

Basics:

$$E(X) = \int x f(x) dx \quad \text{Var}(X) = E(X^2) - E(X)^2$$

$$E(f(X)) = \int f(x) E(X) dx = E[f(X) - E(X)]^2$$

$$E(h(X)) = \int h(x) f(x) dx \quad \text{Var}(f(X)) = E^2 \text{Var}(X)$$

$$\text{Cov}(X, Y) = E[(X - E(X))(Y - E(Y))], \text{corr}(X, Y) = \frac{\text{Cov}(X, Y)}{\sqrt{\text{Var}(X)\text{Var}(Y)}}$$

Basics: $f(\theta|x) = \frac{f(x|\theta)f(\theta)}{\int f(x|\theta)f(\theta)d\theta}$

Linear Regression

$$y = \beta_0 + \beta_1 x + \epsilon, E(\epsilon) = 0, \text{Var}(\epsilon) = \sigma^2$$

$$Y = XB + \epsilon, Y_i | X_i \sim N(\beta^T X_i, \sigma^2)$$

$$\beta = \arg \min_{\beta} \|Y - XB\|_2^2, \hat{\beta}_0 = \bar{y}, \hat{\beta}_1 = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2}$$

$$P(X_i | \beta) = \frac{1}{1 + e^{-\beta^T x_i}}$$

$$\text{MLE for } \beta: \ln L(\beta) = \sum \ln p(x_i | \beta)$$

$$= \sum \left[\left(\frac{1}{1 + e^{-\beta^T x_i}} \right)^{y_i} \left(\frac{e^{-\beta^T x_i}}{1 + e^{-\beta^T x_i}} \right)^{1 - y_i} \right]$$

$$L(\beta) = \prod p(x_i | \beta) = \prod \frac{e^{\beta^T x_i}}{1 + e^{\beta^T x_i}}$$

Markov's: $E(X) \geq t P(X \geq t)$

Delta Method

$$\frac{1}{\sqrt{n}}(X_n - \theta) \sim N(0, \sigma^2), \frac{d}{dx} f(x) \big|_{x=\theta} = f'(\theta)$$

Normal Equation: $\|Y - XB\|_2^2 = (Y - XB)^T (Y - XB)$

$$\hat{\beta} = (X^T X)^{-1} X^T Y, \hat{y} = X(X^T X)^{-1} X^T Y$$

MLE for $\beta: \ln L(\beta) = \sum \ln p(x_i | \beta)$

Large Numbers: $\bar{X}_n = \frac{1}{n} \sum X_i \rightarrow E(X)$

i.e. $\forall \epsilon, P(|\bar{X}_n - E(X)| > \epsilon) \rightarrow 0, n \rightarrow \infty$

Normal Distribution: $f(x) = \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}, X \sim N(\mu, \sigma^2)$

MLE for $\beta: \ln L(\beta) = \sum \ln p(x_i | \beta)$

MLE for $\beta: \ln L(\beta) = \sum \ln p(x_i | \beta)$

CDF: $F(t) = P(Z \leq t) = \int_{-\infty}^t \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{z^2}{2\sigma^2}} dz$

Efficiency: $\frac{\text{Var}(\hat{\theta})}{\text{Var}(\hat{\theta})} = \frac{1}{\text{Var}(\hat{\theta})}$

MLE for $\beta: \ln L(\beta) = \sum \ln p(x_i | \beta)$

MLE for $\beta: \ln L(\beta) = \sum \ln p(x_i | \beta)$

Central Limit Theorem: $\frac{1}{\sqrt{n}} \sum_{i=1}^n (X_i - \mu) \sim N(0, 1)$

Efficiency: $\frac{\text{Var}(\hat{\theta})}{\text{Var}(\hat{\theta})} = \frac{1}{\text{Var}(\hat{\theta})}$

MLE for $\beta: \ln L(\beta) = \sum \ln p(x_i | \beta)$

MLE for $\beta: \ln L(\beta) = \sum \ln p(x_i | \beta)$

Binomial: $X \sim B(n, p), P(X=k) = \binom{n}{k} p^k (1-p)^{n-k}, E(X) = np, \text{Var}(X) = np(1-p)$

Poisson: $X \sim \text{Pois}(\lambda), P(X=k) = \frac{\lambda^k e^{-\lambda}}{k!}, E(X) = \lambda, \text{Var}(X) = \lambda$

Poisson: $X \sim \text{Pois}(\lambda), P(X=k) = \frac{\lambda^k e^{-\lambda}}{k!}, E(X) = \lambda, \text{Var}(X) = \lambda$

Estimator $\hat{\theta}$ unbiased if $E(\hat{\theta}) = \theta$

Estimator $\hat{\theta}$ consistent if $\hat{\theta}_n \rightarrow \theta$ as $n \rightarrow \infty$

Estimator $\hat{\theta}$ consistent if $\hat{\theta}_n \rightarrow \theta$ as $n \rightarrow \infty$

Unbiased, inconsistent: $\hat{\theta} = X_1$

Unbiased, inconsistent: $\hat{\theta} = X_1$

Unbiased, inconsistent: $\hat{\theta} = X_1$

Consistent, unbiased: Sample Variance MLE, $\hat{\sigma}^2 = \frac{1}{n} \sum (X_i - \bar{X})^2$

Consistent, unbiased: Sample Variance MLE, $\hat{\sigma}^2 = \frac{1}{n} \sum (X_i - \bar{X})^2$

Consistent, unbiased: Sample Variance MLE, $\hat{\sigma}^2 = \frac{1}{n} \sum (X_i - \bar{X})^2$

MLE 1. Calculate joint likelihood: $L(\theta | x_1, x_2, \dots, x_n) = \prod_{i=1}^n f(x_i | \theta)$

MLE 1. Calculate joint likelihood: $L(\theta | x_1, x_2, \dots, x_n) = \prod_{i=1}^n f(x_i | \theta)$

MLE 1. Calculate joint likelihood: $L(\theta | x_1, x_2, \dots, x_n) = \prod_{i=1}^n f(x_i | \theta)$

2. Take log: $\ell(\theta | x_1, x_2, \dots, x_n) = \sum_{i=1}^n \ln f(x_i | \theta)$

2. Take log: $\ell(\theta | x_1, x_2, \dots, x_n) = \sum_{i=1}^n \ln f(x_i | \theta)$

2. Take log: $\ell(\theta | x_1, x_2, \dots, x_n) = \sum_{i=1}^n \ln f(x_i | \theta)$

3. Differentiate $\ell(\theta)$ w.r.t. θ : $\frac{d\ell}{d\theta} = 0$, solve

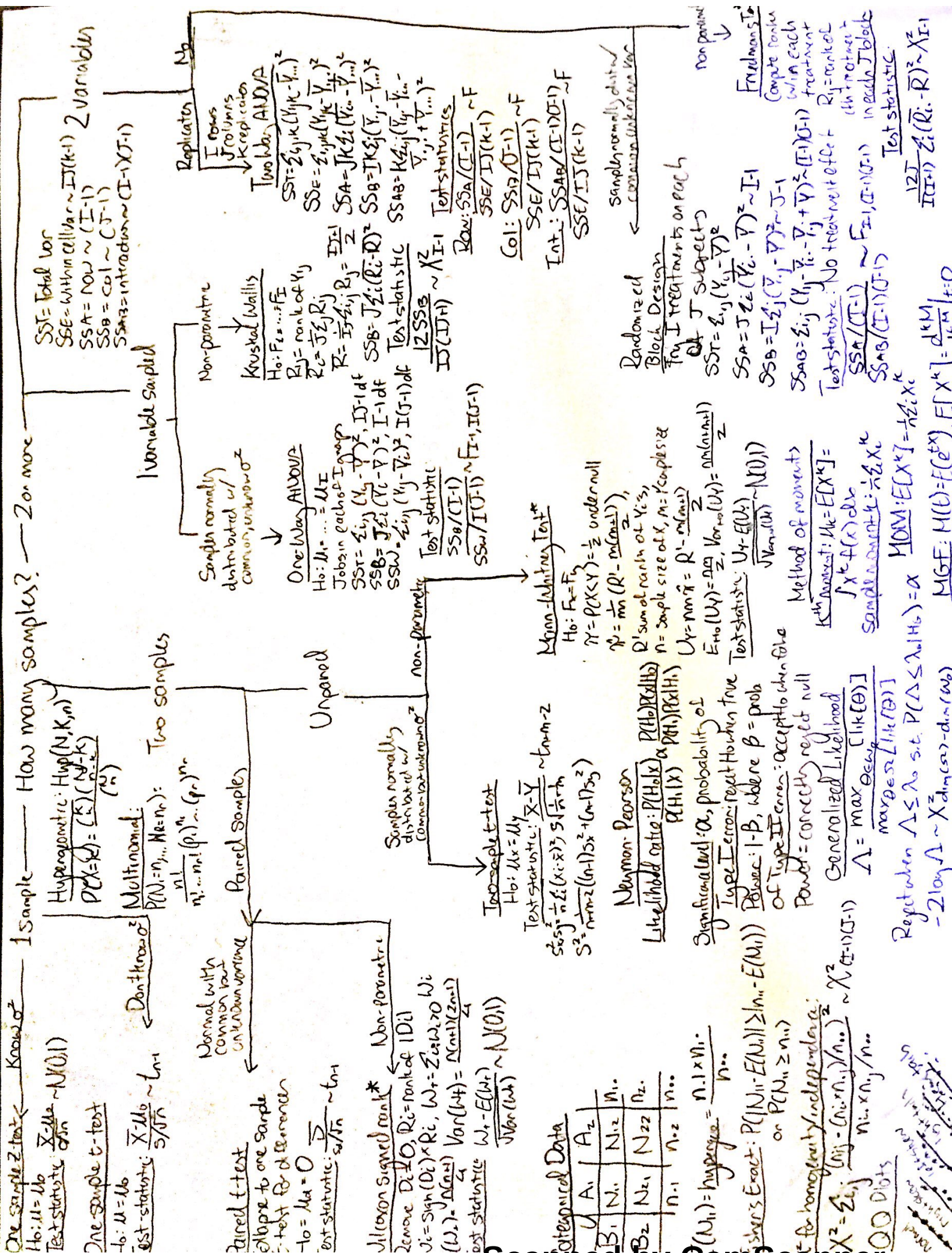
3. Differentiate $\ell(\theta)$ w.r.t. θ : $\frac{d\ell}{d\theta} = 0$, solve

3. Differentiate $\ell(\theta)$ w.r.t. θ : $\frac{d\ell}{d\theta} = 0$, solve

Confidence Interval: $q(\alpha)$ is α th quantile of Φ : $P(Z \leq q(\alpha)) = \alpha$

Confidence Interval: $q(\alpha)$ is α th quantile of Φ : $P(Z \leq q(\alpha)) = \alpha$

Confidence Interval: $q(\alpha)$ is α th quantile of Φ : $P(Z \leq q(\alpha)) = \alpha$



Expectation: $E(X) = \sum_{i=1}^n a_i P(X=a_i) = \int x f(x) dx$

Expectation is Linear: $E(c+X) = c + E(X)$; $E(cX) = cE(X)$

$$E(X+Y) = E(X) + E(Y); E(\sum_i X_i) = \sum_i E(X_i)$$

Expectation of function: $E(h(X)) = \int h(x) f(x) dx$

Independence: X, Y independent if: $P(X \leq x, Y \leq y) = P(X \leq x) P(Y \leq y)$

If independent: $E(f(X)g(Y)) = E(f(X)) E(g(Y))$; $E(XY) = E(X)E(Y)$

Variance: $\text{Var}(X) = E(X^2) - E(X)^2$; $\text{Var}(c+X) = \text{Var}(X)$; $\text{Var}(cX) = c^2 \text{Var}(X)$

If X, Y independent: $\text{Var}(X+Y) = \text{Var}(X) + \text{Var}(Y)$; $\text{Var}(\sum_i X_i) = \sum_i \text{Var}(X_i)$

Covariance between $X, Y = \text{cov}(X, Y) = E[(X - E(X))(Y - E(Y))]$

Correlation: $\text{corr}(X, Y) = \text{cov}(X, Y) / \sqrt{\text{var}(X)\text{var}(Y)}$

Markov's Inequality: If X non-neg. r.v., for $t \geq 0$, then
 $P(X \geq t) \leq E(X)/t$

Chebyshev's Inequality: $P(|X - E(X)| \geq t) \leq \text{Var}(X)/t^2$

Law of Large Numbers: $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i \rightarrow E(X_1)$ i.e.

$$\forall \epsilon > 0, P(|\bar{X}_n - E(X_1)| > \epsilon) \rightarrow 0$$

Normal Distribution: $f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$, $X \sim N(\mu, \sigma^2)$

$$Z \sim N(0, 1), X = \mu + \sigma Z, X \sim N(\mu, \sigma^2)$$

$$X_1 \sim N(\mu_1, \sigma_1^2), X_2 \sim N(\mu_2, \sigma_2^2), X_1, X_2 \text{ indep.}$$

$$X_1 + X_2 \sim N(\mu_1 + \mu_2, \sigma_1^2 + \sigma_2^2)$$

cdf: $\Phi(t) = P(Z \leq t) = \int_{-\infty}^t \frac{1}{\sqrt{2\pi}} e^{-x^2/2} dx$ of standard normal

Central Limit Theorem: $\sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \Rightarrow Z, Z \sim N(0, 1)$

$$P(\sqrt{n} \frac{\bar{X}_n - \mu}{\sigma} \leq t) \rightarrow \Phi(t)$$

Bernoulli: $X \sim B(p)$ if $X = \begin{cases} 1 & \text{w. prob. } p \\ 0 & \text{otherwise} \end{cases}$ $E(X) = p, \text{Var}(X) = p(1-p)$

Binomial: $X \sim \text{Bin}(n, p)$ if $P(X=k) = \binom{n}{k} p^k (1-p)^{n-k}$ $E(X) = np$

If $X \sim \text{Bin}(n, p)$, $X = \sum_{i=1}^n Y_i, Y_i \sim B(p)$ i.i.d. $\text{Var}(X) = np(1-p)$

Poisson: $X \sim \text{Pois}(\lambda), P(X=k) = \lambda^k \frac{e^{-\lambda}}{k!}, E(X) = \lambda, \text{Var}(X) = \lambda$

Estimator $\hat{\theta}$ is function of $X = (X_1, \dots, X_n)$ approximating parameter θ
 $\hat{\theta}$ unbiased if $E(\hat{\theta}) = \theta$. Risk of $\hat{\theta}$ is $R(\hat{\theta}) = E(|\hat{\theta} - \theta|^2)$
 Given $X_1, \dots, X_n$ i.i.d w. mean μ , var σ^2 : $E(\bar{X}) = \mu$, $E(\hat{\sigma}^2) = \sigma^2$
 Sample mean $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$, Sample var $\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$

Maximum Likelihood Estimator: Calculate single value of θ that maximizes likelihood function w.r. to θ

1. Calculate likelihood function $lik(\theta) = \prod_{i=1}^n f_{\theta}(X_i)$
2. Log-likelihood: $l(\theta) = \log(lik(\theta)) = \sum_{i=1}^n \log f_{\theta}(X_i)$
3. Differentiate $l(\theta)$ w.r. to θ , $\frac{dl}{d\theta}$
4. Set to 0, solve for θ

Confidence Interval: Interval containing true θ w. some coverage prob.

If var σ^2 of X_i known, $100(1-\alpha)\%$ CI for μ based on $\bar{X}_n$

is: $[\bar{X}_n - q(1-\alpha/2) \frac{\sigma}{\sqrt{n}}, \bar{X}_n + q(1-\alpha/2) \frac{\sigma}{\sqrt{n}}]$, where $q(\alpha)$ is α -th quantile of $\Phi = P(Z \leq q(\alpha)) = \alpha$.

When var σ^2 unknown, estimator is $\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2$

Then $\sqrt{n} \frac{\bar{X}_n - \mu}{\hat{\sigma}} \Rightarrow N(0,1)$. $100(1-\alpha)\%$ CI is now:

$[\bar{X}_n - q(1-\alpha/2) \frac{\hat{\sigma}}{\sqrt{n}}, \bar{X}_n + q(1-\alpha/2) \frac{\hat{\sigma}}{\sqrt{n}}]$

Consistent estimator: $\hat{\theta}_n \rightarrow \theta$ as $n \rightarrow \infty$. Estimator can be unbiased

w.o. being consistent ($\hat{\theta} = X_1$) or consistent w/o being unbiased ($\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2$)

Unbiased means expected value is parameter, consistent means approaches parameter $\vee$ more samples

Asymptotic MLE: Consistency: as $n \rightarrow \infty$, MLE approaches θ

Normality: as $n \rightarrow \infty$, distribution of MLE estimate tends to normal

$\hat{\theta}_{MLE} \sim N(\theta_0, \sigma_{\theta_0}^2)$. Fisher Information: $I(\theta_0) = E[(\frac{\partial}{\partial \theta} \log(f_{\theta}(x)))^2 | \theta_0]$

$= -E[\frac{\partial^2}{\partial \theta^2} \log(f_{\theta}(x)) | \theta_0]$. Asymptotic mean of MLE

$\theta_0 = \theta$ because $E(\hat{\theta}_{MLE}) \rightarrow E(\theta_0) = \theta$ (unbiased)

Asymptotic variance of MLE $\sigma_{\theta_0}^2 = \frac{1}{n I(\theta_0)}$. So $\hat{\theta}_{MLE} \sim N(\theta_0, \frac{1}{n I(\theta_0)})$

Method of Moments: k th moment $\mu_k = E[X^k] = g_k(\theta)$
 MOM equates theoretical moments to sample moments to solve
 Sample moment $k = \frac{1}{n} \sum_{i=1}^n X_i^k$. Mom: $E[X^k] = \frac{1}{n} \sum_{i=1}^n X_i^k$.

Cramer-Rao Lower Bound: If $\hat{\theta}$ unbiased est., $\text{Var}(\hat{\theta}) \geq \frac{1}{nI(\theta)}$
Efficiency of 2 estimators: $\text{eff}(\hat{\theta}, \tilde{\theta}) = \text{Var}(\tilde{\theta}) / \text{Var}(\hat{\theta})$.
 Unbiased estimator efficient if variance = Cramer Rao.
 MLE is asymptotically efficient.

Sufficiency: $X_1, \dots, X_n \sim F(\theta)$. $T(X_1, \dots, X_n)$ is sufficient statistic for θ if conditional dist. of $X_1, \dots, X_n$ given $T=t$ does not depend on θ for any $t = X_1, \dots, X_n | T(X_1, \dots, X_n) = t$, t not dependent on θ .
 Basically, knowledge of t is good on knowledge of sample.
Factorization Theorem: $f_\theta(x_1, \dots, x_n) = g_\theta(T)h(x_1, \dots, x_n)$ i.e. density can be factored into product s.t. h does not depend on θ , and g depends on θ depends on $(x_1, \dots, x_n)$ only through T .

Bias-Variance Tradeoff: Want to minimize bias (want $\hat{\theta}$ to be close to θ) and variance (don't want values to far from $E(\hat{\theta})$).

Mean Squared Error: $\text{MSE}(\hat{\theta}) = E[(\hat{\theta} - \theta)^2] = E[\hat{\theta}^2 - 2\theta\hat{\theta} + \theta^2] = E(\hat{\theta}^2) - 2\theta E(\hat{\theta}) + \theta^2 = \text{Var}(\hat{\theta}) + (E(\hat{\theta}) - \theta)^2 = \text{Var}(\hat{\theta}) + \text{Bias}(\hat{\theta})^2$.

Rao Blackwell: $\hat{\theta}$ estimator for θ , $E(\hat{\theta}^2) < \infty$, T sufficient for θ , $\tilde{\theta} = E(\hat{\theta} | T)$, then $\text{MSE}(\tilde{\theta}) \leq \text{MSE}(\hat{\theta})$.

Bayesian Inference:

Frequentist

Unknown params θ fixed but unknown
 Prob event is freq. it occurs if repeated
 Params fixed

Bayesian

Unknown param θ r.v. w/ prior dist.
 Prob event is degree belief it occurs
 Params fixed

Bayesian: Prior dist of θ , use data to "update knowledge" & get posterior
 $f(\theta | x) = \frac{f(x | \theta)f(\theta)}{\int f(x | \theta')f(\theta')d\theta'}$

Parametric Bootstrap: Estimate θ from original sample, generate samples $X_1^*, \dots, X_n^*$ from F_θ which approximates F_θ . (Assumed dist. of θ)

Nonparametric Bootstrap: Take samples $X_1, \dots, X_n$ w/ rep. from orig. sample. Once there are B bootstrap samples, generate B estimates of θ :
 $X_1^{*(1)}, \dots, X_n^{*(1)} \rightarrow \hat{\theta}^{*(1)}, \dots, X_1^{*(B)}, \dots, X_n^{*(B)} \rightarrow \hat{\theta}^{*(B)}$

Bootstrap CI: Normal: $\hat{\theta} \pm q(1-\alpha/2) \sqrt{\text{Var}(\hat{\theta})}$. Percentile: $(\hat{\theta}_{(1-\alpha/2)}, \hat{\theta}_{(1-\alpha/2)})$
Pivotal: $(2\hat{\theta} - \hat{\theta}_{(1-\alpha/2)}, 2\hat{\theta} - \hat{\theta}_{(\alpha/2)})$.

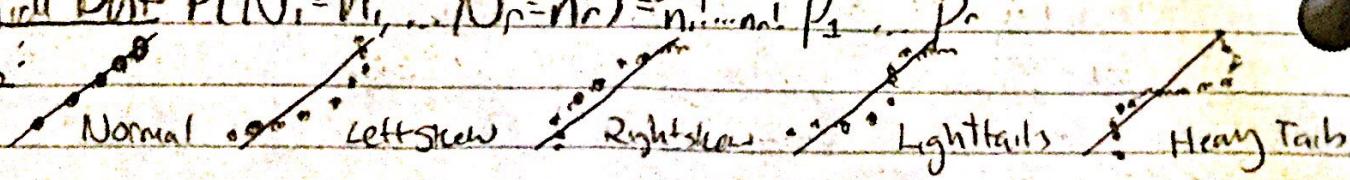
Hypothesis Testing: Neyman-Pearson: Make decision whether to reject H_0 in favor of H_1 based on statistic $T(X_1, \dots, X_n)$.
Acceptance Region: range of T where H_0 not rejected.
Rejection region: range of T s.t. null rejected.
Type I error: reject H_0 when true, probability = α $\rightarrow$ significance level
Type II error: accept H_0 when false, prob β . Power = $1 - \beta$.
Neyman Pearson approach: Fix α , use likelihood ratio $\frac{f_0(x)}{f_1(x)}$ to find test w/ highest power, where f_0 dist under H_0 , f_1 distribution under alternative, both simple (one value).
Neyman Pearson Lemma: Best test rejects H_0 for H_1 when $\frac{f_1(x)}{f_0(x)} < c(\alpha)$.

Generalized Likelihood Ratio Tests: If there is composite hypothesis: $\theta \in \omega_0$ under null, $\theta \in \omega_1$ for alternative. $\Omega = \omega_0 \cup \omega_1$. Then:
 $\Lambda = \max_{\theta \in \omega_0} \text{lik}(\theta) / \max_{\theta \in \Omega} \text{lik}(\theta)$. Reject when $\Lambda \leq \lambda_0$, such that $P(\Lambda \leq \lambda_0 | H_0) = \alpha$. Assume H_0 true, asymptotically:
 $2 \log(\Lambda) \sim \chi^2_{df}$, $df = \# \text{ free param under } \Omega - \# \text{ free param under } \omega_0$.
Thus if F_{χ^2} CDF of χ^2_{df} : $\alpha = P(\Lambda \leq \lambda_0) = P(-2 \log \Lambda \geq -2 \log \lambda_0) = F_{\chi^2_{df}}(-2 \log \lambda_0)$, so $\lambda_0 = \exp(-F_{\chi^2_{df}}^{-1}(1-\alpha))$, $F_{\chi^2_{df}}^{-1}(1-\alpha)$ is $(1-\alpha)$ quantile χ^2_{df} .

Pearson's χ^2 test: $\chi^2 = \sum_{i=1}^n \frac{(O_i - E_i)^2}{E_i} = \sum_{i=1}^n \frac{O_i^2}{E_i} - N$, O_i = observed # E_i expected.
Under null, n is # groups, N is sum of observed, $\chi^2 = \chi^2_{n - \dim(\theta)}$.

Multinomial Dist: $P(N_1 = n_1, \dots, N_r = n_r) = \frac{n!}{n_1! \dots n_r!} p_1^{n_1} \dots p_r^{n_r}$

QQ Plots:



Normal Left skew Right skew Light tails Heavy Tails

χ^2 distribution: $Z_i \sim N(0,1)$, $Z_1^2 + \dots + Z_n^2 \sim \chi_n^2$. $U \sim \chi_n^2$, $V \sim \chi_m^2$, $U+V \sim \chi_{n+m}^2$
 $U \sim \chi_n^2 \Leftrightarrow U \sim \text{Gamma}(\frac{n}{2}, \frac{1}{2})$ so $f(\chi_n^2) = \frac{1}{2^{n/2} \Gamma(\frac{n}{2})} x^{\frac{n}{2}-1} e^{-x/2}$ $x \geq 0$,
 where $\Gamma(a, \lambda)$ has density $g(t) = \frac{1}{\Gamma(a)} \lambda^a t^{a-1} e^{-\lambda t}$, $\Gamma(x) = \int_0^\infty u^{x-1} e^{-u} du$
t-distribution: If $Z \sim N(0,1)$ independent of $U \sim \chi_n^2$, then $\frac{Z}{\sqrt{U/n}} \sim t_n$.
 t_n has density $f_n(x) = \frac{\Gamma((n+1)/2)}{\Gamma(n/2)} \frac{1}{\sqrt{n\pi}} \frac{1}{(1+x^2/n)^{(n+1)/2}}$
 If $X_i \sim N(\mu, \sigma^2)$, $\frac{X_n - \mu}{\sigma/\sqrt{n}} \sim N(0,1)$. If don't know σ , estimate
 by $s^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$, distribution in $\frac{X_n - \mu}{s/\sqrt{n}} \sim t_{n-1}$
F-distribution: $U \sim \chi_m^2$, $V \sim \chi_n^2$, $F = \frac{U/m}{V/n} \sim F_{m,n}$.

Moment Generating Functions: $M(t) = E(e^{tx})$ is MGF of X
 $E(X^k) = \frac{d^k M}{dt^k} \Big|_{t=0}$. If X has MGF $M_X(t)$, $Y = a + bX$, $M_Y(t) = e^{at} M_X(bt)$
 If X, Y indep., $Z = X + Y$, $M_Z(t) = M_X(t) M_Y(t)$.

Two-sample t-test: $X_1 = x_1, \dots, X_n = x_n$ from pop 1 w. unknown var σ_x^2 .
 $Y_1 = y_1, \dots, Y_m = y_m$ from pop 2 w. unknown var σ_y^2 . Assume $\sigma_x^2 = \sigma_y^2$.
 To test $H_0: \mu_1 = \mu_2$ against $H_1: \mu_1 > \mu_2$, use test statistic
 $t = \frac{\bar{X} - \bar{Y}}{s/\sqrt{n+m}} \sim t_{n+m-2}$, $s^2 = \frac{1}{n+m-2} ((n-1)s_x^2 + (m-1)s_y^2)$. P-value is
 $P_{H_0}(\bar{X} - \bar{Y} > \bar{x} - \bar{y}) = P\left(\frac{\bar{X} - \bar{Y}}{s/\sqrt{n+m}} > \frac{(\bar{x} - \bar{y})}{s/\sqrt{n+m}}\right) = P(t_{n+m-2} > \frac{(\bar{x} - \bar{y})}{s/\sqrt{n+m}})$.

Mann-Whitney Test: Nonparametric, $X \sim F$, $Y \sim G$. $H_0: F = G$. Test ranks
 1. Concatenate X, Y into single vector Z . 2. Let n_1 be smaller sample size.
 3. Compute R = sum of ranks of smaller sample. 4. $R' = n_1(m+n+1) - R$
 5. If null true, $\pi = P(X < Y) = \frac{1}{2}$, so $\hat{\pi} = \frac{\# \text{times } X_i < Y_j \text{ all } i,j}{mn} = \frac{1}{mn} (R' - \frac{m(m+1)}{2})$
 If $\hat{\pi}$ much different from $\frac{1}{2}$, reject null.

Paired z-test: Know σ^2 , use z-test: $\mu_0 = 0$, $H_1: \mu_0 < 0$. P-value:
 $P_{H_0}(\bar{D} \leq \bar{d}) = P_{H_0}(\frac{\bar{D}}{\sigma/\sqrt{n}} \leq \frac{\bar{d}}{\sigma/\sqrt{n}}) = \Phi(\frac{\bar{d}}{\sigma/\sqrt{n}})$

Paired t-test: Don't know σ^2 but assume normal: $P(t_{n-1} \leq \frac{\bar{d}}{s/\sqrt{n}})$

Wilcoxon signed rank: Paired, nonparametric. Remove $D_i = 0$, R_i = rank of $|D_i|$. $W_i = \text{sign}(D_i) \times R_i$. $W_r = \sum_{i=1}^n W_i$. H_0 true, then
 $E(W_r) = \frac{n(n+1)}{4}$, $\text{Var}(W_r) = \frac{n(n+1)(2n+1)}{24}$

ANOVA: Obs. indep, normal, σ^2 . Have J obs. in each of I groups.
 $H_0: \mu_1 = \dots = \mu_I$. H_1 : @ least 1 diff. $Y_{ij} = \mu + \alpha_i + \epsilon_{ij}$, Y_{ij} = j th in i th group.
 μ = global mean, $\mu_i = \mu + \alpha_i$, ϵ_{ij} random errors, $\bar{Y}_i$ = group mean, $\bar{Y}$ = global mean.
 $\bar{Y}_i = \frac{1}{J} \sum_j Y_{ij}$, $\bar{Y} = \frac{1}{IJ} \sum_i \sum_j Y_{ij}$. Total Sum of Squares $SST = \sum_{ij} (Y_{ij} - \bar{Y})^2$
Between $SSB = J \sum_i (\bar{Y}_i - \bar{Y})^2$. Within $SSW = \sum_{ij} (Y_{ij} - \bar{Y}_i)^2$. $SST = SSB + SSW$
 SST has $IJ-1$ df, SSB $I-1$ df, SSW $I(J-1)$ df. Test statistic is
 $F = \frac{SSB/(I-1)}{SSW/(I(J-1))} \sim F_{I-1, I(J-1)}$.

Kruskal Wallis: Nonparametric one way ANOVA. R_{ij} rank of Y_{ij} .
Average rank of i th group $\bar{R}_i = \frac{1}{J} \sum_j R_{ij}$. Global: $\bar{R} = \frac{1}{IJ} \sum_{ij} R_{ij} = \frac{IJ+1}{2}$.
 $SSB = \sum_i J(\bar{R}_i - \bar{R})^2$. Test stat $K = \frac{12}{IJ(I+1)} SSB \sim \chi^2_{I-1}$.

Two Way ANOVA: 2 factors, I & J levels each, IJ treatments, $K > 1$ obs/treatment.
 $Y_{ijk} = \mu + \alpha_i + \beta_j + \delta_{ij} + \epsilon_{ijk}$. $SST = \sum_{ijk} (Y_{ijk} - \bar{Y}_{...})^2$. Within cell var
 $SSE = \sum_{ijk} (Y_{ijk} - \bar{Y}_{ij.})^2$. $SSA(\text{Row}) = Jk \sum_i (\bar{Y}_{i..} - \bar{Y}_{...})^2$. $SSB(\text{Col}) =$
 $Ik \sum_j (\bar{Y}_{.j.} - \bar{Y}_{...})^2$. $SSAB(\text{Interaction}) = K \sum_i \sum_j (\bar{Y}_{ij.} - \bar{Y}_{i..} - \bar{Y}_{.j.} + \bar{Y}_{...})^2$. $MS = \frac{SS}{df}$.
A-Test: $P(F_{I-1, IJ(K-1)}) \geq MSA/MSE$. B-Test: $P(F_{J-1, IJ(K-1)}) \geq MSB/MSE$.
Interaction: $P(F_{(I-1)(J-1), IJ(K-1)}) \geq MSAB/MSE$.

Randomized Block Design: Split I treatments across J different groups;
each $j \in J$ has samples in each $i \in I$. Model as $Y_{ij} = \mu + \alpha_i + \beta_j + \epsilon_{ij}$, Gaussian noise.
 $SST = \sum_{ij} (Y_{ij} - \bar{Y})^2$. Treatment: $SSA = J \sum_i (\bar{Y}_{i.} - \bar{Y})^2$. Subject: $SSB = I \sum_j (\bar{Y}_{.j} - \bar{Y})^2$.
Unexplained: $SSAB = \sum_{ij} (Y_{ij} - \bar{Y}_{i.} - \bar{Y}_{.j} + \bar{Y})^2$. $SST = SSA + SSB + SSAB$. $MSA = \frac{SSA}{I-1}$. $MSB = \frac{SSB}{J-1}$. $MSAB = \frac{SSAB}{(I-1)(J-1)}$.
P-value: $P(F_{(I-1)(J-1), (I-1)(J-1)}) \geq \frac{MSA}{MSAB}$.

Friedman's Test: Non-parametric RBS. Calculate ranks R_{ij} within each block.
Compute $Q = \frac{12}{IJ(I+1)} SSA = \frac{12}{IJ(I+1)} \sum_i (R_{i.} - \bar{R})^2 \sim \chi^2_{I-1}$. P-val: $P(\chi^2_{I-1} \geq Q)$.

Fisher's Exact Test: For table w/ 2 cat factors: $\begin{matrix} & A_1 & A_2 \\ B_1 & n_{11} & n_{12} \\ B_2 & n_{21} & n_{22} \end{matrix}$, expect two factors to be independent and n 's hypergeometric: $\text{Hyp}(N, K, n)$ describes k successes in n draws from pop. w/ K good: $P(X=k) = \frac{\binom{K}{k} \binom{N-K}{n-k}}{\binom{N}{n}}$, $E(X) = \frac{nK}{N}$. Expectation of n_{11} cells thus: $\frac{n_{1.} \cdot n_{.1}}{n}$. Two sided P-val: $P(|N_{11} - E(N_{11})| \geq |n_{11} - E(N_{11})|)$.

Linear Regression, MU RV, Prediction, Logistic Regression, O-Measures

Linear Regression: $y = \beta_0 + \beta_1 x + \epsilon$, $E(\epsilon) = 0$, $Var(\epsilon) = \sigma^2$
 Estimate $\beta_0, \beta_1 \rightarrow \hat{y} = \hat{\beta}_0 + \hat{\beta}_1 x$, $Y = XB + \epsilon$, $Y_i | X_i \sim N(\beta^T x_i, \sigma^2)$
 $\hat{B} = \arg \min_B \|Y - XB\|_2^2 = (X^T X)^{-1} X^T Y$, $E(\hat{B}) = \beta$, $Var(\hat{B}) = \sigma^2 (X^T X)^{-1}$
 Residuals: $e_i = y_i - \hat{y}_i = y_i - \hat{\beta}_0 - \hat{\beta}_1 x_i$
 Different than $\epsilon_i = y_i - \beta_0 - \beta_1 x$

Multivariate Random Variables: $Y_1, \dots, Y_n \rightarrow E(Y_i) = \mu$, $Cov(Y_i, Y_j) = \sigma_{ij}$

$Y = (Y_1, \dots, Y_n)$ Mean vector: $E(Y) = \mu = (\mu_1, \dots, \mu_n)$

Covariance Matrix: $Cov(Y) = \Sigma_{Y,Y} =$

$Z = b + AY$, b is vector, A is matrix

$E(Z) = b + A\mu$ $Cov(Z) = Cov(AY) = A(Cov(Y))A^T = A \Sigma A^T$

Quadratic transformation for vector: $X = Y^T A Y \in \mathbb{R}$ (scalars, $x = \text{any}$)

$X = Y^T A Y$ is quadratic form in Y .

To calculate $E(\text{quadratic form})$, use following:

1. Trace and Expectation are linear operators, so can be interchanged. Trace of $X = \begin{bmatrix} x_{11} & & x_{1n} \\ & \ddots & \\ x_{n1} & & x_{nn} \end{bmatrix}$,

$Tr(X) = \text{sum of diag} = \sum_{i=1}^n x_{ii}$

2. $X = Y^T A Y \in \mathbb{R}$ is real number, not matrix/vector, and if $x \in \mathbb{R}$, $Tr(X) = x$.

3. If A, B matrices, $tr(AB) = tr(BA)$

So: $E(Y^T A Y) = E(tr(Y^T A Y)) \rightarrow$ b/c $Y^T A Y$ scalar, trace scalar is

$= E[tr(A Y Y^T)] \rightarrow \text{fact 3}$

$= tr(E[A Y Y^T]) \rightarrow \text{fact 1}$

$= tr(A E[Y Y^T]) \rightarrow \text{fact 1}$

$= tr(A(\Sigma + \mu \mu^T)) \rightarrow E(X^2) = Var(X) + (E(X))^2$

$= tr(A \Sigma) + tr(A \mu \mu^T) \rightarrow \text{fact 1}$

$= tr(A \Sigma) + tr(\mu^T A \mu) \rightarrow \text{fact 3}$

$E(Y^T A Y) = tr(A \Sigma) + \mu^T A \mu \rightarrow \text{fact 2}$

Unbiased estimator for $\sigma^2 = \frac{RSS}{n-p}$, $RSS = \|e\|_2^2 = e^T e$

$e = Y - \hat{Y} = Y - X \hat{B} = Y - X(X^T X)^{-1} X^T Y = (I - X(X^T X)^{-1} X^T) Y = P_X Y$

where $P_X := I - X(X^T X)^{-1} X^T$ $RSS = (P_X Y)^T (P_X Y) = Y^T P_X^T P_X Y = Y^T P_X Y$

so $E(RSS) = E(Y^T P_{\perp} Y) = \sigma^2 \text{tr}(P_{\perp}) + E(Y)^T P_{\perp} E(Y)$ ($\Sigma = \sigma^2 I_n$)
 $= \sigma^2(n-p)$ ($E\{Y^T A Y\} = \text{tr}(A \Sigma) + \mu^T A \mu$) ($\text{Var}(Y_i) = \sigma^2$)
 $\uparrow$ $P_{\perp}$ is projection matrix orthogonal to X , (Y 's independent)
 thus $P_{\perp} E(Y) = P_{\perp} X \beta = 0$, so $E(Y)^T P_{\perp} E(Y) = 0$
 $\rightarrow$ trace of I_n is n , so $\text{tr}(P_{\perp}) = \text{tr}(I_n - X(X^T X)^{-1} X^T)$
 $= n - \text{tr}(X(X^T X)^{-1} X^T) = n - \text{tr}((X^T X)^{-1} X^T X) = n - \text{tr}(I_p) = n - p$

Predictor: Predicted $\hat{Y}_i$'s: $\hat{Y}_i = x_i^T \hat{\beta} = x_i^T (X^T X)^{-1} X^T Y$

Since $\hat{Y}_i$'s are RV, interested in variance of predictions

$$\text{Var}(\hat{Y}_i) = \text{Var}(x_i^T (X^T X)^{-1} X^T Y) = x_i^T (X^T X)^{-1} X^T \text{Var}(Y) X (X^T X)^{-1} x_i^T$$

$$= \sigma^2 x_i^T (X^T X)^{-1} (X^T X) (X^T X)^{-1} x_i^T = \sigma^2 x_i^T (X^T X)^{-1} x_i^T$$

Avg. variance is $\frac{1}{n} \sum_i \text{Var}(\hat{Y}_i) = \frac{\sigma^2}{n} \sum_i x_i^T (X^T X)^{-1} x_i^T =$
 $\frac{\sigma^2}{n} \text{tr}(X(X^T X)^{-1} X^T) = \frac{\sigma^2}{n} \text{tr}((X^T X)^{-1} X^T X) = \frac{\sigma^2}{n} \text{tr}(I_p) = \sigma^2 \frac{p}{n}$

More variables in model $\rightarrow$ more variable predicted values will be

Logistic Regression: $Y_i \in \{0, 1\}$. $Y_i | x_i \sim \text{Bern}(p(x_i, \beta))$.

$p(x_i, \beta) \mapsto \frac{1}{1 + e^{-\beta^T x_i}}$. $\beta^T x_i \geq 0$ means $p(x_i, \beta) \geq 0.5$, so $Y_i = 1$ like
 $p(x_i, \beta)$ basically probability that $Y = 1$.

MLE for β : $\text{lik}(\beta) = \prod_{i=1}^n P(Y_i | X_i = x_i, \beta) = \prod_{i=1}^n \left(\frac{1}{1 + e^{-\beta^T x_i}} \right)^{y_i} \left(\frac{1}{1 + e^{\beta^T x_i}} \right)^{1 - y_i}$
 $\ell(\beta) = \sum_{i=1}^n \beta^T x_i y_i - \log(1 + e^{\beta^T x_i})$; $\nabla \ell(\beta) = \sum_{i=1}^n x_i y_i - \frac{x_i e^{\beta^T x_i}}{1 + e^{\beta^T x_i}}$
 $= \sum_{i=1}^n x_i (y_i - p(x_i, \beta)) = X^T (y - p)$, $p = (p(x_1, \beta), p(x_2, \beta), \dots, p(x_n, \beta))$

Setting $\nabla \ell(\beta) = 0$ cannot yield closed form solution, so use iterative approach such as Newton's method for β .

Newton's Method: To find $f(x) = 0$, $x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}$.

To find β s.t. $\nabla \ell(\beta) = 0$: $\beta^{n+1} = \beta^n - \left(\frac{\partial^2 \ell(\beta)}{\partial \beta \partial \beta^T} \right)^{-1} \nabla \ell(\beta)$, hence

δ -method: If we have seq. RV's X_n , s.t.
 $\sqrt{n}(X_n - \theta) \rightarrow N(0, \sigma^2)$, then for any differentiable
 and non-zero function $g(\cdot)$, we have $\sqrt{n}(g(X_n) - g(\theta)) \rightarrow N(0, \sigma^2 g'(\theta)^2)$
Example: Suppose $X_n \sim \text{Bin}(n, p)$. Since $(X_n/n)_n$ is
 mean of n iid Bern(p) r.v.s, CLT says $\sqrt{n}(\frac{X_n}{n} - p) \rightarrow N(0, p(1-p))$
 and $\text{Var}(\frac{X_n}{n}) = \frac{p(1-p)}{n}$
 δ -method: $\sqrt{n}(g(\frac{X_n}{n}) - g(p)) \rightarrow N(0, p(1-p)(g'(p))^2)$
 Can use to calculate $\text{Var}(\log(\frac{X_n}{n}))$: $g(x) = \log(x)$
 $g'(x) = \frac{1}{x}$ δ -method: $\sqrt{n}(\log(\frac{X_n}{n}) - \log(p)) \rightarrow N(0, \frac{p(1-p)}{p^2})$
 $\text{Var}(\log(\frac{X_n}{n})) = \frac{1-p}{pn}$

Summary Part 4

χ^2 test for homogeneity/independence. To check null, equal
 distributions, $H_0: \pi_{11} = \pi_{12} = \dots = \pi_{1J}$, use $\chi^2 = \sum_{i,j} \frac{(O_{ij} - E_{ij})^2}{E_{ij}} = \sum_{i,j} \frac{(n_{ij} - (n_{i.} n_{.j})/n_{..})^2}{n_{i.} n_{.j} / n_{..}}$
 $df = (I-1)(J-1)$ so p -val: $P(\chi^2_{(I-1)(J-1)} \geq \chi^2)$.

Linear Regression Con't: Min vals are: $\hat{\beta}_0 = \bar{y} - \hat{\beta}_1 \bar{x}$, $\hat{\beta}_1 = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2}$
 $\hat{\beta}_1$ also: $\frac{\text{Cov}(x, y)}{\text{Var}(x)} = \frac{s_{xy}}{s_x^2} = \rho \sqrt{\frac{s_y^2}{s_x^2}}$, ρ = correlation coefficient = $\frac{s_{xy}}{\sqrt{s_x^2 s_y^2}}$, $|\rho| \leq 1$
 $\hat{\beta}_0, \hat{\beta}_1$ unbiased: $E(\hat{\beta}_0, \hat{\beta}_1) = (\beta_0, \beta_1)$ $\text{Var}(\hat{\beta}_0) = \frac{\sigma^2 \sum x_i^2}{n \sum x_i^2 - (\sum x_i)^2}$, $\text{Var}(\hat{\beta}_1) = \frac{\sigma^2}{n \sum x_i^2 - (\sum x_i)^2}$
 $\text{Cov}(\hat{\beta}_0, \hat{\beta}_1) = \frac{-\sigma^2 \sum x_i}{n \sum x_i^2 - (\sum x_i)^2}$, Estimate σ^2 by variance of residuals = $\frac{DSS}{n-p} = \frac{\sum \epsilon_i^2}{n-p}$

Multiple Linear Regression: $y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_{p-1} x_{p-1} + \epsilon_i$, $E(\epsilon_i) = 0$, $\text{Var}(\epsilon_i) = \sigma^2$
 Want to minimize $S(\beta_0, \beta_1, \dots, \beta_{p-1}) = \sum \epsilon_i^2 = \sum (y_i - \beta_0 - \beta_1 x_{i1} - \beta_2 x_{i2} - \dots - \beta_{p-1} x_{i,p-1})^2$
Matrix notation: $Y = XB + \epsilon$, $X \in \mathbb{R}^{n \times p}$, $B \in \mathbb{R}^{p \times 1}$, $\text{Cov}(\epsilon) = \sigma^2 I_{n \times n}$
 Thus $S(B) = \|Y - XB\|_2^2 = (Y - XB)^T (Y - XB)$, $\hat{B} = (X^T X)^{-1} X^T Y$, $\hat{Y} = X(X^T X)^{-1} X^T Y$
 $E(\hat{B}) = B$, $\text{Cov}(\hat{B}) = \sigma^2 (X^T X)^{-1}$, $\hat{\sigma}^2 = \frac{RSS}{n-p} = \frac{\|Y - X\hat{B}\|_2^2}{n-p} = \frac{\|e\|_2^2}{n-p}$