

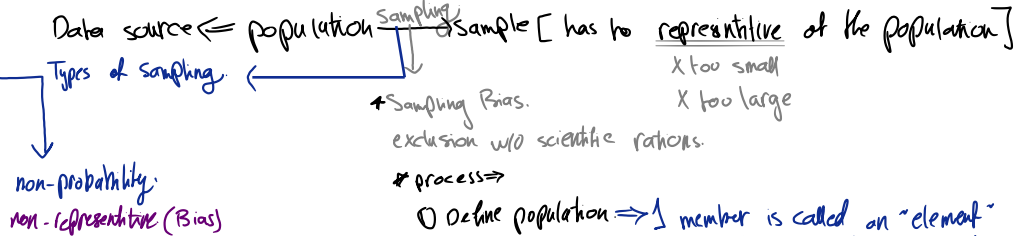
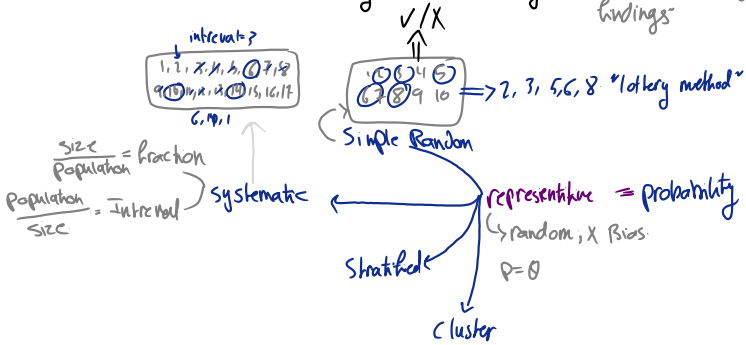
Biostatistics unit 1:

- Time consuming
- expensive
- fast & efficient
- not always possible
- strong external validity ~ allows for generalizing our findings

* Data collection:

① Primary $\Rightarrow$ "you" collect it for a specific use
 $\hookrightarrow$ questionnaire, survey, ...

② secondary $\Rightarrow$ they're already collected; medical records, published research
 $\hookrightarrow$ might be collected for a different reason.



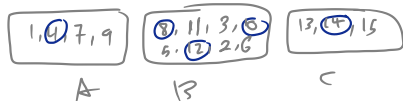
* process $\Rightarrow$

- ① Define population $\Rightarrow$ 1 member is called an "element"
- ② Frame $\Rightarrow$ all accessible members / population elements ~
- ③ Method
- ④ Size $\Rightarrow$ Cochran's formula = $Z^2 \cdot \frac{pq}{n} \cdot (1-P)$
 $\hookrightarrow$ depends on the CI, e^2 margin of error "precision", p estimate, q table.
- ⑤ execute the process



1. external validity
2. Representation.

$x/V \Leftarrow$



A, B, C: characteristics needed for the study & it has to be proportional to the population.

○: selected by randomly/systemically

1. x sampling frame $\checkmark/x \Leftarrow$
2. costly
3. external validity.
4. speed.
5. not always possible

cluster
↓
predefined groups:



Data from selected clusters

↳ all units who meet the criteria will be sampled.

usually used in interventional studies
↳ eg: immunization coverage

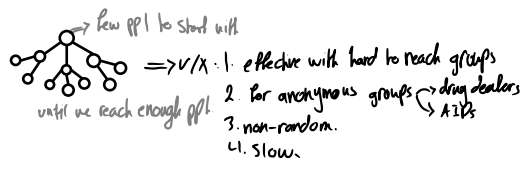
saves money & Time
 : $\sqrt{X}$ = easy to research
 ↳ $\times$ representation

non-probability Sampling:
 $\times$ Generalized findings

representative of the accessible population
 ↳ experts of their fields
 ↳ handpicking
 ↳ qualitative researchers use it

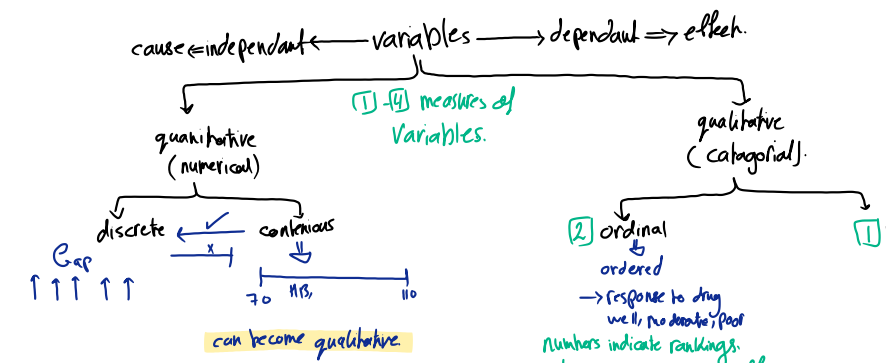
④ Purposive:

② snowball $\Rightarrow$



③ Quota:

Groups that you think are important
 $\hookrightarrow \sqrt{X} \Rightarrow$ 1. easy & quick
 2. saves money & time
 3. $\times$ representation.



② ordinal

ordered
 $\rightarrow$ response to drug well, no change, poor
 numbers indicate rankings.
 $\hookrightarrow$ we can't know the difference between them.
 $\Rightarrow$ mode, frequency, median.
 found in the middle
 V: sad, sad, ok, happy, happy

① nominal

unorganized
 Sex: m, f
 Blood type: A, B, O, AB
 $\Rightarrow$ least effective (no meaning)
 only mode, frequency
 most repeated
 how many times it was repeated

③ Interval (numerical).

- ordered, compared, no true Zero, no ratios.
 eg: Temp $\Rightarrow$ Temp = 0 $\Rightarrow \checkmark$ Temp
 Temp = 15 $\Rightarrow \times$ 15 $\neq \frac{1}{2}$ 30
 - can know the difference between ranks.
 $\Rightarrow$ mode, frequency, median, mean, can (+, -)

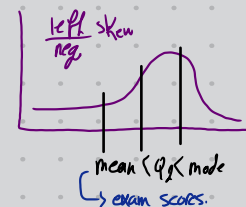
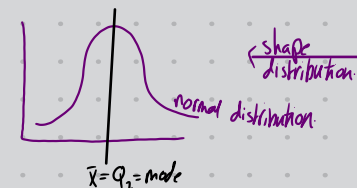
④ Ratio.

- ordered, compared, True Zero, $\checkmark$ ratio
 - can know the difference between ranks.
 $\hookrightarrow$ eg. weight = 0 $\Rightarrow$ no weight
 weight = 15 $\Rightarrow$ 15 $\neq \frac{1}{2}$ 30 kg
 $\Rightarrow$ mode, frequency, median, mean, can ($\times$, $\div$)

Discriptive statistics

parameter	Sample
μ	$\bar{x}$
σ^2	S^2
σ	S

parameter $\leftarrow$ population $\leftarrow$ Data $\rightarrow$ organized & classified into groups $\sim$ organized
 Sample $\leftarrow$ Statistics $\leftarrow$ $\rightarrow$ unorganized nor summarized $\sim$ raw data



Tendency

[1] Mean ($M, \bar{x}$)

$$\bar{x} = \frac{\sum x}{n}$$

$$\mu = \frac{\sum X}{N}$$

always present
 $+, -, 0$
 unique
 sensitive to extreme values

[2] Median

- middle value
 $\# \text{ odd} \Rightarrow$ always in the middle
 $\# \text{ even} \Rightarrow$ two middle digits
 median is a more robust statistic in the case of extreme values.

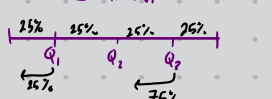
Frequency unaffected by extreme values
 unique

[3] Mode

- most repeated value
 not unique
 could have no mode

location

[1] Quartiles



[2] Deciles (10 parts)

[3] Percentile (100 parts)



Dispersion

(spread of value)

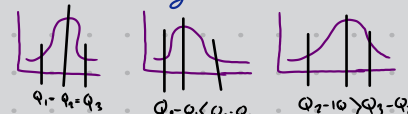
[1] Range (Max-Min)

- influenced by extreme values

[2] IQR ($Q_3 - Q_1$)



- not influenced by outliers



[3] standard deviation ($\sigma = \sqrt{\text{var}}$)

$\uparrow \text{std} \rightarrow \uparrow \text{variability} \Rightarrow$ not a v. reliable data

[4] Variation ($\frac{\sum (X - \bar{x})^2}{n-1} / \frac{\sum (X - \bar{x})^2}{N}$)

* Standard error (e):

$$\frac{S}{\sqrt{n}} \rightarrow \uparrow \text{error}$$

data classification

[1] Spatial

A	5
B	8
C	7
Total	20

A	5
B	8
C	7
Total	20

Frequency

$$\left(\frac{x}{T}\right) \text{ Relative frequency}$$

$$\left(\frac{x}{T} \times 100\right) \text{ percent}$$

[2] Temporal

over a period of time
 "chronological"

[3] Qualitative

Simple
 2 classes

Black White

Frequency

$$\left(\frac{x}{T}\right) \text{ Relative frequency}$$

$$\left(\frac{x}{T} \times 100\right) \text{ percent}$$

[4] Quantitative

frequency

Daily price (\$)	Frequency	largest - smallest
50-60	2	T
60-70	13	
70-80	16	
80-90	7	
90-100	7	
100-110	5	
Total	50	

Approximate Class Width = $(109 - 50) / 6 = 9.83 \approx 10$

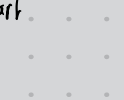
Manifold
 different classes

dark grey light grey soft white ivory

Frequency

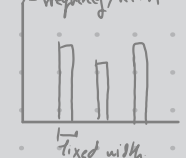
$$\left(\frac{x}{T}\right) \text{ Relative frequency}$$

$$\left(\frac{x}{T} \times 100\right) \text{ percent}$$

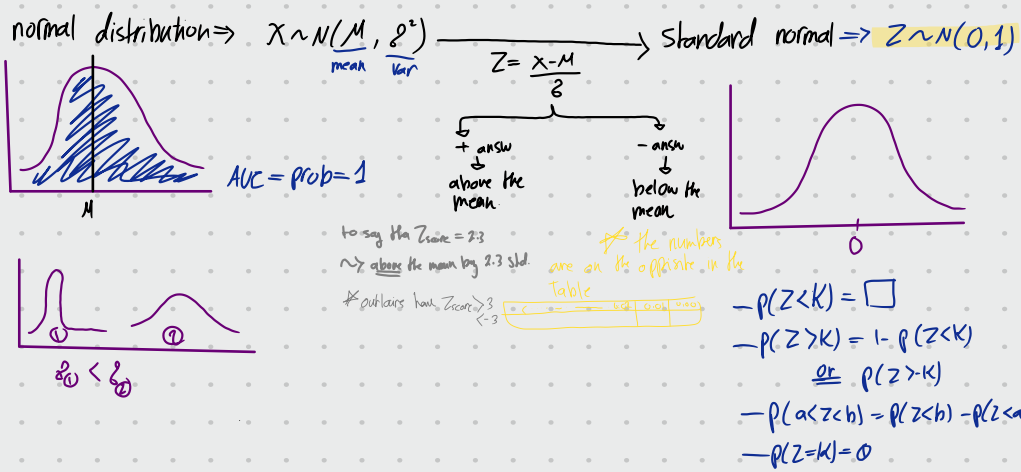


relative frequency

Bar graph



Cumulative Frequency Table			
Price (\$)	Cumulative Frequency	Relative Frequency	Cumulative Percent
≤ 50	2	0.04	4
≤ 60	14	0.28	28
≤ 70	27	0.54	54
≤ 80	43	0.86	86
≤ 90	50	1.00	100



Central limit theorem:-
 sample is distributed normally
 $\bar{X} \sim N(M, \frac{\sigma^2}{n}) \Rightarrow$ std. $= \frac{\sigma}{\sqrt{n}}$
 $\hookrightarrow$ from a normally distributed population.
 ① $n > 30 \Rightarrow$ because of skewed distribution



Inferential statistics $\begin{cases} \text{relation between variables} \\ \text{confidence interval} \\ \text{Test of hypothesis} \end{cases}$

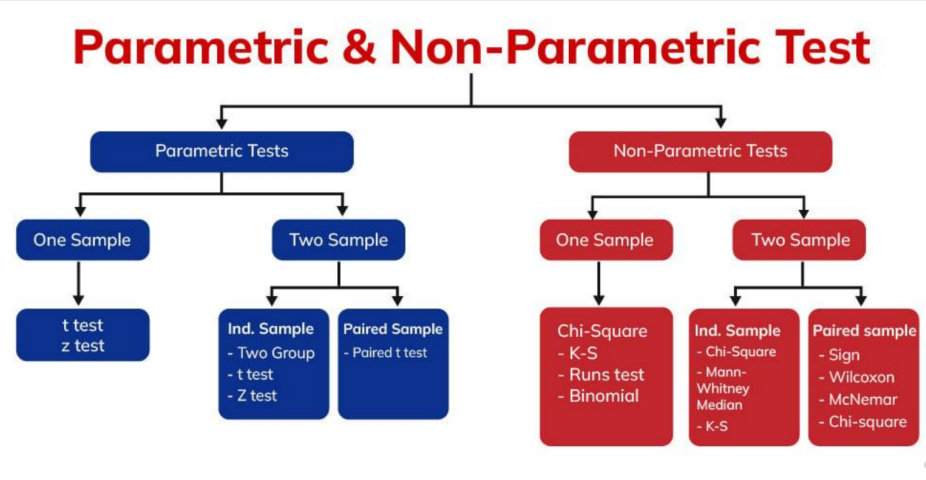
* Research hypothesis:
 * qualitative has no hypothesis
 - predict between dependent Vs independent variable
 - contain the population

$\leadsto$ difference $\rightarrow$ significant $\Rightarrow$ we reject the null hypothesis

H_0 : null $\leadsto$ hypothesized parameter "fail to reject"
 H_1 : alternative $\leadsto$ "accepted" if H_0 is rejected

	H_0 true	H_0 false
Rej H_0	Type I (α) error	power of test $(1 - \beta)$
Acc H_0		Type II (β) error

type I, α : $P(\text{rej } H_0 \text{ when it's true})$
 type II, β : $P(\text{acc } H_0 \text{ when it's false}) \Rightarrow$ more "wrong"
 $1 - \beta$ = the power of test
 rej H_0 when it's false



Difference between Parametric & Non-parametric test

	Parametric test	Non parametric test
1.	Used for ratio or interval data	For ordinal or nominal data
2.	Used for Normal distribution	Any distribution
3.	Mean is usual central measure	Median is usual central measure
4.	Information about population is completely known	No information available
5.	Specific assumptions made regarding population	Assumption free test

$\hookrightarrow$ T test $\begin{cases} \text{one sample} \\ \text{2 samples with unknown variances} \end{cases}$

$\hookrightarrow$ ANOVA $\leadsto$ 3 $\ll$ samples
 $\hookrightarrow$ data as an interval/ratio (quantitative).

Chi-square:-
 H_0 : Variables are independent (no association)
 H_1 : Variables are dependent (✓ association)

observed values $\Rightarrow$
 expected value = $\frac{\text{ROW} \times \text{col}}{\text{Total}}$

$\chi^2 = \sum \frac{(O - E)^2}{E}$ distance $0 \neq E \uparrow, \chi^2 \uparrow \Rightarrow$ rej H_0
 $0 \neq E \downarrow, \chi^2 \downarrow \Rightarrow$ acc H_0

* p value ≤ 0.05 $\Rightarrow$ we reject $H_0 \Rightarrow$ ✓ association
 significance

p value $> 0.05 \Rightarrow$ we fail to reject $H_0 \Rightarrow$ ✗ association