

①

Hypothesis Testing

Goal: To assess the validity of certain claims or answer questions about a population based on data.

- Which of the two claims do the data support? (i.e. which of the claims is more consistent with the observed data?)

Ques Is treatment A effective at treating Covid?

B ? " more effective than breastnut

- Consider modeling the score on final^{exam} as a linear regression as:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \dots + \beta_m X_m + \text{error}$$

↑ score on the final

↑ score on MT1

↑ score on MT2

↑ score on HWs etc

Ques: Is variable X_k helpful in predicting Y ?

Ques. Is the coin fair? (i.e. a) $p = \text{Prob}(H) = 0.5$ or
b) $p \neq 0.5$)

If (i) is true, then $X_i = \text{outcome on flip } i \sim \text{Bernoulli}(0.5)$
 & for odd X_i , $Y_n = \sum_{i=1}^n X_i \sim \text{Binomial}(n, 0.5)$

Ex:

- (a) $D = \{4H, 6T\}$

- (b) $D = \{40H, 60T\}$

- (c) $D = \{400H, 600T\}$

Which provide strongest evidence for (ii)
" " " " against (i)?

(2)

A HT helps us assess validity of a claim or assertion (i.e. hypothesis) based on a data set (taking sampling variability & distr. into account).

- A HT consists of two hypotheses (non-overlapping)
 H_0 : null hypothesis, status quo, usually requires no change
vs
 H_a : alternative hypothesis, the claim to be assessed compared to H_0 , usually requires substantial action.

- A HT is asymmetric in deciding the validity of H_0 vs H_a .
 - If there is sufficient evidence for H_a in data, you reject H_0 and accept H_a .
 - If there is no sufficient evidence for H_a in data, you fail to reject H_0 (or retain H_0).
 - One assumes population is as in H_0 , and tests this assumption with evidence in data.
If data (evidence) contradicts "substantially" with H_0 , then s/he rejects H_0 in favor of H_a .

Ex: Can fair or not? $H_0: p = 0.5$ vs $H_a: p \neq 0.5$
 $D = \{4HT, 6T\}$ reject or retain H_0 ?

Some examples of H_0 & H_a :

Situation	Null Hypothesis	Alternative Hyp., ③
• Effect of a treatment on a disease	not effective (i.e., like a placebo)	effective
• medical testing (for a disease)	healthy / not infected	infected
• Drug test, A vs B	no difference	difference (e.g., A better than B)
• extraterrestrial life in exoplanets	none (i.e., no Earthlike life)	Earthlike life exists
* US Criminal System	innocent (not guilty)	guilty
• medical mask vs cloth " etc.	no difference	medical masks better

Criminal Trial Analogy:

H_0 : not guilty

H_a : guilty

• need evidence beyond reasonable doubt to declare the suspect guilty.

If no such evidence exists, suspect is not guilty (i.e. evidence did not corroborate the guilt verdict)

(4)

Outline of a generic HT:

0. Set up H_0 & H_1 (if not provided)

1. Determine a statistic (as a measure of evidence from data) $T = T(D) = T(X_1, \dots, X_n)$ and compute it.

2. Determine a region of rejection (i.e. rejection region (RR) for when or which values of T constitute "enough" evidence for H_1 .

3. Check if T falls in RR or not.

If T falls in RR, reject H_0

" " does not fall in RR, fail to reject H_0 or retain H_0 .

In general Hypotheses are written in terms of parameters
Ex: $X_1, \dots, X_n \sim F(x; \theta)$, θ is unknown parameter

simple null $\rightarrow H_0: \theta = \theta_0$ vs $H_1: \theta \neq \theta_0$ (2-sided alter)

one-sided tests $\left\{ \begin{array}{l} \text{or} \\ H_0: \theta \leq \theta_0 \end{array} \right.$ vs $H_1: \theta > \theta_0$ (right-sided alternative)
 or $H_0: \theta \geq \theta_0$ vs $H_1: \theta < \theta_0$ (left-sided alternative)

Ex: (Coin example) $H_0: p = 0.5$ vs $H_1: p \neq 0.5$ (2-sided)

if $H_0: p \leq 0.5$ vs $H_1: p > 0.5$ (1-sided)

- In a H.T, one can make two types of error: ⑤

	Truth	Decision	
		Retain H_0	Reject H_0
H_0		✓ true -ive	X Type I error false +ive
H_a		X Type II error false -ive	✓ true +ive

Type I error: Rejecting a true H_0

" II " : Retaining a false H_0 (when H_a is true)

$P(\text{Type I error}) = \alpha = \text{level of significance}$

$P(\text{Type II error}) = \beta$

$P(\text{reject } H_0 | H_0 \text{ is false}) = \text{power} = 1 - \beta$

- Want α & β small i.e. want small α & high power
- may not be possible to attain both goals, instead one controls $P(\text{type I error}) \leq \alpha$ and aims for highest power.

Types of Tests

1) Wald's Test:
(aka Normal-Based Test) $D = \{X_1, \dots, X_n\}$ $X_i \stackrel{\text{iid}}{\sim} F(x, \theta)$ known
 $H_0: \theta = \theta_0$ vs $H_a: \theta \neq \theta_0$

$\hat{\theta}$ is an estimator for θ

$\hat{\theta}$ is asymptotically normal, i.e. $\hat{\theta} \stackrel{\text{approx}}{\sim} N(\theta, se^2)$

Assume H_0 is true, then $\hat{\theta} \overset{\text{approx}}{\sim} N(\theta_0, se^2)$ (6)
 (i.e. under H_0 , $\hat{\theta} \overset{\text{approx}}{\sim} N(\theta_0, se^2)$).

Statistic: $\hat{\theta}$ or better than this

$$W = \frac{\hat{\theta} - \theta_0}{se} \overset{\text{approx}}{\sim} N(0, 1)$$

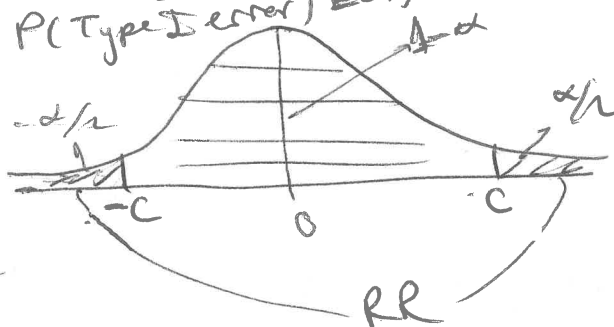
$se \leftarrow$ if not known, replace with a consistent estimator $\hat{se}$.

Decision Rule (DR) Reject H_0 if $|\hat{\theta} - \theta_0|$ is large, i.e.
 $|W| > c$

(that is $RR: \{W: W > c \text{ or } W < -c\}$)

But to have
 we need to

$$P(W > c | H_0) + P(W < -c | H_0) = \alpha$$



$$\text{Hence, } P(W > c) = P(W < -c) = \alpha/2$$

$$\Rightarrow c = z_{\alpha/2}$$

Thus, DR: Reject H_0 if $|W| > z_{\alpha/2}$

• With this DR, $P(\text{Type I error}) = P(\text{reject } H_0 | H_0 \text{ true})$

$$= P(|W| > z_{\alpha/2}) = P(W < -z_{\alpha/2}) + P(W > z_{\alpha/2}) = \frac{\alpha}{2} + \frac{\alpha}{2} = \alpha$$

Notes:

1) W is measuring how far away is $\hat{\theta}$ from θ_0 in terms of se when H_0 is true

i.e. if H_0 is true $\hat{\theta} - \theta_0 \approx 0$ (for large n) since $\hat{\theta} \rightarrow \theta = \theta_0$

else $|\hat{\theta} - \theta_0| > 0$

(7)

2) But precision or variability of $\hat{\theta}$ is important.

eg. $\hat{\theta} - \theta_0 = 0.1$

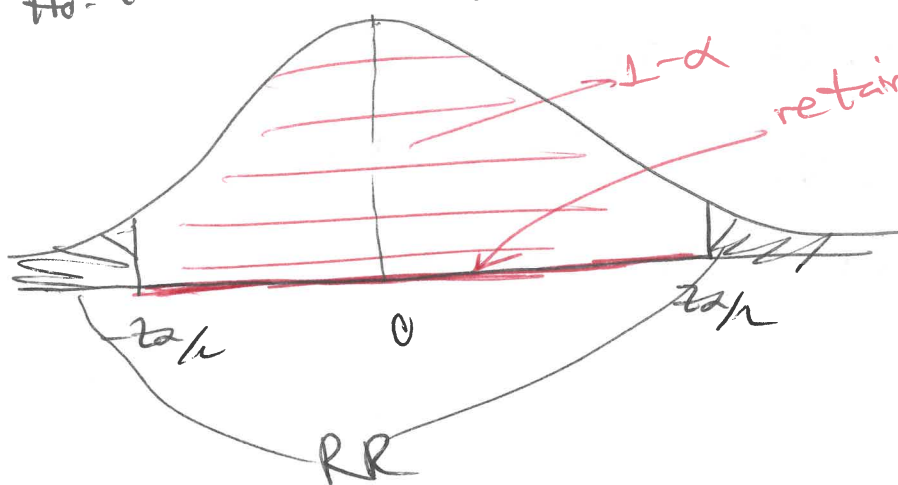
case (i) $se(\hat{\theta}) = 0.01 \Rightarrow W = \frac{0.1}{0.01} = 10$ strong evidence for H_0 .

case (ii) $se(\hat{\theta}) = 100 \Rightarrow W = \frac{0.1}{100} = 0.001$ no evidence for H_0 .
 (highly consistent with H_0)

Thus, how many se's is the difference $\hat{\theta} - \theta_0$ is important.

• Connection ^(of Two-sided Tests) with $(1-\alpha)$ CI

If $H_0: \theta = \theta_0$ is true, then $W \sim Z \sim N(0,1)$



• Retain H_0 if $-z_{\alpha/2} < W < z_{\alpha/2}$

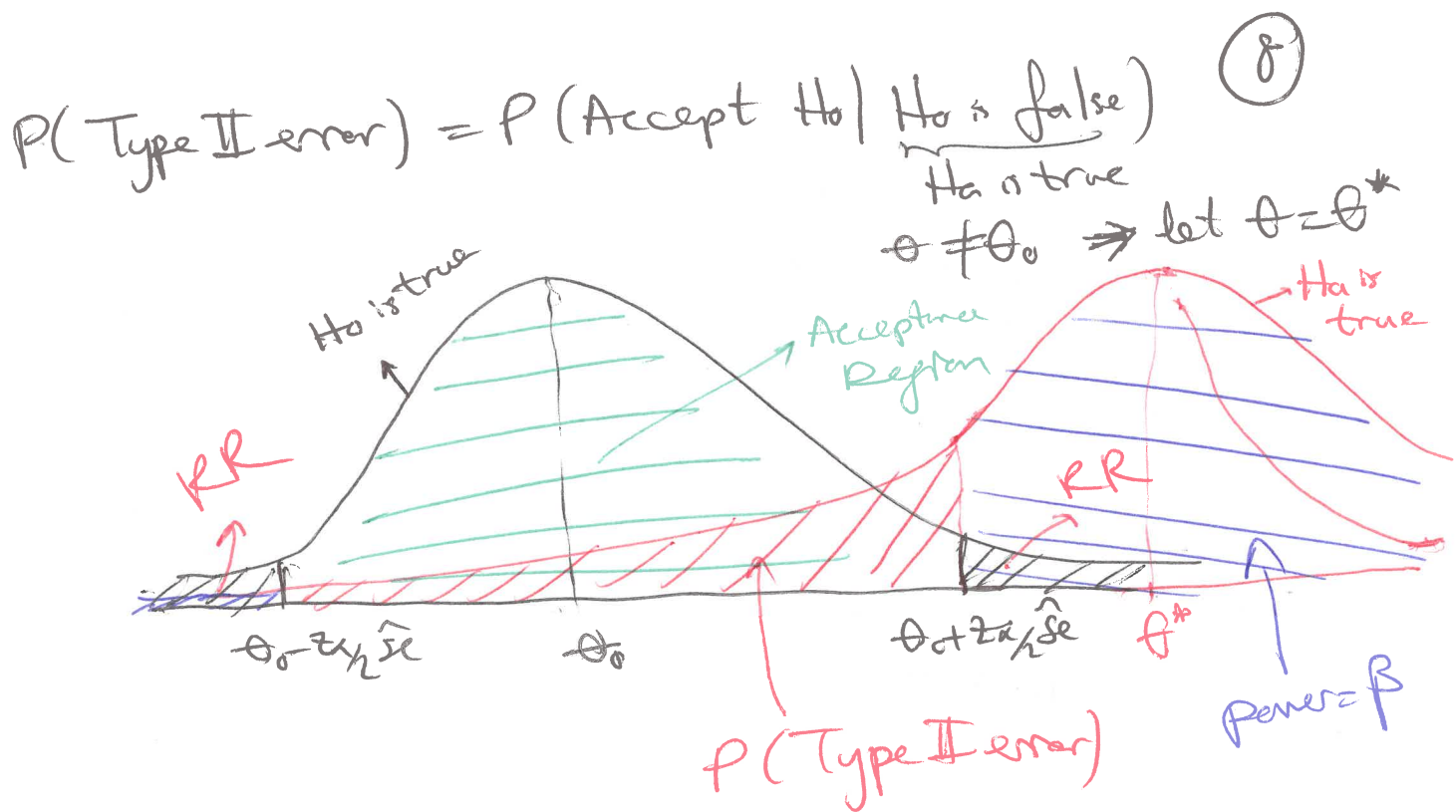
$$\Leftrightarrow -z_{\alpha/2} < \frac{\hat{\theta} - \theta_0}{se} < z_{\alpha/2}$$

$$\Leftrightarrow \hat{\theta} - z_{\alpha/2} se < \theta_0 < \hat{\theta} + z_{\alpha/2} se$$

$$\Leftrightarrow \theta_0 \in (\hat{\theta} - z_{\alpha/2} se, \hat{\theta} + z_{\alpha/2} se)$$

$(1-\alpha)$ CI

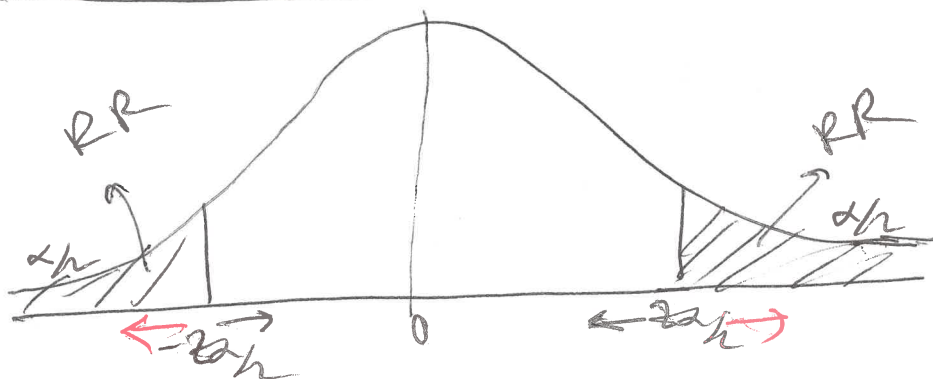
So, reject $H_0: \theta = \theta_0$, if θ_0 is not in $(1-\alpha)$ CI
 retain $H_0: \theta = \theta_0$ if θ_0 is in the $(1-\alpha)$ CI.



One-Sided Alternatives

- If $H_0: \theta \leq \theta_0$ vs $H_a: \theta > \theta_0$
 DR: Reject H_0 , if $W > z_\alpha$
- If $H_0: \theta \geq \theta_0$ vs $H_a: \theta < \theta_0$
 DR: Reject H_0 , if $W < -z_\alpha$.

Dependence of 2-sided test on α



- If $\alpha \uparrow$, $RR \uparrow$
- If $\alpha \downarrow$, $RR \downarrow$

(9)

$P(\text{Type I error}) = \alpha$ "our tolerance prob. for rejecting a true H_0 "

common choices $\alpha = 0.01$, 0.05 or 0.10

Ex: 64 tosses of a coin, 25 H, 39 T.

Test whether coin is fair or not at $\alpha = 0.05$ level.

Sol'n: $H_0: p = 0.5$ vs $H_a: p \neq 0.5$

$\hat{p} = \bar{X} = \frac{\sum X_i}{n} = \frac{25}{64}$ where $X_i = \mathbb{I}(\text{coin shows H}) \sim \text{Bernoulli}(p)$ with $p = P(H)$.

$n = 64$ is large, so by CLT,

$\hat{p}$ approx Normal. $(N(p, \text{se}^2(\hat{p})))$

So, Wald's test is applicable, with

test stat: $W = \frac{\hat{p} - p_0}{\text{se}(\hat{p})} = \frac{0.39 - 0.5}{\text{se}(\hat{p})}$

$\text{se}(\hat{p}) = \sqrt{\frac{p(1-p)}{n}}$ but under $H_0: p = 0.5$

$\text{se}(\hat{p}) = \sqrt{\frac{0.5(1-0.5)}{64}} = \frac{1}{16} = 0.0625 \Rightarrow W = \frac{-0.11}{0.0625} = -1.76$

DR: Reject H_0 if $|W| > z_{0.025} = 1.96$

Here, we retain H_0 , since $1.76 < 1.96$.

• one could also estimate $\text{se}(\hat{p})$ as $\sqrt{\frac{\hat{p}(1-\hat{p})}{n}} \approx 0.061$ (since $\hat{p}$ is consistent)
In fact if $H_0: p \leq 0.5$ or $H_0: p \geq 0.5$, we have to use this se.

Ex: What if we observe 10 H & 54 T in the previous exam?

$$\text{So, } W = \frac{\hat{p} - p_0}{\text{se}(\hat{p})} = \frac{10/64 - 0.5}{0.0625} = -5.5$$

Since $5.5 > 1.96$, we reject H_0 .

Wald Test for Two Independent Populations

$D_1 = \{X_1, X_2, \dots, X_n\}$ with $X_i \text{ i.i.d. } F(x; \theta_1)$

$D_2 = \{Y_1, Y_2, \dots, Y_m\}$ " $Y_j \text{ i.i.d. } F(y; \theta_2)$

$X_i \perp Y_j$ for all i, j . & $\hat{\theta}_1$ & $\hat{\theta}_2$ are both asymptotically normal

HT: $H_0: \theta_1 = \theta_2$ vs $H_a: \theta_1 \neq \theta_2$

$$\theta_1 = \theta_2 \Leftrightarrow \delta = \theta_1 - \theta_2 = 0$$

$$\Rightarrow H_0: \delta = 0 \leftarrow \delta_0 \text{ vs } H_a: \delta \neq 0$$

and $\hat{\delta} = \hat{\theta}_1 - \hat{\theta}_2$ A.N. ($\hat{\theta}_1 - \hat{\theta}_2 \text{ approx } N(\theta_1 - \theta_2, \text{se}^2(\hat{\theta}_1 - \hat{\theta}_2))$)

$$W = \frac{\hat{\delta} - \delta_0}{\text{se}(\hat{\delta})} = \frac{(\hat{\theta}_1 - \hat{\theta}_2) - 0}{\text{se}(\hat{\theta}_1 - \hat{\theta}_2)} = \frac{\hat{\theta}_1 - \hat{\theta}_2}{\text{se}(\hat{\theta}_1 - \hat{\theta}_2)} \Rightarrow |W| \stackrel{?}{>} Z_{\alpha/2}$$

$$\text{Here, } \text{Var}(\hat{\theta}_1 - \hat{\theta}_2) = \text{Var}(\hat{\theta}_1) + \text{Var}(\hat{\theta}_2) \\ = \text{se}^2(\hat{\theta}_1) + \text{se}^2(\hat{\theta}_2)$$

$$\Rightarrow \text{se}(\hat{\theta}_1 - \hat{\theta}_2) = \sqrt{\text{se}^2(\hat{\theta}_1) + \text{se}^2(\hat{\theta}_2)}$$

(11)

Z-TestCase (I) $D = \{X_1, X_2, \dots, X_n\}$, $X_i \stackrel{\text{iid}}{\sim} N(\mu, \sigma^2)$ μ unknown,
 σ^2 known. $H_0: \mu = \mu_0$ vs $H_a: \mu \neq \mu_0$.Under H_0 ,

$$Z = \frac{\bar{X} - \mu_0}{\sigma/\sqrt{n}} \sim N(0, 1)$$

DR: Reject H_0 , if $|Z| > z_{\alpha/2}$.Case II: D as above, $X_i \stackrel{\text{iid}}{\sim} F(x)$ where $E[X_i] = \mu$ unknown,
 $\text{Var}(X_i) = \sigma^2$ knownThen for large n ($n > 30$),by CLT, $\bar{X} \stackrel{\text{approx}}{\sim} N(\mu, \frac{\sigma^2}{n})$ so, under H_0 ,

$$Z = \frac{\bar{X} - \mu_0}{\sigma/\sqrt{n}} \stackrel{\text{approx}}{\sim} N(0, 1)$$

DR: Reject H_0 , if $|Z| > z_{\alpha/2}$ Note:Wald test

vs

Z-test σ is assumed known

$$W = \frac{\hat{\mu} - \mu_0}{\text{se}(\hat{\mu})}$$

$$Z = \frac{\hat{\mu} - \mu_0}{\sigma/\sqrt{n}}$$

$$\text{or } W = \frac{\hat{\mu} - \mu_0}{\hat{\text{se}}(\hat{\mu})}$$

 σ is known F is not known or nonnormal, but $n > 30$ or $F \equiv \text{Normal}$, any n . $\hat{\mu} \approx \text{A.N.}$

- One-sided Z-tests

- $H_0: \mu \leq \mu_0$ vs $H_a: \mu > \mu_0$

DR: Reject H_0 if $Z > z_\alpha$

- $H_0: \mu > \mu_0$ vs $H_a: \mu \leq \mu_0$

DR: Reject H_0 if $Z < -z_\alpha$

Z-Test for 2 Indep. Populations

Case I: $D_1 = \{X_1, X_2, \dots, X_n\}$, $X_i \stackrel{\text{iid}}{\sim} N(\mu_1, \sigma_1^2)$
 $D_2 = \{Y_1, Y_2, \dots, Y_m\}$, $Y_j \stackrel{\text{iid}}{\sim} N(\mu_2, \sigma_2^2)$ $\left(\begin{matrix} X_i + Y_j \\ \text{for all } i, j \end{matrix} \right)$

μ_1, μ_2 unknown, σ_1^2 & σ_2^2 known

$H_0: \mu_1 = \mu_2$ vs $H_a: \mu_1 \neq \mu_2$

Let $\delta = \mu_1 - \mu_2 \Rightarrow H_0: \delta = 0$ vs $H_a: \delta \neq 0$

$$Z = \frac{(\bar{X} - \bar{Y}) - \delta_0}{\text{se}(\bar{X} - \bar{Y})} \leftarrow ?$$

Here, $\text{Var}(\bar{X} - \bar{Y}) = \text{Var}(\bar{X}) + \text{Var}(\bar{Y}) = \frac{\sigma_1^2}{n} + \frac{\sigma_2^2}{m}$

$$\Rightarrow \text{se}(\bar{X} - \bar{Y}) = \sqrt{\frac{\sigma_1^2}{n} + \frac{\sigma_2^2}{m}}$$

DR: Reject H_0 , if $|Z| > z_{\alpha/2}$

(13)

Case II $X_i \stackrel{\text{iid}}{\sim} F_X$ $Y_j \stackrel{\text{iid}}{\sim} F_Y$

$$w/ E[X_i] = \mu_1 \quad E[Y_j] = \mu_2$$

$$\text{Var}(X_i) = \sigma_1^2$$

$$\text{Var}(Y_j) = \sigma_2^2$$

 σ_1^2, σ_2^2 known

$$n \geq 30$$

$$m \geq 30$$

By CLT, $Z = \frac{(\bar{X} - \bar{Y}) - \delta_0}{\sqrt{\frac{\sigma_1^2}{n} + \frac{\sigma_2^2}{m}}} \approx N(0,1)$

DR: Reject H_0 , if $|Z| > z_{\alpha/2}$.

t-Test

$D = \{X_1, X_2, \dots, X_n\}$, $X_i \stackrel{\text{iid}}{\sim} N(\mu, \sigma^2)$ μ unknown, σ^2 unknown

$H_0: \mu = \mu_0$ vs $H_a: \mu \neq \mu_0$

$$T = \frac{\bar{X} - \mu_0}{s/\sqrt{n}} \quad \text{where} \quad s^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$$

Under H_0 , $T \sim t_{n-1}$

Notes

1) If n is large ($n \geq 30$), by CLT, $\frac{\bar{X} - \mu_0}{s/\sqrt{n}} \approx N(0,1)$ & Slutsky

so, one can use a Z-test.

2) So, t-test is most useful when n is small ($n < 30$)

3) normality of X_i is a weak assumption, if not severely nonnormal t-test can still be used.

D.R. for t-testReject H_0 , if $|T| > t_{\alpha/2, n-1}$ ← table look-up
or use R, Python, etc.One-Tailed t-Test• $H_0: \mu \leq \mu_0$ vs $H_a: \mu > \mu_0$ DR: Reject H_0 , if $T > t_{\alpha, n-1}$ • $H_0: \mu \geq \mu_0$ vs $H_a: \mu < \mu_0$ DR: Reject H_0 , if $T < -t_{\alpha, n-1}$ t-Test for Two Indep. Populations (Unpaired) $D_1 = \{X_1, X_2, \dots, X_n\}$, $X_i \stackrel{\text{iid}}{\sim} N(\mu_1, \sigma_1^2)$ • σ_1^2, σ_2^2 unknown $D_2 = \{Y_1, Y_2, \dots, Y_m\}$, $Y_j \stackrel{\text{iid}}{\sim} N(\mu_2, \sigma_2^2)$ • $\min(n, m) < 30$ $X_i \perp Y_j$ for all i, j $H_0: \mu_1 = \mu_2$ vs $H_a: \mu_1 \neq \mu_2$

$$\delta = \mu_1 - \mu_2$$

$$T = \frac{\bar{X} - \bar{Y} - \delta_0}{\widehat{se}(\bar{X} - \bar{Y})}$$

Case I: $\sigma_1^2 = \sigma_2^2 = \sigma^2$

$$s_p^2 = \frac{(n-1)s_1^2 + (m-1)s_2^2}{n+m-2} \quad \& \quad \widehat{se}(\bar{X} - \bar{Y}) = s_p \sqrt{\frac{1}{n} + \frac{1}{m}}$$

DR: Reject H_0 , if $|T| > t_{\alpha/2, n+m-2}$

Case II: $\sigma_1^2 \neq \sigma_2^2$

$$\widehat{SE}(\bar{X} - \bar{Y}) = \sqrt{\frac{s_1^2}{n} + \frac{s_2^2}{m}}$$

$$\text{where } s_1^2 = \frac{1}{n-1} \sum (x_i - \bar{x})^2$$

$$s_2^2 = \frac{1}{m-1} \sum (y_i - \bar{y})^2$$

DR: Reject H_0 if $|T| > t_{\alpha/2, df}$ where df is adjusted (Satterthwaite approximation)
 $\max(n-1, m-1) \leq df \leq n+m-2$. (from software)• One-tailed versions are possible with $t_{\alpha, df}$ or $-t_{\alpha, df}$ in DR.eg $H_0: \mu_1 \leq \mu_2$ vs $H_a: \mu_1 > \mu_2$ DR: Reject H_0 if $T > t_{\alpha, df}$.Paired t-Test $D_1 = \{x_1, x_2, \dots, x_n\}$, $x_i \stackrel{iid}{\sim} N(\mu_1, \sigma_1^2)$ σ_1^2, σ_2^2 unknown. $D_2 = \{y_1, y_2, \dots, y_n\}$, $y_j \stackrel{iid}{\sim} N(\mu_2, \sigma_2^2)$ x_i & y_i are dependent (there is reason to pair x_i & y_i) $H_0: \mu_1 = \mu_2$ vs $H_a: \mu_1 \neq \mu_2$

take differences

 $D = \{d_1, d_2, \dots, d_n\}$ where $d_i = x_i - y_i$ $i=1, 2, \dots, n$ $d_i \stackrel{iid}{\sim} N(\mu_1 - \mu_2, \sigma_1^2 + \sigma_2^2) \equiv N(\delta, \sigma_\delta^2)$ $H_0: \delta = 0$ vs $H_a: \delta \neq 0$

(16)

$$T = \frac{\bar{d} - 0}{s_d/\sqrt{n}} = \frac{\bar{d}}{s_d/\sqrt{n}} \quad \text{where } \bar{d} = \frac{\sum d_i}{n} \quad \& \quad s_d^2 = \frac{1}{n-1} \sum (d_i - \bar{d})^2$$

DR: Reject H_0 if $|T| > t_{\alpha/2, n-1}$

p-value (Observed Significance Level)

- An alternate way of reporting the result/decision of the HT.
- So far the conclusion was binary with only "reject H_0 " or "retain H_0 " at a given α level.
- Let T_{obs} be the observed (computed) value of the test stat.
- p-value: Probability of getting the current value of the statistic or more "extreme" values than the observed statistic.

$$p\text{-value} = \text{Prob}(\text{statistic } T \text{ more extreme than } T_{obs} | H_0 \text{ is true})$$

"more extreme" means "towards" or "in favor of H_0 " i.e. further "away from H_0 ".

DR (with p-value):

Reject H_0 if $p\text{-value} < \alpha$.

Ex: Recall Wald Test:

$$D = \{X_1, X_2, \dots, X_n\} \quad X_i \stackrel{iid}{\sim} F(x; \theta), \quad \hat{\theta} \text{ an estimator for } \theta$$

and $\hat{\theta}$ is A.N.

$$W = \frac{\hat{\theta} - \theta_0}{se(\hat{\theta})} \quad \text{or} \quad \frac{\hat{\theta} - \theta_0}{\hat{se}(\hat{\theta})}, \quad W \approx N(0, 1)$$

Let the observed or computed test stat be W_{obs}

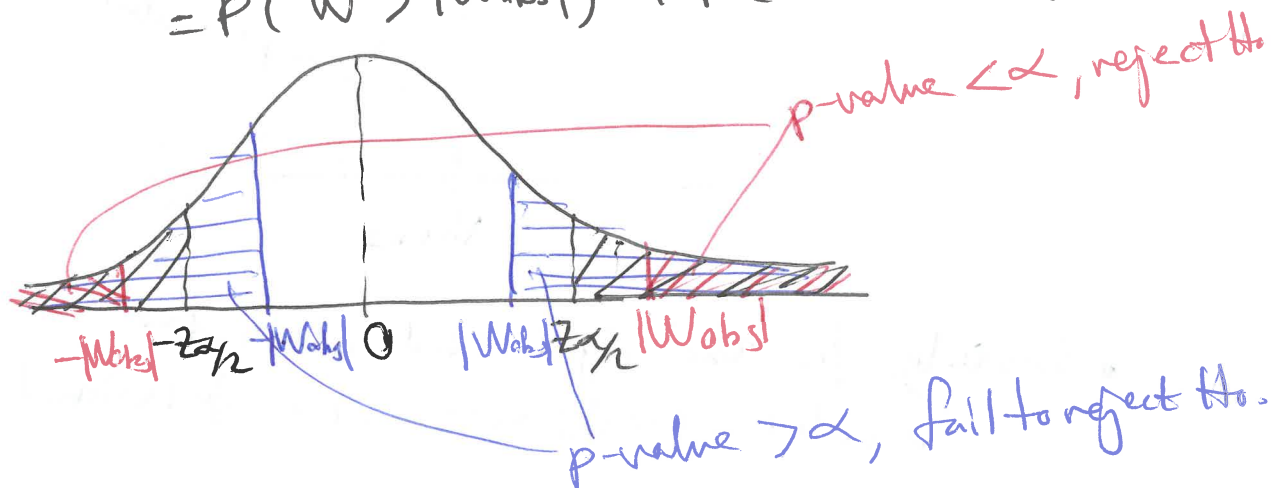
(17)

Ques Which values of W would be more extreme than W_{obs} (under H_0)?

Ans. $|W| > |W_{obs}|$

so $p\text{-value} = P(|W| > |W_{obs}|)$

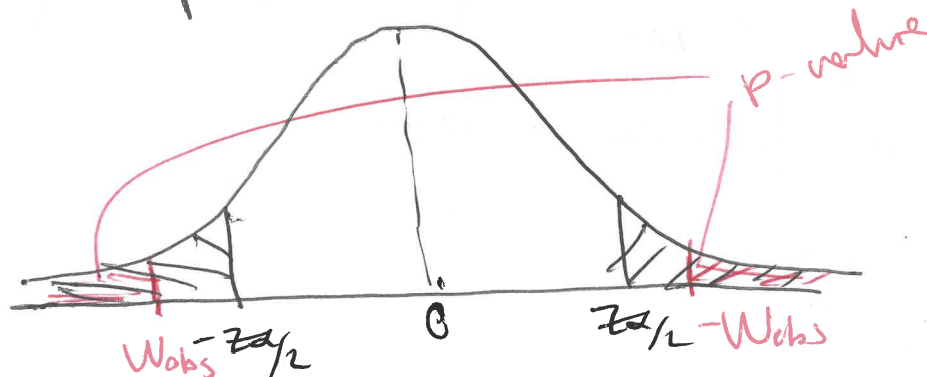
$$= P(W > |W_{obs}|) + P(W < -|W_{obs}|)$$



Since $W \overset{\text{approx}}{\sim} N(0, 1)$,

$$p\text{-value} \approx 2(1 - \Phi(|W_{obs}|)) = 2\Phi(-|W_{obs}|)$$

Ex: If $W_{obs} < 0$, $p\text{-value} = 2\Phi(W_{obs})$



p-value for One-Sided Wald Test

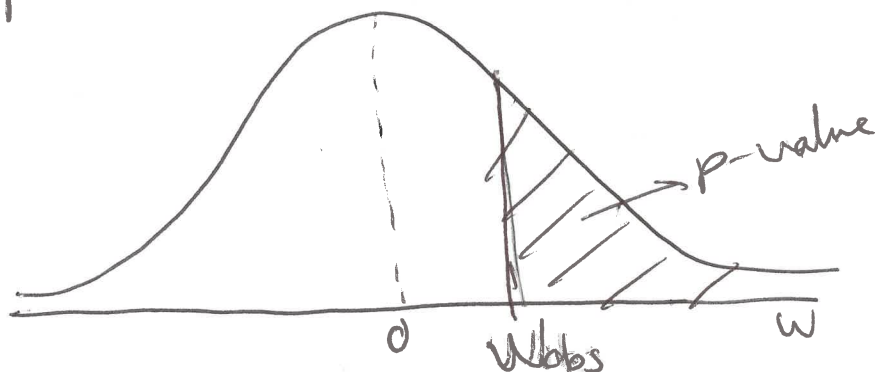
(18)

$H_0: \theta \leq \theta_0$ vs $H_a: \theta > \theta_0$, p-value=? (right-sided)

$$W_{obs} = \frac{\hat{\theta} - \theta_0}{se(\hat{\theta})}$$

$W > W_{obs}$ would be more extreme
W values than observed.

So, p-value = $P(W > W_{obs}) = 1 - \Phi(W_{obs})$



- Similarly, for the left-sided alternative,
p-value = $P(W < W_{obs}) = \Phi(W_{obs})$

- p-values for Z-tests & t-tests are similar.

Ex: For 2-sided Z-test with $Z_{obs} > 0$

$$\begin{aligned} \text{p-value} &= P(Z > Z_{obs}) + P(Z < -Z_{obs}) \\ &= 2(1 - \Phi(Z_{obs})) \end{aligned}$$

- p-values for one-sided Z-tests can also be obtained by replacing W, W_{obs} with Z, Z_{obs} , respectively.

- Likewise for the t-tests.

Two Caveats: (i) A large p-value is not strong evidence in favor of H_0 . " " can occur when (a) H_0 is true or (ii) H_0 is false but test has low power
(2) p-value $\neq$ Prob(H_0 is true | data)

(Comparing correct prediction rates of 2 ^{prediction} algorithms) (19)

Ex: $D_1 = \{X_1, X_2, \dots, X_n\}$, $X_i \stackrel{iid}{\sim} \text{Bernoulli}(p_1)$
 $D_2 = \{Y_1, Y_2, \dots, Y_m\}$, $Y_j \stackrel{iid}{\sim} \text{Bernoulli}(p_2)$ ($X_i \perp Y_j$ for all i, j)

Is $p_1 = p_2$? (n, m are large)

Sol'n: $H_0: p_1 = p_2$ vs $H_a: p_1 \neq p_2$
 equivalently

$H_0: p_1 - p_2 = 0$ vs $H_a: p_1 - p_2 \neq 0$

The MLE of $p_1 - p_2$ is $\hat{p}_1 - \hat{p}_2$, since

" " of p_1 is $\hat{p}_1$ &

" " of p_2 is $\hat{p}_2$.

$\hat{p}_1 \approx N(p_1, \text{Var}(\hat{p}_1))$
 $\hat{p}_2 \approx N(p_2, \text{Var}(\hat{p}_2))$ } $\Rightarrow \hat{p}_1 - \hat{p}_2 \approx N(p_1 - p_2, \text{Var}(\hat{p}_1) + \text{Var}(\hat{p}_2))$
 where $\text{Var}(\hat{p}_1) = \frac{p_1(1-p_1)}{n}$
 & $\text{Var}(\hat{p}_2) = \frac{p_2(1-p_2)}{m}$

So, Wald test stat is

$$W = \frac{(\hat{p}_1 - \hat{p}_2) - 0}{\text{se}(\hat{p}_1 - \hat{p}_2)} \quad \text{where } \text{se}(\hat{p}_1 - \hat{p}_2) = \sqrt{\text{Var}(\hat{p}_1) + \text{Var}(\hat{p}_2)}$$

$$= \sqrt{\frac{p_1(1-p_1)}{n} + \frac{p_2(1-p_2)}{m}}$$

But p_1, p_2 unknown

$\hat{p}_1$ is consistent for p_1 &
 $\hat{p}_2$ " " " p_2

$$\Rightarrow \hat{\text{se}}(\hat{p}_1 - \hat{p}_2) = \sqrt{\frac{\hat{p}_1(1-\hat{p}_1)}{n} + \frac{\hat{p}_2(1-\hat{p}_2)}{m}}$$

$$\Rightarrow W = \frac{\hat{p}_1 - \hat{p}_2}{\hat{\text{se}}(\hat{p}_1 - \hat{p}_2)}$$

(20)

D.R. Reject H_0 if $|W| > z_{\alpha/2}$.

OR compute
 $p\text{-value} = 2 \min(P(W > W_{obs}), P(W < W_{obs}))$

D.R. Reject H_0 if $p\text{-value} \leq \alpha$.

Likelihood Ratio Test (LRT)

$H_0: \theta \in \Omega_0$ vs $H_a: \theta \in \Omega \setminus \Omega_0$

LRT stat is

$$T_{LRT} = 2 \log \left(\frac{\sup_{\theta \in \Omega} L(\theta)}{\sup_{\theta \in \Omega_0} L(\theta)} \right) \quad \text{where } L(\theta) = f(x_1, x_2, \dots, x_n; \theta)$$

$\uparrow$ likelihood $\uparrow$ joint pdf

$$= 2 \log \left(\frac{L(\hat{\theta}_{MLE})}{L(\hat{\theta}_0)} \right)$$

where $\hat{\theta}_{MLE}$ is the (global) MLE &

$\hat{\theta}_0$ is " MLE under H_0 .

D.R. Reject H_0 if $T_{LRT} > c$ $\nwarrow$ critical value based on distr. of T_{LRT} .

Note: Almost all tests we have seen are special cases of LRT (i.e. can be derived from T_{LRT}).