete and minimal sufficient statistic, then $T(\vec{X})$ is independent of every ancillary statistic.

Note that the word "minimal" is unnecessary since complete statistics are minimal by definition.

UMVUE

A statistic used as an estimator of θ , $T(\vec{X})$, is unbiased if $E[T(\vec{X})] = \theta$.

If $W^*(\vec{X})$ is an unbiased estimator of θ , and $W^*(\vec{X})$ has less variance than any other unbiased estimator of θ , that is

$$\text{Var}_\theta(W^*(\vec{X})) \leq \text{Var}_\theta(W(\vec{X})) \quad \forall \theta, W_{\text{unbiased}}$$

then we call $W^*(\vec{X})$ the Uniform Minimum Variance Unbiased estimator of θ , or the UMVUE of θ .

MLE & Log-Likelihood

The likelihood $L(\theta|\vec{X})$ is the joint pdf $L(\theta|\vec{X}) = f(\vec{x}|\theta)$ or PMF $L(\theta|\vec{X}) = P(\vec{X}=\vec{x}|\theta)$ of an i.i.d. random sample as a function of the parameter θ . Since the function depends on random variables, the likelihood is a "random function".

Often, we are concerned with finding the θ that, given $X_1, \dots, X_n$, maximizes $L(\theta|\vec{X})$. This value for θ is called the maximum likelihood estimator of MLE of θ , $\hat{\theta}_{MLE}$.

Since $\log(x)$ is monotone, the θ that maximizes $L(\theta|\vec{X})$ also maximizes $\log(L(\theta|\vec{X}))$. $\log(L(\theta|\vec{X}))$ is typically easier to work with, and is called the Log-Likelihood, notated $\ell(\theta|\vec{X})$.

Often, MLE is found by equating the score to zero, checking boundaries (if any), and checking to see that second deriv of $\log-L$ is negative. i.e. Use calculus to find a global max of $\log-L$.

Score Function

The derivative of the Log-likelihood $\ell(\theta | \vec{x})$ w.r.t. θ , $\frac{d}{d\theta} \ell(\theta | \vec{x})$, is called the Score-function. When $\theta = \hat{\theta}_{MLE}$, the score is always zero, since any maximizer (or minimizer) of a function must have a tangent of slope zero.

For any value of θ , the score is a random function because $L(\theta | \vec{X})$ is a random function depending on $\vec{X}$.

The expected value of the score across $\vec{X}$ is a function of θ but always is equal to zero,
$$E_{\theta} \left[\frac{d}{d\theta} \log L(\theta | \vec{X}) \right] = 0 \quad \forall \theta$$

This is true because:
$$E_{\theta} \left[\frac{d}{d\theta} \log L(\theta | \vec{X}) \right] = \int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} \log(f(\vec{x}|\theta)) f(\vec{x}|\theta) d\vec{x} = \int_{-\infty}^{\infty} \left(\frac{\frac{\partial}{\partial \theta} f(\vec{x}|\theta)}{f(\vec{x}|\theta)} \right) f(\vec{x}|\theta) d\vec{x} = \int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} f(\vec{x}|\theta) d\vec{x} = \frac{\partial}{\partial \theta} \int_{-\infty}^{\infty} f(\vec{x}|\theta) d\vec{x} = \frac{\partial}{\partial \theta} 1 = 0$$

$E[\text{score}] = 0$ for discrete r.v.s too, just replace $\int_{-\infty}^{\infty}$ with $\sum_{i=1}^{\infty}$
Note that $g(x) = \log f(x) \Rightarrow g'(x) = \frac{f'(x)}{f(x)}$ by chain rule

Fisher Information

Recall that the Score is a random function, $\frac{d}{d\theta} \ell(\theta | \vec{X})$. The Fisher Information $I(\theta)$ is the Variance of the Score,

$$\begin{aligned} I(\theta) &\equiv E \left[\left(\frac{d}{d\theta} \ell(\theta | \vec{X}) - \underbrace{E \left[\frac{d}{d\theta} \ell(\theta | \vec{X}) \right]}_{=0} \right)^2 \right] \\ &= E \left[\left(\frac{d}{d\theta} \ell(\theta | \vec{X}) \right)^2 \right] \end{aligned}$$

The Fisher information can also be shown to be equivalent to the negative of the Hessian, which is the second derivative of the log-likelihood:

$$\begin{aligned} \text{Proof: } E[-H] &= E \left[- \frac{\partial^2}{\partial \theta^2} \log(f(\vec{X} | \theta)) \right] \\ &= E \left[- \frac{\partial}{\partial \theta} \left(\frac{\partial}{\partial \theta} \log(f(\vec{X} | \theta)) \right) \right] = E \left[- \frac{\partial}{\partial \theta} \left(\frac{\frac{\partial}{\partial \theta} f(\vec{X} | \theta)}{f(\vec{X} | \theta)} \right) \right] \\ &= E \left[\frac{\frac{\partial}{\partial \theta} f(\vec{X} | \theta) \frac{\partial}{\partial \theta} f(\vec{X} | \theta) - f(\vec{X} | \theta) \frac{\partial^2}{\partial \theta^2} f(\vec{X} | \theta)}{f(\vec{X} | \theta)^2} \right] \\ &= E \left[\left(\frac{\frac{\partial}{\partial \theta} f(\vec{X} | \theta)}{f(\vec{X} | \theta)} \right)^2 \right] - E \left[\frac{\frac{\partial^2}{\partial \theta^2} f(\vec{X} | \theta)}{f(\vec{X} | \theta)} \right] \end{aligned}$$

Fisher Information Continued

Recall that $\frac{d}{dx} \log(f(x)) = \frac{f'(x)}{f(x)}$, so

$$= E\left[\left(\frac{\partial}{\partial \theta} \log(f(\vec{x}|\theta))\right)^2\right] = \int_{-\infty}^{\infty} x \frac{\frac{\partial^2}{\partial^2 \theta} f(\vec{x}|\theta)}{f(\vec{x}|\theta)} dx$$

Integrate $\int_{-\infty}^{\infty} x \frac{\frac{\partial^2}{\partial^2 \theta} f(\vec{x}|\theta)}{f(\vec{x}|\theta)} dx$ by parts:

$$u = \frac{\partial^2}{\partial^2 \theta} f(\vec{x}|\theta), \quad v = \frac{x}{f(\vec{x}|\theta)}, \quad \int u v dx = u v dx - \int u' v dx$$

$$\text{Since } \int v dx = \int_{-\infty}^{\infty} \frac{\partial^2}{\partial^2 \theta} f(\vec{x}|\theta) = \int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} \left(\frac{\partial}{\partial \theta} f(\vec{x}|\theta) \right) dx$$

$$= \frac{\partial}{\partial \theta} \left(\int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} f(\vec{x}|\theta) dx \right) = \frac{\partial}{\partial \theta} \left(\frac{\partial}{\partial \theta} \left(\int_{-\infty}^{\infty} f(\vec{x}|\theta) dx \right) \right) = \frac{\partial^2}{\partial^2 \theta} (1) = 0$$

this becomes $= u \cdot 0 - \int u' \cdot 0 dx = 0$, thus

$$\int_{-\infty}^{\infty} x \frac{\frac{\partial^2}{\partial^2 \theta} f(\vec{x}|\theta)}{f(\vec{x}|\theta)} dx = 0, \text{ thus}$$

$$E[-H] = E\left[\left(\frac{\partial}{\partial \theta} \log(f(\theta|\vec{x}))\right)^2\right] = I(\theta)$$

Fisher Information Matrix

The Fisher information, The Variance of the Score for Parameter θ , has a natural extension to multiple Parameters: For multiple Parameters, the information matrix is the variance covariance matrix of the scores of these Parameters.

For Parameters $\theta_1, \theta_2, \dots, \theta_n \in \vec{\theta}$ each entry $I_{i,j}$ is $\text{Cov}(\text{score}(\theta_i), \text{score}(\theta_j)) =$
$$E\left[\left(\frac{\partial}{\partial \theta_i} \ell(\vec{\theta} | \vec{X})\right) \left(\frac{\partial}{\partial \theta_j} \ell(\vec{\theta} | \vec{X})\right)\right]$$

Additionally, since information matrix is the negative expectation of the Hessian:

$$I_{i,j} = -E\left[\frac{\partial^2 \ell(\vec{\theta} | \vec{X})}{\partial \theta_i \partial \theta_j}\right]$$

Fisher Scoring

SUPPOSE WE WANT TO FIND $\hat{\theta}_{MLE}$ OF A
an i.i.d. Sample FROM A DISTRIBUTION FOR
WHICH THE CORRECT $W(\vec{X}) = \hat{\theta}_{MLE}$ CANNOT BE
SOLVED FOR ANALYTICALLY, OR WE ARE TOO LAZY TO
DETERMINE IT ANALYTICALLY. WE CAN USE A
TAYLOR EXPANSION AND $l(\theta|\vec{X})$ AND $I(\theta)$ TO
APPROXIMATE $\hat{\theta}_{MLE}$ ARBITRARILY CLOSELY. THIS IS CALLED
Fisher-Scoring. IF WE DECIDE NOT TO REPLACE
 H WITH $I(\theta)$, THIS IS CALLED THE Newton-Raphson METHOD.

RECALL FROM CALC 2 THAT A TAYLOR EXPANSION OF
A FUNCTION $F(x)$ APPROXIMATES THE FUNCTION'S
VALUE AT A POINT 'a' WITH $F(x) \approx F(a) + F'(a)(x-a) + \dots$
USING THE 'FIRST ORDER' EXPANSION TO APPROXIMATE $F(x)$:
 $F(x) \approx F(a) + F'(a)(x-a)$ (LINEAR APPROXIMATION AT a)

IF WE ARE INTERESTED IN FINDING THE ROOT OF $F(x)$
(AS WE ARE IN THIS CASE FOR REASONS TO BE EXPLAINED SOON)
WE SET $F(x)=0$: $0 \approx F(a) + F'(a)(x-a)$
 $x \approx a - \frac{F(a)}{F'(a)}$

Fisher Scoring Continued

We can use the approximation $x \approx a - \frac{f(a)}{f'(a)}$ to define an iterative procedure that finds the root of $f(x)$ by repeatedly replacing a with the previous approximation of the root:

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$$
 If you run this for many iterations, and f is not something really really weird, after a while $x_{n+1} \approx x_n \approx x$ s.t. $f(x) = 0$

Now recall that the MLE has to maximize the likelihood, so its first derivative must be zero.

Thus, $\hat{\theta}_{MLE}$ can be approximated using the above Newton-Raphson technique where $f(x) = \frac{d}{d\theta} \ell(\theta | \vec{x})$, the Score

$$\theta_{n+1} = \theta_n - \frac{\frac{d}{d\theta} \ell(\theta_n | \vec{x}) \rightarrow \text{Score}}{\frac{d^2}{d\theta^2} \ell(\theta_n | \vec{x}) \rightarrow \text{Hessian}}$$

Since the Score and Hessian are both random this can be unstable, so replace H with $E[H | \theta = \theta_n]$ to get Fisher Scoring

$$\theta_{n+1} = \theta_n - \frac{\frac{d}{d\theta} \ell(\theta_n | \vec{x})}{E[H | \theta = \theta_n]} \Rightarrow \theta_{n+1} = \theta_n + \frac{\frac{d}{d\theta} \ell(\theta_n | \vec{x})}{I(\theta_n)}$$

Cramer - Rao Lower Bound

Let $X_1, \dots, X_n$ be a sample with PDF $f(x|\theta)$, and let $W(\vec{X})$ be an estimator of θ with finite variance and expectation satisfying

$$\frac{d}{d\theta} E[W(\vec{X})] = \int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} W(\vec{X}) f(\vec{X}|\theta) d\vec{x}$$

$$\text{Then } \text{Var}(W(\vec{X})) \geq \frac{\left(\frac{d}{d\theta} E[W(\vec{X})]\right)^2}{I(\theta)}$$

The rhs is known as the Cramer - Rao lower bound because it is the lowest possible variance of an estimator $W(\vec{X})$

When we constrain $W(\vec{X})$ to be unbiased $E[W(\vec{X})] = \theta$, then

$$\text{Var}(W(\vec{X})) \geq \frac{\left(\frac{d}{d\theta} \theta\right)^2}{I(\theta)} \Rightarrow \text{Var}(W(\vec{X})) \geq \frac{1}{I(\theta)}$$

If $W(\vec{X})$ is unbiased, its variance is bounded by the inverse of the variance of the score

Cramer-Rao Lower Bound

Proof: Recall $\text{Cov}(U, V)^2 \leq \text{Var}(U) \text{Var}(V)$
and $\frac{d}{d\theta} \log(f(x)) = \frac{f'(x)}{f(x)}$.

Let $\vec{X}$ be an iid r.v. from a distribution with parameter θ and let $W(\vec{X})$ be an estimator of θ .

$$\text{Let } U = W(\vec{X}), \quad V = \frac{d}{d\theta} \log(f(\theta | \vec{X})) \text{ the score}$$
$$\text{Cov}(U, V)^2 \leq \text{Var}(U) \text{Var}(V) \Rightarrow \text{Var}(W(\vec{X})) \geq \frac{\text{Cov}(W(\vec{X}), \frac{d}{d\theta} \log(f(\theta | \vec{X})))^2}{\text{Var}(\frac{d}{d\theta} \log(f(\theta | \vec{X})))}$$

$\text{Var}(\frac{d}{d\theta} \log(f(\theta | \vec{X}))) = I(\theta)$ by definition, and

$$\text{Cov}(W(\vec{X}), \frac{d}{d\theta} \log(f(\theta | \vec{X}))) = E[W(\vec{X}) \frac{d}{d\theta} \log(f(\theta | \vec{X}))] - \underbrace{E[W(\vec{X})]}_0 \underbrace{E[\frac{d}{d\theta} \log(f(\theta | \vec{X}))]}_0$$

$$= \int_{-\infty}^{\infty} W(\vec{x}) f(\vec{x} | \theta) \left(\frac{\partial}{\partial \theta} \log(f(\vec{x} | \theta)) \right) d\vec{x}$$

$$= \int_{-\infty}^{\infty} W(\vec{x}) \cancel{f(\vec{x} | \theta)} \left(\frac{\frac{\partial}{\partial \theta} f(\vec{x} | \theta)}{f(\vec{x} | \theta)} \right) d\vec{x} = \int_{-\infty}^{\infty} \frac{\partial}{\partial \theta} W(\vec{x}) f(\vec{x} | \theta) d\vec{x}$$

$$= \frac{\partial}{\partial \theta} \int_{-\infty}^{\infty} W(\vec{x}) f(\vec{x} | \theta) d\vec{x} = \frac{\partial}{\partial \theta} E[W(\vec{X})] \quad \text{a.e.d.}$$

Attainment & Lehman-Scheffé

Let $\vec{X}$ be an i.i.d. r.s. from $f(x|\theta)$, where $f(x|\theta)$ satisfies the conditions of CRLB. If $W(\vec{X})$ is any unbiased estimator of $\psi(\theta)$ (a function of θ , usually just $\psi(\theta) = \theta$) then $W(\vec{X})$ attains CRLB for unbiased estimators if and only if:

$$\exists a(x) \text{ s.t. } a(\theta)[W(\vec{X}) - \psi(\theta)] = \frac{\partial}{\partial \theta} \ln(\theta|\vec{X})$$

Lehman-Scheffé Theorem:

Again, let $\vec{X}$ be an i.i.d. r.s. from a distribution with unknown parameter θ . Suppose $T(\vec{X})$ is an estimator of $\psi(\theta)$, and suppose $T(\vec{X})$ is a function of a complete and sufficient statistic for θ and is unbiased for $\psi(\theta)$ ($E[T(\vec{X})] = \psi(\theta)$). Then $T(\vec{X})$ is UMVUE for $\psi(\theta)$.

Further: $T(\vec{X})$ will also be a function of $\hat{\theta}_{MLE}$.
Hence: Unbiased + Complete + Sufficient = UMVUE

Rao-Blackwell & Unbiased Theorem

Let $\vec{X}$ be i.i.d. rs from dist w/ unknown θ .

Let W be any unbiased estimator of $\tau(\theta)$, and let $T(\vec{X})$ be a sufficient statistic for θ . Define $\phi(T(\vec{X})) = E[W(\vec{X}) | T(\vec{X})]$. Then $E[\phi(T(\vec{X}))] = \tau(\theta)$ and

$$\text{Var}(\phi(T(\vec{X}))) \leq \text{Var}(W(\vec{X})) \quad \forall \theta$$

Proof: Using Adam's rule and Eve's rule,
 $E[\phi(T(\vec{X}))] = E[E[W(\vec{X}) | T(\vec{X})]] = E[W(\vec{X})] = \tau(\theta)$

$$\begin{aligned} \text{Var}(W(\vec{X})) &= \text{Var}(E[W(\vec{X}) | T(\vec{X})]) + E[\text{Var}(W(\vec{X}) | T(\vec{X}))] \\ &= \text{Var}(\phi(T(\vec{X}))) + E[\text{Var}(W(\vec{X}) | T(\vec{X}))] \end{aligned}$$

$$\geq \text{Var}(\phi(T(\vec{X}))). \text{ This is the Rao-Blackwell Theorem.}$$

Unbiased Theorem

Suppose $W(\vec{X})$ is UMVUE. Then $W(\vec{X})$ is Unbiased.

Hypothesis Testing

A hypothesis is a statement, often about a population parameter of an assumed distribution

The two complementary hypotheses in a testing procedure are called the null hypothesis and the alternative hypothesis, denoted H_0 and H_1 .

If θ is the population parameter, the general format of the test is usually

$$H_0: \theta \in \Theta_0 \quad H_1: \theta \in \Theta_0^c$$

You can also use a hypothesis test to test the hypothesis that $X_1, X_2, \dots, X_n$ come from a particular distribution, e.g. Normal

$$H_0: X_1, \dots, X_n \sim \text{Normal} \quad H_a: X_1, \dots, X_n \not\sim \text{Normal}$$

You can also test other estimands of interest, e.g.

$$H_0: E[X] = 0 \quad H_a: E[X] \neq 0$$

Hypothesis Procedures

A hypothesis testing procedure is a rule that specifies for which sample values the decision is made to reject H_0 and accept H_1 as true and for which sample values this rejection of H_0 is not possible.

The subset of the sample space for which H_0 is rejected is the "rejection region" and the subset for which H_0 is not rejected is the "acceptance region".

Typically, a hypothesis test is specified in terms of a test statistic of the sample, $T(X_1, \dots, X_n)$

Types OF Error

A hypothesis test can make one of two kinds of error.

A type I error takes place when the null hypothesis is rejected even when it is actually true. False rejection.

A type II error takes place when the null hypothesis is accepted even when it is actually false. False Acceptance.

Since the null hypothesis doesn't depend on the data, the probability of making a type I for any particular test can be calculated explicitly without seeing the data.

For this reason, the goal of hypothesis testing is to minimize type II error while holding the type I error rate fixed at some level α , $0 < \alpha < 1$

Likelihood Ratio Tests

Likelihood Ratio Testing (LRTs) tests based on the ratio of the maximum likelihood under the null hypothesis to the overall maximum likelihood. The test stat is written

$$\lambda(x) = \frac{\sup_{\theta_0} L(\theta | x)}{\sup_{\theta} L(\theta | x)}$$

And the test must have a rejection region of the form $\{x: \lambda(x) \leq c\}$, $0 \leq c \leq 1$

Likelihood Ratio Test Example

Let $X_1, \dots, X_n$ be a r.s. from exponential population with pdf $f(x|\theta) = e^{-x/\theta}$ for $x \geq \theta$, $f(x|\theta) = 0$ $x < \theta$. Test $H_0: \theta \leq \theta_0$.

The likelihood function is thus

$$L(\theta|x) = \begin{cases} e^{-\sum x_i + n\theta} & \theta \leq X_{(1)} \\ 0 & \theta > X_{(1)} \end{cases}$$

Since L is an increasing function of θ , the denominator of the LRT uses the highest possible value of θ , $\theta = X_{(1)}$. The numerator will also be as high as possible, which is either $X_{(1)}$ or θ_0 . So

$$\lambda(x) = \begin{cases} 1 & X_{(1)} \leq \theta_0 \\ \frac{e^{-\sum x_i + n\theta_0}}{e^{-\sum x_i + nX_{(1)}}} & X_{(1)} > \theta_0 \end{cases}$$

When $\theta_0 > X_{(1)}$, the null $\theta \leq \theta_0$ includes all θ less than $X_{(1)}$, so it includes all possible θ , so immediately fail to reject. Reject when θ_0 is too small.

Power

The probability of making a type 2 error under the alternative hypothesis usually necessitates a particular assumption of an "alternative distribution"

If we are testing $\mu = \mu_0$ for $X_1, \dots, X_n$ iid $N(\mu, 1)$, the type II error is substantially different if $\mu = \mu_0 + 1$ vs $\mu = \mu_0 + 10$.

The power function of a hypothesis test with rejection region R is the function of θ defined by

$$B(\theta) = P_{\theta}(X \in R) = \begin{cases} \text{Prob of type I error} & \theta \in \theta_0 \\ 1 - \text{Prob of type II error} & \theta \in \theta_0^c \end{cases}$$

Ideally, a test should have probability near 0 (e.g. 0.05) $\forall \theta \in \theta_0$ and probability near 1 $\forall \theta \in \theta_0^c$

Neyman-Pearson Lemma

If C is a class of tests for $H_0: \theta \in \Theta_0$ vs $\theta \in \Theta_0^c$, a test in C with power function $\beta(\theta)$ is a uniformly most powerful test in C if $\beta(\theta) \geq \beta'(\theta) \forall \theta \in \Theta_0^c, \forall \beta' \in C$

Consider testing $H_0: \theta = \theta_0$ vs $H_1: \theta = \theta_1$. Where the pdf or pmf corresponding to θ_i is $f(x|\theta_i)$, $i=0,1$, using a test with rejection region R satisfying

$$\begin{aligned} \text{a. } & X \in R \text{ if } f(x|\theta_1) > K f(x|\theta_0) \text{ and} \\ & X \in R^c \text{ if } f(x|\theta_1) < K f(x|\theta_0) \text{ for} \\ & \text{some } K \geq 0, \text{ and } \text{b. } \alpha = P_{\theta_0}(X \in R) \end{aligned}$$

I. Any test that satisfies a & b is a Uniformly most powerful level α test

II. If there exists a test satisfying a & b with $K > 0$, then every UMP level α test is a size α test (satisfies b.), except perhaps $P_{\theta_0}(X \in A) = P_{\theta_1}(X \in A) = 0$

Sufficient Neyman-Pearson Lemma

Consider the hypothesis problem posed in the Neyman Pearson. Suppose $T(X)$ is a sufficient statistic for θ and $g(t|\theta_i)$ is the pdf or pmf of T corresponding to θ_i , $i=0,1$. Then any test based on T with rejection region S (a subset of the sample space of T) is a UMP level α test if it satisfies

$$t \in S \quad \text{if} \quad g(t|\theta_1) > k g(t|\theta_0)$$

and

$$t \in S^c \quad \text{if} \quad g(t|\theta_1) < k g(t|\theta_0)$$

for $k \geq 0$ where $\alpha = P_{\theta_0}(T \in S)$

i.e. A UMP level α test based on a sufficient statistic is just as good as one based on the whole sample

Karlin - Rubin Lemma

Consider testing $H_0: \theta \leq \theta_0$ vs $H_1: \theta > \theta_0$. Suppose T is a sufficient statistic for θ and the family of pdfs or pmfs $\{g(t|\theta): \theta \in \Theta\}$ of T has a monotone likelihood ratio.

That is, for every $\theta_2 > \theta_1 \in \Theta$ for $g(t|\theta): \theta \in \Theta$, the ratio $g(t|\theta_2)/g(t|\theta_1)$ is monotone (noninc or nondec) func of t on $t: g(t|\theta_1) > 0$ or $g(t|\theta_2) > 0$ with $\frac{c}{d}$ defined as ∞ if $c > 0$

Then for any t_0 the test that rejects H_0 if and only if $T > t_0$ is a UMP level α test where $\alpha = P_{\theta_0}(T > t_0)$

P-Values

A P-Value $P(X)$ is a test statistic s.t. $0 \leq P(X) \leq 1$ for every sample point X , and s.t. the type I error of the test $P(X) \leq \alpha$ is less than or equal to α .

That is, a valid P-value has $\forall \alpha, 0 \leq \alpha \leq 1$

$$P_0(P(X) \leq \alpha) \leq \alpha$$

An easy way to construct a level α test based on $P(X)$ is $P(X) \leq \alpha$. Since a P-value can be anywhere from 0 to 1, it gives more information than "Accept" or "Reject".

If $W(X)$ is a test statistic s.t. large values of W give evidence H_1 is true, then let

$$P(X) = \sup_{\theta \in \Theta_0} P_{\theta}(W(X) \geq W(x))$$

Probability of a

Random Var

Realizations

result "as significant or more significant" than the one you observed

Confidence Interval

An Interval Estimate of a real-valued parameter θ is any pair of functions $L(X_1, \dots, X_n)$ and $U(X_1, \dots, X_n)$ of a sample that satisfy $L(\vec{x}) \leq U(\vec{x}) \quad \forall x$. The inference $L(\vec{x}) \leq \theta \leq U(\vec{x})$, written $[L(x), U(x)]$ is an interval estimate.

The most common interval estimate is known as a Confidence Interval, and is associated with both a test and a sig. level α . The CI or Confidence Set for a set of observations $\vec{X}$ is the set of θ_0 's for which the level α test of θ_0 would NOT result in a rejection of the null hypothesis. This is also called a $100\alpha\%$ CI.

In other words, a $100\alpha\%$ CI is the set of "plausible" null hypothesized θ_0 s would it result in the null being rejected.

Pivot

Often, a 100% CI can be found using a pivot. A pivot is a random variable whose distribution does not depend on any parameters.

Many distributions have the location-scale type, and these have simple pivots.

Form of PDF	Type of PDF	Pivot
$f(x-u)$	Location	$\bar{X} - u$
$\frac{1}{\sigma} f\left(\frac{x}{\sigma}\right)$	Scale	$\bar{X}/\sigma$
$\frac{1}{\sigma} f\left(\frac{x-u}{\sigma}\right)$	Location-Scale	$\frac{\bar{X} - u}{\sigma}$

The Cdf can be found by putting the $\frac{\alpha}{2}$ and $\frac{1-\alpha}{2}$ % quantiles of the pivot distribution on each side and solving for the parameter(s).

For example, $X_1, \dots, X_n \sim N(u, 4)$, $\bar{X} = 10$, $\alpha = 0.05$, pivot distribution is $N(0, 1)$, so

$$-1.96 \leq \frac{10 - u}{2/\sqrt{n}} \leq 1.96 \Rightarrow \frac{1.96 \cdot 2}{\sqrt{n}} + 10 \geq u \geq \frac{-1.96 \cdot 2}{\sqrt{n}} + 10$$

Unbiasedness & Length

In the previous notes, we have split α evenly between both sides of the distribution of the test statistic T , but this is not always optimal (although it often is).

One way to define an 'optimal' interval is through its length. For example, a 95% interval that is $(1, 2)$ is probably more helpful than one that is $(1.2, 4)$.

Finding the interval of minimum length depends on the observed sample. Instead, we often minimize the Expected Length, $E[U(\vec{X})] - E[L(\vec{X})]$.

Sometimes, it is more convenient to simply find the shortest pivotal interval rather than the shortest interval.

A confidence set is said to be unbiased if

$$P_{\theta}(\theta' \in C(\vec{X})) \leq 1 - \alpha \quad \forall \theta \neq \theta'$$

Pivoting a CDF

Another common pivot, other than location-scale, is to pivot using the CDF. Let T be a statistic with continuous CDF $F_T(t|\theta)$. If $F_T(t|\theta)$ is decreasing function of $\theta \forall t$, solve for θ for

$$\text{Upper: } F_T(t|\theta_U(t)) = \frac{\alpha}{2}, \text{ Lower: } F_T(t|\theta_L(t)) = 1 - \frac{\alpha}{2}$$

If $F_T(t|\theta)$ is an increasing func of $\theta \forall t$, do

$$\text{Upper: } F_T(t|\theta_U(t)) = 1 - \frac{\alpha}{2}, \text{ Lower: } F_T(t|\theta_L(t)) = \frac{\alpha}{2}$$

Central Limit Theorem

Suppose $X_1, \dots$ is a sequence of i.i.d. random variables with $E[X] = \mu$ and $\text{Var}(X) = \sigma^2 < \infty$.

Then $\sqrt{n} \left(\frac{\bar{X}_n - \mu}{\sigma} \right) \xrightarrow{d} \text{Normal}(0, 1)$

Convergence in distribution: $X_1, \dots$ converges in distribution to X , $X_n \xrightarrow{d} X$, if

$$\lim_{n \rightarrow \infty} F_{X_n}(x) = F_X(x) \quad \forall x \text{ s.t. } F_X(x) \text{ is cont.}$$

Proof: First, let us rewrite the theorem explicitly, for $X_1, X_2, \dots$ iid $\text{Var} < \infty$, WLOG replace each r.v. $X_1, \dots$ with $\frac{X_1 - \mu}{\sqrt{\sigma^2}}, \frac{X_2 - \mu}{\sqrt{\sigma^2}}, \dots$, call these $Y_1, \dots$. Obviously, $Y_1, Y_2, \dots$, have mean zero and variance 1. CLT says

$$\sqrt{n} \bar{Y}_n = \frac{\sum_{i=1}^n Y_i}{\sqrt{n}} \xrightarrow{d} N(0, 1), \quad \text{that is}$$

$$\lim_{n \rightarrow \infty} F_{\sum_{i=1}^n Y_i / \sqrt{n}}(y) = \int_{-\infty}^y \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx \quad (1)$$

Proof of Central Limit Theorem

Recall from the MGF notes that the normal satisfies $X_1 \sim N + X_2 \sim N \Rightarrow X_1 + X_2 \sim N(\mu_1 + \mu_2, \sigma_1^2 + \sigma_2^2)$ for independent X_1, X_2

Recall also that the MGF uniquely identifies a distribution (two rvs with the same MGF have the same distribution). Recall that the normal has MGF $\exp(\mu t + \frac{\sigma^2}{2} t^2)$, where $M_X(t) = E[e^{tx}] = \int_{-\infty}^{\infty} e^{tx} f(x) dx$. Thus, $N(0,1)$ has MGF $\exp(\frac{t^2}{2})$. As a result, we can show (1) by showing (2) below:

$$\lim_{n \rightarrow \infty} M_{\sum_{i=1}^n Y_i / \sqrt{n}}(t) = \exp\left(\frac{t^2}{2}\right)$$

Call $\sum_{i=1}^n Y_i$ to be Z_n . By properties of MGFs (summing indep vars multiplies MGFs) it must be true that

$$M_{\sum_{i=1}^n Y_i}(t) = (M_Y(t))^n \Rightarrow M_{Z_n}(t) = \left(M_Y\left(\frac{t}{\sqrt{n}}\right)\right)^n$$

Proof of Central Limit Theorem

$M_Y(t)$ has a Taylor series expansion about zero given by

$$M_Y(t) = M_Y(0) + t M_Y'(0) + \frac{1}{2} t^2 M_Y''(0) + \varepsilon$$

Where $\frac{\varepsilon}{t^2}$ has a t and goes to zero as $t \rightarrow 0$.

Since $E[Y] = 0$, $M_Y'(0) = 0$, and since $E[Y^2] = \text{Var}(Y) + E[Y]^2 = \text{Var}(Y) = 1$, so as $n \rightarrow \infty$, $\frac{t}{\sqrt{n}} \rightarrow 0$, thus

$$\lim_{n \rightarrow \infty} M_{Z_n}(t) = \lim_{n \rightarrow \infty} \left(M_Y\left(\frac{t}{\sqrt{n}}\right) \right)^n =$$

$$M(0) = E[X^0] = 1 \quad t M'(0) = t E[X] = 0 \quad \frac{1}{2} t^2 M_Y''(0)$$

$$\lim_{n \rightarrow \infty} \left(\downarrow 1 + \downarrow 0 + \frac{1}{2} \left(\frac{t}{\sqrt{n}}\right)^2 \cdot 1 + \varepsilon \right)^n$$

$$= \lim_{n \rightarrow \infty} \left(1 + \frac{t^2}{2n} + \varepsilon \right)^n = \lim_{n \rightarrow \infty} \left(1 + \frac{t^2/2}{n} + \varepsilon \right)^n$$

Where all terms in ε are $t^3 = \frac{t^3}{n^{3/2}}$ or greater and disappear in the limit. Thus,

Proof of CLT

$$\lim_{n \rightarrow \infty} \left(1 + \frac{t^2/2}{n} + \varepsilon\right)^n = \lim_{n \rightarrow \infty} \left(1 + \frac{t^2/2}{n}\right)^n \quad \text{and}$$

by definition of e^x this is

$$\lim_{n \rightarrow \infty} \left(1 + \frac{t^2/2}{n}\right)^n = e^{t^2/2}$$

Thus $M_{Z_n(t)} \rightarrow e^{t^2/2}$ as $n \rightarrow \infty$

Thus a set of r.v.s $Y_1, \dots, Y_n$ that are standardized $\left(\frac{X - E[X]}{\sqrt{\text{Var}(X)}}\right)$ and each X has $\text{Var}(X) < \infty$, it is true that $\bar{Y} \xrightarrow{d} N(0, 1)$, since

$$M_{\bar{Y}}(t) \rightarrow M_{N(0,1)}(t), \text{ and}$$

Convergence of MGF is equivalent to convergence in distribution, since both MGF and distribution uniquely identify a probability distribution.

Lemma for CLT Proof

Let X be a standard normally distributed r.v., that is, it has pdf $f(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}x^2}$

$$\begin{aligned}\text{Then the MGF is given by } E[e^{tx}] \\&= \int_{-\infty}^{\infty} e^{tx} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}x^2} dx = \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x^2 - 2tx)} dx \\&= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x^2 - 2tx + t^2 - t^2)} dx = \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}((x-t)^2 - t^2)} dx \\&= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-t)^2} e^{t^2/2} dx = e^{t^2/2} \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-t)^2} dx\end{aligned}$$

Since inside the integral is the $\text{Normal}(t, 1)$ pdf, it integrates to 1, so as a result we have

$$e^{t^2/2} \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}(x-t)^2} dx = e^{t^2/2}$$

Thus standard normal has MGF $e^{t^2/2}$

Wald & Score Interval

In large samples, any distribution of finite variance has mean approaching a normal, $\bar{X} \rightarrow N(\mu, \frac{\sigma^2}{n})$. Wald and Score intervals are based on this large sample approximation.

Let $\hat{\theta}$ be a point estimate of θ based on the mean. Since $\hat{\theta}$ is a statistic and is based on the mean, as $n \rightarrow \infty$ $\hat{\theta} \rightarrow N(\theta, \text{Var}(\hat{\theta}))$. We want to construct the Wald interval as

$$\hat{\theta} \pm \overset{\substack{\uparrow \text{CDF of normal dist at } \frac{\alpha}{2}, \text{ eg } -1.96 \text{ for } \alpha = 0.05}}{Z_{\alpha/2}} \text{SE}(\hat{\theta})$$

Where Standard Error (SE) is just the Standard Deviation of the estimator $\hat{\theta}$, for $\bar{X}$ is $\sqrt{\frac{\text{Var}(\bar{X})}{n}} = \frac{\text{SD}(\bar{X})}{\sqrt{n}}$. For Wald, the $\text{SD}(\hat{\theta})$ is estimated too, leading to the below Wald Interval

$$\hat{\theta} \pm Z_{\alpha/2} \widehat{\text{SE}}(\hat{\theta})$$

Wald Score Interval

The Score interval uses the same normal approximation, but is "smart" enough to use the true variance rather than the approximation based on $\hat{\theta}$. This is done by making a normal approximation to the score function, $\frac{\partial}{\partial \theta} \ell(\theta | \vec{x})$. This can be done because, for a set of i.i.d. r.v.s $X_1, \dots, X_n$, the joint pdf is $\prod (f(x_i | \theta))^n$, and thus the log likelihood is $\sum \log(f(x_i | \theta))$ and the derivative is $\frac{\partial}{\partial \theta} \sum \log(L(\theta | x_i)) = \sum \frac{\partial}{\partial \theta} \log(L(\theta | x_i))$. So the joint score function $\frac{\partial}{\partial \theta} \ell(\theta | \vec{x})$ is the sum of i.i.d. r.v.s $\frac{\partial}{\partial \theta} \log(L(\theta | x_i))$ and thus the CLT applies. The Score interval is given by solving

$$\left(\frac{\partial}{\partial \theta} \ell(\hat{\theta} | \vec{x}) \right) / \sqrt{I(\theta)} = \pm Z_{\alpha/2} \quad \text{for } \theta$$

Recalling that $I(\theta)$ is just $\text{Var}\left(\frac{\partial}{\partial \theta} \ell(\theta | \vec{x})\right)$.
Solve individually for upper and lower intervals.

Wald & Score Example

let $X_1, \dots, X_n$ be iid bernoulli. Find Wald and score intervals for p .

$$P(X=x) = p^x (1-p)^{1-x} \text{ for } X=0 \text{ or } X=1$$

$$\text{Wald: } \hat{p} \pm Z_{\alpha/2} \widehat{SE}(\hat{p})$$

$\hat{p}$ is just $\frac{\sum X}{n} = \bar{X}$, and $SD[\bar{X}]$ is known to be $\frac{\sqrt{p(1-p)}}{\sqrt{n}}$, so $SE(\bar{X}) = \frac{SD(\bar{X})}{\sqrt{n}} = \frac{\sqrt{p(1-p)}}{\sqrt{n}}$, filling $\hat{p}$ for Wald interval is

$$\bar{X} \pm Z_{\alpha/2} \sqrt{\frac{\bar{X}(1-\bar{X})}{n}}$$

Score: $\frac{\partial}{\partial \theta} \ell(\theta | \vec{X}) = \pm Z_{\alpha/2}$, first, let's find score and info.

$$\begin{aligned} \frac{\partial}{\partial \theta} \log \left(\prod p^{x_i} (1-p)^{1-x_i} \right) &= \frac{\partial}{\partial p} \sum x_i \log(p) + (1-x_i) \log(1-p) \\ &= \frac{\sum x_i}{p} - \frac{n - \sum x_i}{1-p} = \frac{\bar{X}}{p} - \frac{n(1-\bar{X})}{1-p} \end{aligned}$$

Wald & Score Example

Fisher info $I(p)$ is also EL-second deriv, so

$$\begin{aligned} &= -E \left[\frac{\partial}{\partial p} \left(\frac{\sum x_i}{p} - \frac{n - \sum x_i}{1-p} \right) \right] = E \left[\frac{\sum x_i}{p^2} + \frac{n - \sum x_i}{(1-p)^2} \right] \\ &= \frac{np}{p^2} + \frac{n - np}{(1-p)^2} = \frac{n}{p} + \frac{n}{1-p} = \frac{n(1-p)}{p(1-p)} + \frac{np}{p(1-p)} = \frac{n}{p(1-p)} \end{aligned}$$

So score interval lower is

$$\frac{\frac{n\bar{x}}{p} - \frac{n(1-\bar{x})}{1-p}}{\sqrt{\frac{n}{p(1-p)}}} = Z_{\alpha/2} \Rightarrow \frac{\frac{n\bar{x}(1-p) - n(1-\bar{x})p}{p(1-p)}}{\sqrt{\frac{n}{p(1-p)}}} = Z_{\alpha/2}$$

$$\Rightarrow \frac{n(\bar{x} - \bar{x}p - p + \bar{x}p)}{\sqrt{n p(1-p)}} = Z_{\alpha/2} \Rightarrow Z_{\alpha/2} = \frac{n(\bar{x} - p)}{\sqrt{n p(1-p)}}$$

$$\Rightarrow \bar{x} - p = Z_{\alpha/2} \sqrt{\frac{p(1-p)}{n}} \Rightarrow p = \bar{x} - Z_{\alpha/2} \sqrt{\frac{p(1-p)}{n}}$$

$$\text{Thus, } p_L = \bar{x} - Z_{\alpha/2} \sqrt{\frac{p(1-p)}{n}}, \quad p_U = \bar{x} + Z_{\alpha/2} \sqrt{\frac{p(1-p)}{n}}$$

Since p is unknown, the score interval here requires an iterative/numerical method to find p_L and p_U

Extended Example

Many concepts from foregoing notes are illustrated in the below example for the normal distribution:

Suppose for all n that $X_1, \dots, X_n$ are normally distributed with unknown mean μ and known variance σ^2 .

a. Show that the sample mean $\bar{X} = \frac{\sum X_i}{n}$ and sample variance $S^2 = \frac{1}{n-1} \sum (X_i - \bar{X})^2$ are independent, without Basu's theorem

b. Show that $\bar{X}$ is UMVUE for μ

c. Derive the score & Fisher info of $\bar{X}$, and show that $\bar{X}$ achieves the CRLB

d. Find LRT of $\mu = \mu_0$ against $\mu \neq \mu_0$

e. Use NP-lemma to find the UMP test

f. Find a CI for μ based on pivot

Part A

Show that $\bar{X} \perp S^2$

Since the normal is a location-scale family, assume WLOG that $\mu=0$ $\sigma^2=1$

$$\begin{aligned} S^2 &= \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2 = \frac{1}{n-1} \left((x_1 - \bar{X})^2 + \sum_{i=2}^n (x_i - \bar{X})^2 \right) \\ &= \frac{1}{n-1} \left(\left(\sum_{i=2}^n (x_i - \bar{X}) \right)^2 + \sum_{i=2}^n (x_i - \bar{X})^2 \right) \end{aligned}$$

Thus S^2 can be written as a function only of $x_2 - \bar{X}, x_3 - \bar{X}, \dots$, and we can show these are independent of $\bar{X}$.

The joint pdf of $x_1, \dots, x_n$ is given by $f(x_1, \dots, x_n) = \frac{1}{(2\pi)^{n/2}} e^{-\frac{1}{2} \sum_{i=1}^n x_i^2}$, $-\infty < x_i < \infty$

We transform $y_1 = \bar{X}$, $y_2 = x_2 - \bar{X}$, $\dots$, $y_n = x_n - \bar{X}$

Since the transformation is linear, we just need to write $\sum_{i=1}^n x_i^2$ in terms of y 's and take the pdf.

Part A

$$\begin{aligned} \sum x_i^2 &= \left(x_1 - \sum_{i=2}^n y_i \right)^2 + \sum_{i=2}^n (y_i + y_i)^2 \\ &\quad \uparrow \qquad \qquad \qquad \uparrow \\ &= \left(\bar{x} - (x_2 - \bar{x}) - (x_3 - \bar{x}) - \dots \right)^2 + \left((x_2 - \bar{x} + \bar{x})^2 + (x_3 - \bar{x} + \bar{x})^2 + \dots \right) \\ &\quad \uparrow \qquad \qquad \qquad \uparrow \\ &= \left(n\bar{x} - \sum_{i=2}^n x_i \right)^2 + (x_2^2 + x_3^2 + x_4^2 + \dots) \\ &\quad \uparrow \qquad \qquad \qquad \uparrow \\ &\quad x_1^2 + \underbrace{x_2^2 + x_3^2 + \dots + x_n^2} \end{aligned}$$

the pdf can then be rewritten

$$\begin{aligned} &\frac{n}{(2\pi)^{n/2}} e^{-\frac{1}{2} \left((x_1 - \sum_{i=2}^n y_i)^2 + \sum_{i=2}^n (y_i + y_i)^2 \right)} \\ &= \frac{n}{(2\pi)^{n/2}} e^{-\frac{1}{2} (x_1 - \sum_{i=2}^n y_i)^2} e^{-\frac{1}{2} \sum_{i=2}^n (y_i + y_i)^2} \\ &= \frac{n}{2\pi^{n/2}} e^{-\frac{1}{2} (x_1^2 - \sum_{i=2}^n y_i y_i - (\sum_{i=2}^n y_i)^2)} e^{-\frac{1}{2} (\sum_{i=2}^n y_i^2 + \sum_{i=2}^n y_i y_i + \sum_{i=2}^n y_i^2)} \\ &= \frac{n}{2\pi^{n/2}} e^{-\frac{1}{2} n x_1^2} e^{-\frac{1}{2} \left(\sum_{i=2}^n y_i^2 - (\sum_{i=2}^n y_i)^2 \right)} \end{aligned}$$

Thus $Y_1 = \bar{X}$ is independent of $Y_2, \dots, Y_n$, and S^2 can be written as a function only of $Y_2, \dots, Y_n$.
Therefore $\bar{X}$ and S^2 are independent.

Part B

Show that $\bar{X}$ is UMVUE of μ

Step 1. Show that $\bar{X}$ is unbiased

$$\begin{aligned} E[\bar{X}] &= E\left[\frac{X_1 + X_2 + \dots + X_n}{n}\right] = E\left[\frac{X_1}{n}\right] + \dots + E\left[\frac{X_n}{n}\right] \\ &= \frac{\mu}{n} + \dots + \frac{\mu}{n} = n\left(\frac{\mu}{n}\right) = \mu \end{aligned}$$

Step 2. Show that $\bar{X}$ is sufficient

$$\begin{aligned} f(\vec{X} | \theta) &= \prod_{i=1}^n f(X_i | \theta) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(X_i - \mu)^2}{2\sigma^2}} \\ &= \prod_{i=1}^n \frac{1}{\sqrt{2\pi\sigma^2}} e^{\frac{-X_i^2 + 2\mu X_i - \mu^2}{2\sigma^2}} = \left(\frac{1}{\sqrt{2\pi\sigma^2}}\right)^n e^{\sum_{i=1}^n \frac{-X_i^2}{2\sigma^2} + \frac{2\mu X_i}{2\sigma^2} - \frac{\mu^2}{2\sigma^2}} \\ &= \left(\frac{1}{\sqrt{2\pi\sigma^2}}\right)^n e^{\sum_{i=1}^n \frac{-X_i^2}{2\sigma^2}} e^{\sum_{i=1}^n \frac{\mu X_i}{\sigma^2} - \frac{n\mu^2}{2\sigma^2}} \\ &= \underbrace{\left(\frac{1}{\sqrt{2\pi\sigma^2}}\right)^n e^{\sum_{i=1}^n \frac{-X_i^2}{2\sigma^2}}}_{h(\vec{X})} e^{\underbrace{\frac{n\mu\bar{X} - n\mu^2}{2\sigma^2}}_{g(\mu | \bar{X})}} \end{aligned}$$

Part B

Step 3. Show $\bar{X}$ is complete for μ

Exp Family: $f(X|\theta) = h(x) c(\theta) \exp\left(\sum_{j=1}^k w_j(\theta) t_j(x)\right)$

For parameters $\theta_1, \dots, \theta_k$. In this case, σ^2 is known. We only are concerned with $\theta_1 = \mu$, and the statistic $\sum X_i$ is sufficient (and therefore $\bar{X}$ is sufficient) if $f(X|\mu)$ is in the Exp Family

$$f(X|\mu) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

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

$$= \underbrace{\frac{1}{\sqrt{2\pi\sigma^2}} e^{\frac{x^2}{\sigma^2}}}_{h(x)} \underbrace{e^{\frac{\mu^2}{\sigma^2}}}_{c(\mu)} \exp\left(-\underbrace{\mu}_{w(\mu)} \cdot \underbrace{\frac{x}{\sigma^2}}_{t(x)}\right)$$

$w(\mu) = -\mu$ which contains open sets in $\mathbb{R}$ (it's not just a constant)

Therefore $\sum X_i$ (& $\bar{X}$) are suff for $\bar{X}$
thus $\bar{X}$ is complete & unbiased so $\bar{X}$ is UMVUE

Part C

First, Derive Score eq & Fisher info

$$L(u | \vec{X}) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}\sigma^2} \exp\left(-\frac{(x_i - u)^2}{2\sigma^2}\right)$$

$$\ell(u | \vec{X}) = n \log\left(\frac{1}{\sqrt{2\pi}\sigma^2}\right) + \sum_{i=1}^n \frac{(x_i - u)^2}{2\sigma^2}$$

$$\frac{d}{du} \ell(u | \vec{X}) = - \sum_{i=1}^n \frac{(x_i - u)}{\sigma^2} = \sum_{i=1}^n \frac{(u - x_i)}{\sigma^2}$$

$$E\left[\left(\frac{d}{du} \ell(u | \vec{X})\right)^2\right] = E\left[\sum_{i=1}^n \frac{(x_i - u)^2}{\sigma^4}\right]$$

$$= \frac{1}{\sigma^4} E\left[\sum_{i=1}^n (x_i - u)^2\right] = \frac{1}{\sigma^4} n \sigma^2 = \frac{n}{\sigma^2}$$

Thus, $I(u) = \frac{n}{\sigma^2}$. Thus CRLB is $\frac{1}{I(u)} = \frac{1}{n/\sigma^2} = \frac{\sigma^2}{n}$

$$\text{Var}(\bar{X}) = \text{Var}\left(\frac{x_1 + x_2 + \dots + x_n}{n}\right) = \frac{1}{n^2} \text{Var}(x_1 + \dots + x_n)$$

Since $x_1 + \dots + x_n$ is the sum of independent normals,

$$= \frac{1}{n^2} \text{Var}\left(N(nu, n\sigma^2)\right) = \frac{n\sigma^2}{n^2} = \frac{\sigma^2}{n}$$

Thus $\bar{X}$ achieves the CRLB for u

PART D

LRT $u = u_0$ vs $u \neq u_0$

LRT format: Reject if $C \geq \frac{L_{H_0}}{L_{H_a}}$

$$\frac{\prod \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x_i - u_0)^2}{2\sigma^2}\right)}{\prod \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x_i - u_{MLE})^2}{2\sigma^2}\right)} = \exp\left(\frac{1}{2\sigma^2} \sum_{i=1}^n \left((x_i - u_{MLE})^2 - (x_i - u_0)^2\right)\right)$$

Now, we just need to derive u_{MLE} , LUCKILY, this is easy to do by simply setting the previously derived score equation to zero.

$$\frac{d}{du} \ell(u | \vec{X}) = \sum_{i=1}^n \frac{(x_i - u)}{\sigma^2} \stackrel{\text{set}}{=} 0$$

$$\Rightarrow \sum_{i=1}^n (x_i - u_{MLE}) = 0 \Rightarrow \left(\sum_{i=1}^n x_i\right) - n u_{MLE} = 0$$

$$\Rightarrow \sum_{i=1}^n x_i = n u_{MLE} \Rightarrow u_{MLE} = \frac{\sum_{i=1}^n x_i}{n} \Rightarrow u_{MLE} = \bar{X}, \text{ so}$$

the LRT is just

$$\text{Reject if } C \geq \exp\left(\frac{1}{2\sigma^2} \sum_{i=1}^n \left((x_i - \bar{X})^2 - (x_i - u_0)^2\right)\right)$$

Part E

Step 1. Show the LRT is VMP

Use Suff. NPLemma: a test based on a sufficient statistic is

UMP if for pdf of S

$$t \in S \quad \text{if} \quad g(t|\theta_1) > k g(t|\theta_0)$$

and

$$t \in S^c \quad \text{if} \quad g(t|\theta_1) < k g(t|\theta_0)$$

for $k \geq 0$ where $\alpha = P_{\theta_0}(T \in S)$

Since $\bar{X}$ is normal with mean μ_0/μ_{MLE} variance $\frac{\sigma^2}{n}$, it has pdf under the null

$$g(t|\mu_0) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(t-\mu_0)^2}{2\sigma^2}\right)$$

$$g(t|\mu_{MLE}) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(t-\bar{X})^2}{2\sigma^2}\right)$$

The test is then

Part E

$$t \in S \quad \text{if} \quad \exp\left(-\frac{(t-\bar{X})^2}{2\sigma^2} - \frac{(t-u_0)^2}{2\sigma^2}\right) > k$$

and

$$t \in S^c \quad \text{if} \quad \exp\left(-\frac{(t-\bar{X})^2}{2\sigma^2} - \frac{(t-u_0)^2}{2\sigma^2}\right) < k$$

$$\text{for } k \geq 0 \quad \text{where} \quad \alpha = P_{\theta_0}(T \in S)$$

Since the GUFF statistic t is $\bar{X}$, this is

$$\bar{X} \in S \quad \text{if} \quad \exp\left(-\frac{(\bar{X}-u_0)^2}{2\sigma^2}\right) > k$$

and

$$\bar{X} \in S^c \quad \text{if} \quad \exp\left(-\frac{(\bar{X}-u_0)^2}{2\sigma^2}\right) < k$$

$$\text{for } k \geq 0 \quad \text{where} \quad \alpha = P_{\theta_0}(\bar{X} \in S)$$

Now, we just need to show that
 $\alpha = P_{u_0}(\bar{X} \in S)$

$$\begin{aligned} P_{u_0}(\bar{X} \in S) &= P_{u_0}\left(\exp\left(-\frac{(\bar{X}-u_0)^2}{2\sigma^2}\right) > k\right) \\ &= P_{u_0}\left(-\sqrt{2\sigma^2 \ln\left(\frac{1}{k}\right)} + u_0 \leq \bar{X}\right) = P_{u_0}\left(c^+ \leq \bar{X}, c^- \geq \bar{X}\right) \end{aligned}$$

for some c^+, c^-

Part E

The rejection probability under the null hypothesis $P(C^+ \leq \bar{X}, C^- \geq \bar{X})$ is known since $\bar{X}$ has normal distribution. So the UMP α -level test can be formulated as

$$\bar{X} \in R \quad \text{if} \quad C Z_{\frac{\alpha}{2}} \leq \bar{X} \leq C Z_{1-\frac{\alpha}{2}}$$

Part F

BC $\bar{X}$ is also normal, and σ^2 is known, a CI for μ can be found by pivoting the standard normal distribution.

$\frac{\bar{X} - \mu_0}{\sigma/\sqrt{n}} \sim Z$, so the 95% CI can be found with the CDF of Z , $Z_{\frac{\alpha}{2}}$, $Z_{1-\frac{\alpha}{2}}$

$$\bar{X} - \frac{\sigma}{\sqrt{n}} Z_{\frac{\alpha}{2}} \leq \mu \leq \bar{X} + \frac{\sigma}{\sqrt{n}} Z_{\frac{\alpha}{2}}$$