

Business Research Design and Methods

Research Methodology :-

→ Research :- Research is a logical and systematic search for new and useful information on a particular topic.

Research Methodology :- RM is a systematic way to solve a problem. It is a science of studying how research is to be carried out. Essentially, the procedure by which researchers go about their work of describing, explaining and predicting phenomena are called research methodology. It is also defined as the study of methods by which knowledge is gained. Its aim is to give the work plan of research.

→ Objectives of research Methodology

- To gain familiarity or achieve a new insight towards a certain topic
- To verify and test important facts
- To analyze an event, process or phenomenon
- To identify the cause and effect relationship
- To find solutions to scientific, non-scientific and social problems
- To determine the frequency at which something occurs

⇒ Types of research

Fundamental Research

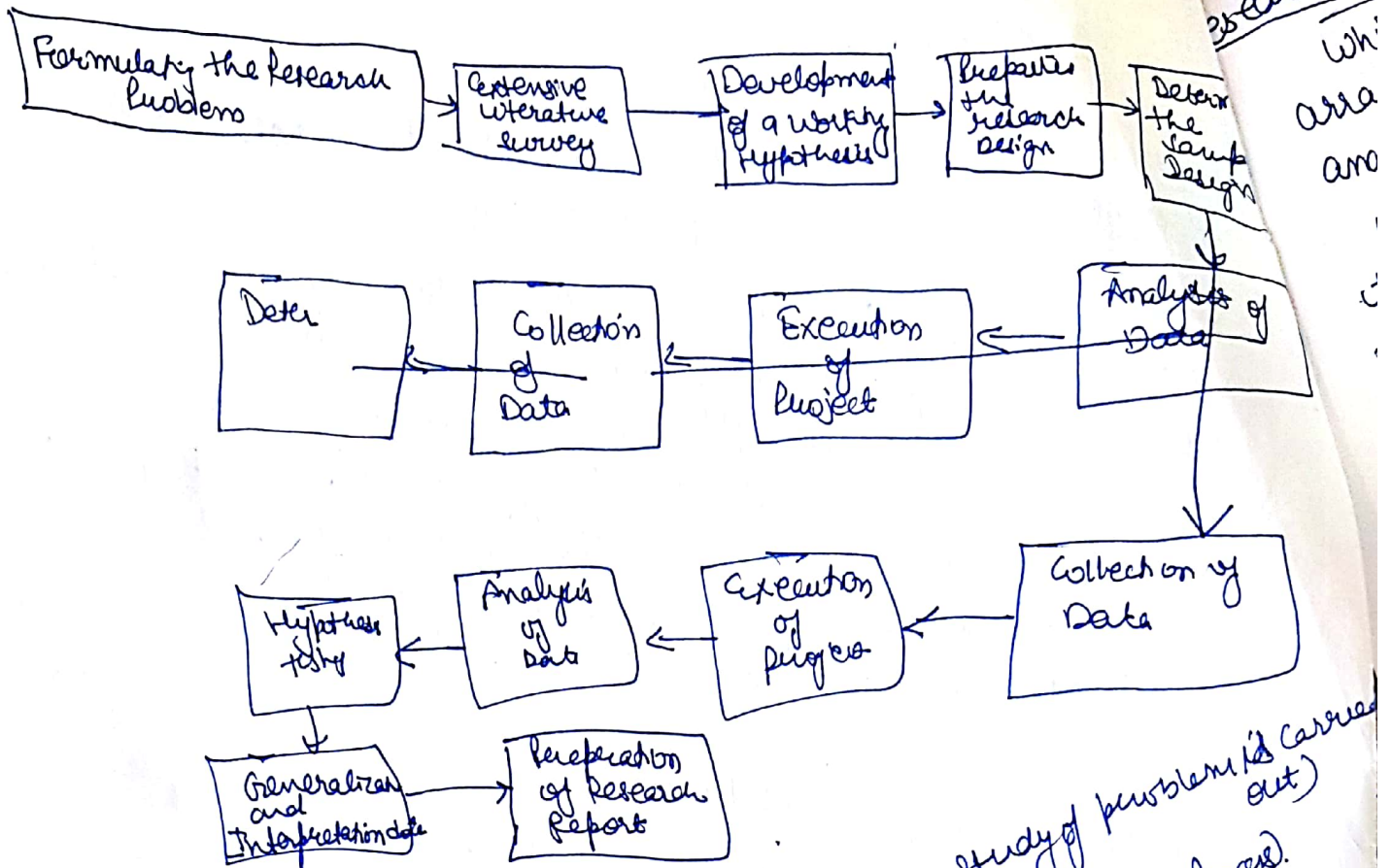
(Pure)
based on question
'Why things happen?'

- focus on one discipline
- descriptive study of research problem
- Reporting - technical language
- generalization & formulation of theory
- solve general problem

Pure Research (action)

- Based on question
'How things happen?'
- takes an interdisciplinary approach
- makes use of problem solving approaches to solve a specific problem
- simple language
- aims to find solution for problem
- deal with practical problem

→ Research Process



⇒ Nature of Research

- Research is a process (in which in-depth study of problem is carried out)
- Research is properly conducted helpful in decision making process.
- The research activities will help us in testing of hypothesis
- The research is a fact finding process.

⇒ Importance of Research

- Research inculcates scientific and inductive thinking and promotes development of logical habits of thinking and organization
- It aids in decision making
- Research plays a dynamic role in several fields and it has increased significance in recent times, it can be related to small business and also to the economy as a whole
- Most of the Government Regulations and Policies are based on and are a result of intensive research
- Its significance lies in solving various planning and operational problems
- It involves the study of cause and effect relationship between various variables & help to identify behavior, patterns, trends in certain variables

Research Design - The research design is the structure with which research is conducted. A research design is the arrangement of conditions for the collection and analysis of data in a manner that aims to combine to the research purpose. It constitutes the collection, measurement and analysis of data. Research design gives an outline of everything from defining the problem in terms of objective to final analysis of data.

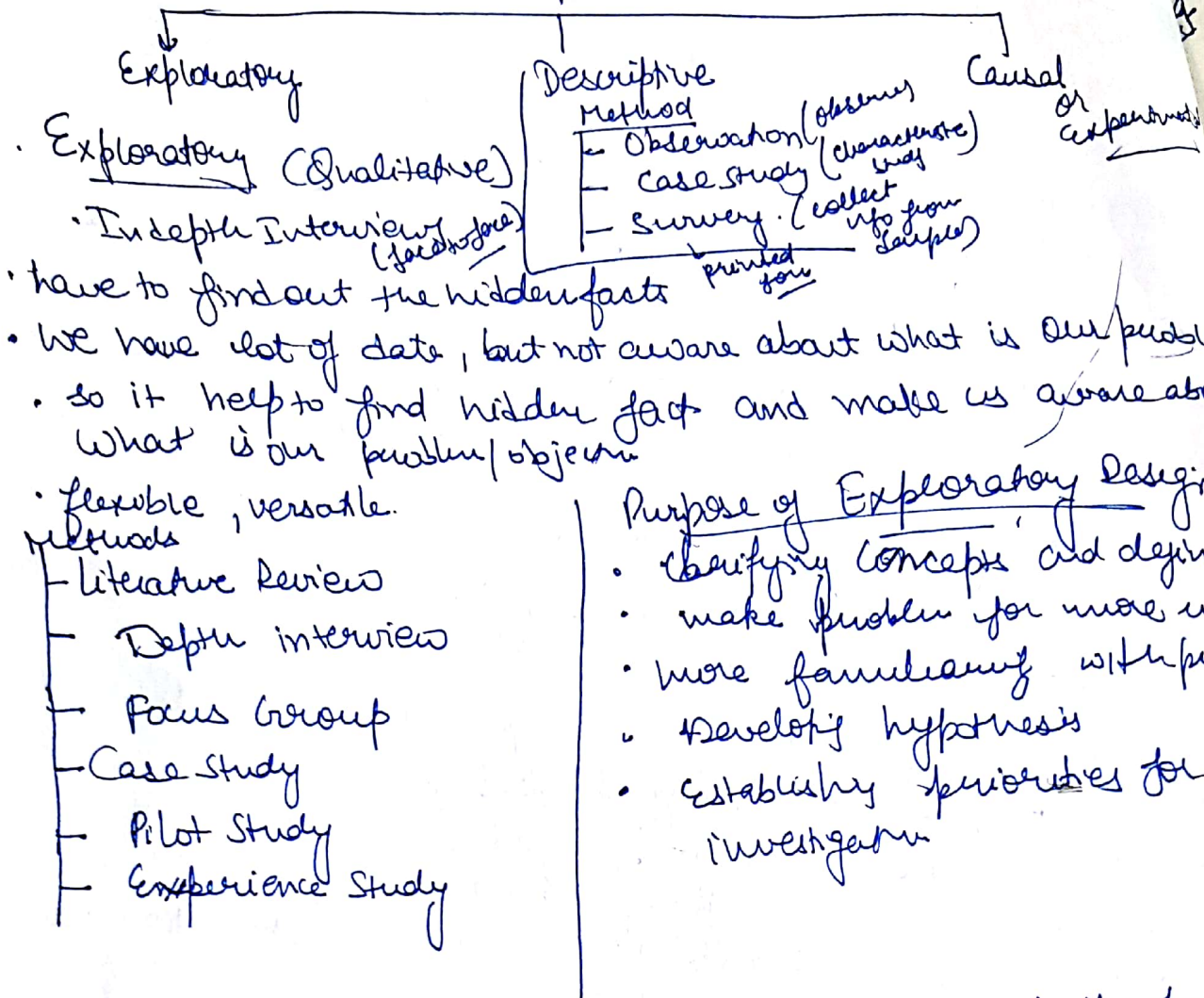
Contents of Research Design

1. Statement of Research objective i.e. why the research project is to be conducted
2. Type of data needed
3. Definition of population and sampling procedure to be followed
4. Time, costs, & responsibility specification
5. Methods, ways, and procedures used for collection of data.
6. Data analysis - tools or methods used to analyze data.
7. Probable output or research outcome and possible action to be taken based on those outcome.

Types of research design

1. Exploratory Research Design
 2. Descriptive Research Design
 3. Experimental Research Design.
1. Exploratory Research Design -
- discover & insights to generate possible explanation.
 - help in exploring problem or situation.
 - break problem into sub parts
 - help to form hypothesis
 - used to increase familiarity with problem under investigation
 - researcher new in area, or when problem is of different type

Types of Research Design



- have to find out the hidden facts
- we have lot of data, but not aware about what is our problem
- so it help to find hidden fact and make us aware about what is our problem/objective
- flexible, versatile.

Purpose of Exploratory Design

- Clarifying concepts and defining them
- make problem for more investigation
- more familiar with problem
- develop hypothesis
- establish priorities for further investigation

2. Descriptive Research Design :- describe characteristics of our population

- concerned with describing problem & its solution
- problem the cleared one, should be in hand
- rests on one or more hypotheses
- not flexible like exploratory
- directed to answer the problem

3. Causal or Experimental Research Design

- determine the cause the and effect relationship
- typical form of experiment
- impact of manipulations on independent variable or dependent variable
- dependent - sales, profit etc
- independent - price, promotion etc

eg sales increase with advertisement
effect cause

basics - experiment - to control conditions
 - manipulate variables

hypothesis - assumption - keep throughout research
 approve/disapprove

Experimentation - to tell the relationship
 between experiment
 +ve or -ve hypothesis

Basic issues

1. Manipulation of independent variables

- Experimental Treatment
- Experimental Group (artificial)
- Control Group (natural state)

2. Selection and measurement of dependent variables

3. Selection and assignment of Test Unit

i) Test Unit (Subject on which we have to conduct experiment)

ii) Sample selection error (biases)

iii) Random Sampling

iv) Matching

(Assign to each group)

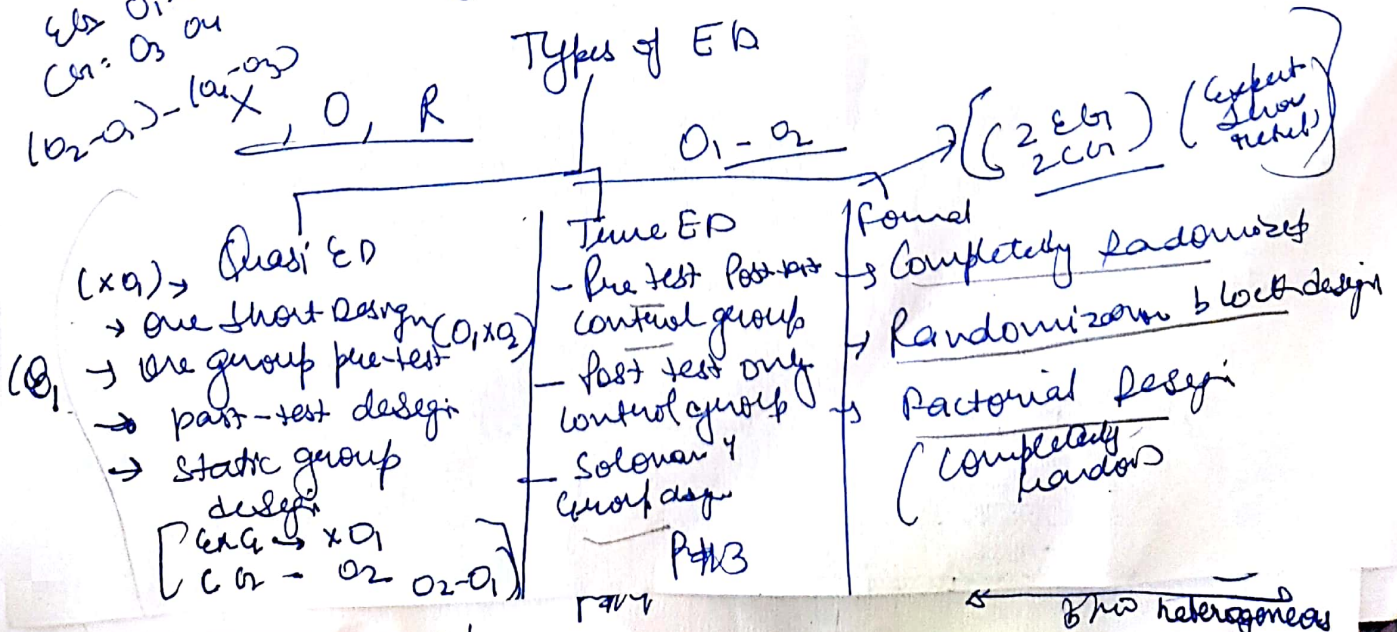
v) Control over extraneous variables

a) Constant error.

(comes every time we do research)

eg $O_1 \times O_2$
 $C_1 = O_3$ or
 $(O_2 - O_1) - (O_1 - O_3)$

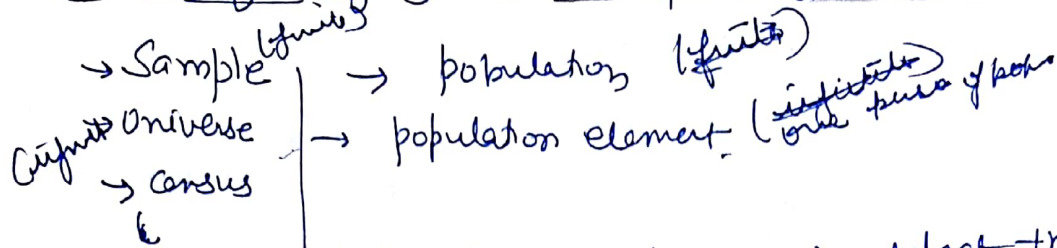
Types of EB



→ Homogeneous

[Sampling Design] → Concept, types and their application

Stages



Sampling is the process of how to select the sample.

Why Sampling is required

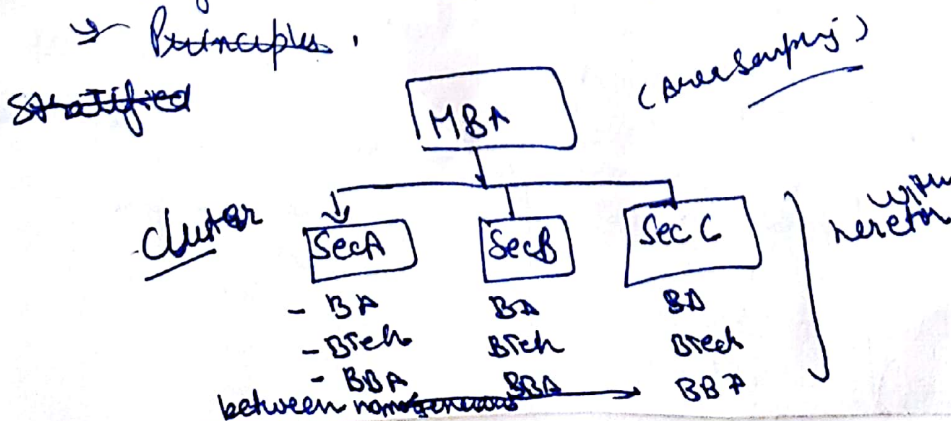
1. Reduce cost
 2. Reduce time
 3. Less effort
 4. Easy to structure Test unit
 5. Organised convenience
- Studies individual
 - Provides accuracy
 - Control unlimited data
 - Gives greater speed / helps to complete in stipulated time

Steps in Sampling process procedures

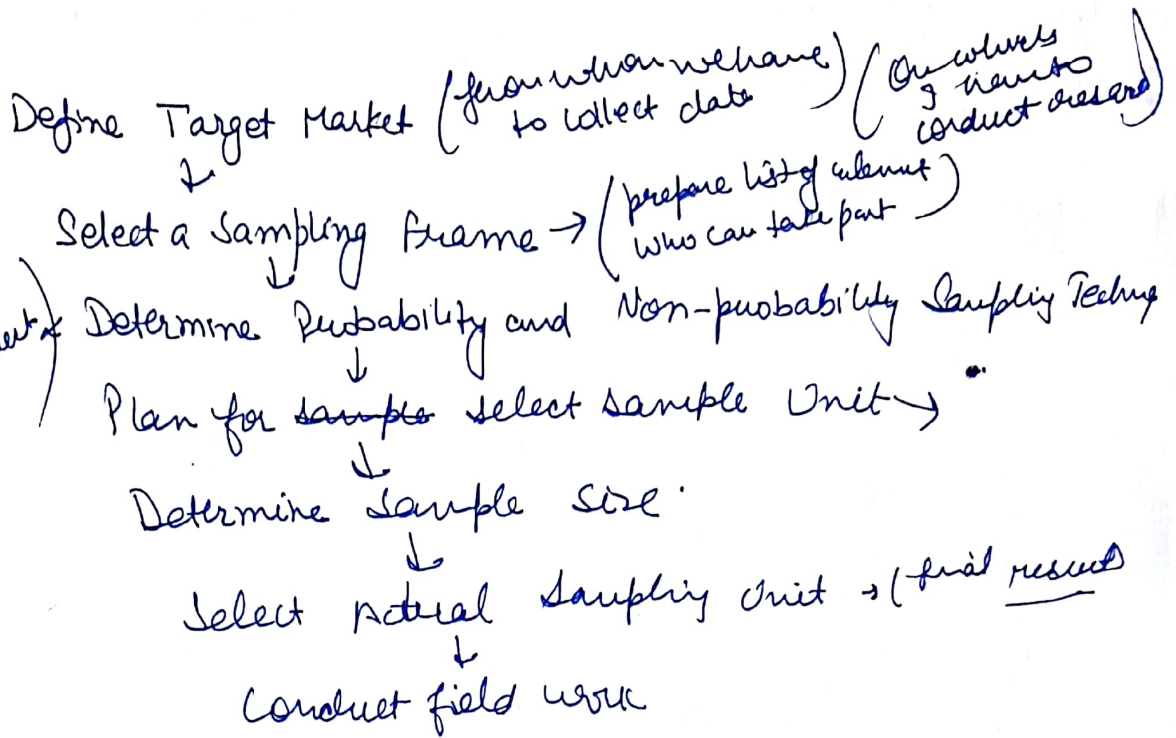
- Define the population (element, units, extent & time)
- Specify sampling frame (Telephone directory)
- Specify sampling unit (retailers, our product, students, unemployed)
- Specify sampling method / technique
- Determine sampling size
- Specify sampling size (optimum sample)
- Specify sampling plan
- Select the sample.

Two Principles

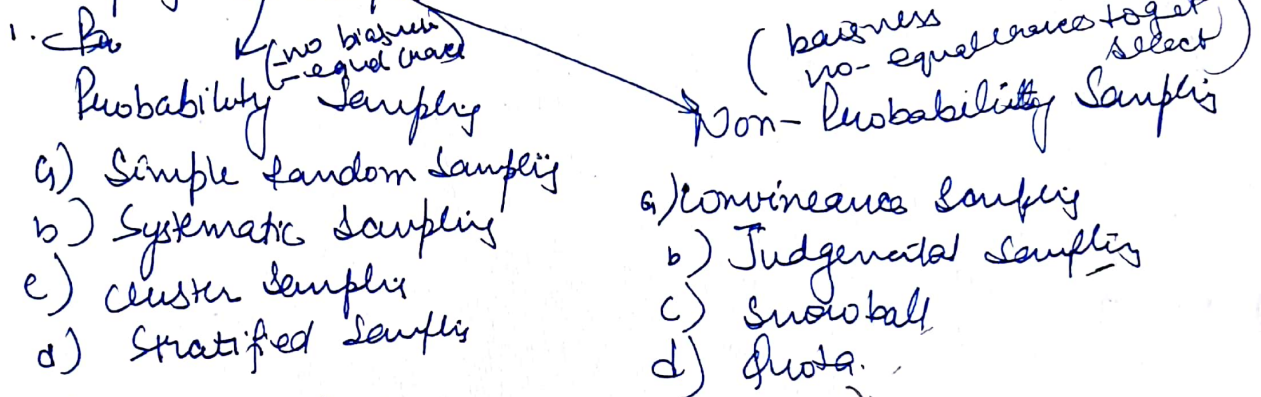
- Statistical regularity - random
- Large number - (Steadiness, stability & consistency)
- Principles:



Stages



Sampling Technique



Systematic Sampling : procedure

1st after 10 after 14
Next Next Next

Proceeds with selection of every k th item

Cluster → divide on certain criteria (within)

Stratified : homogeneous

Factor Influences Sample Size

Stratified - More representative

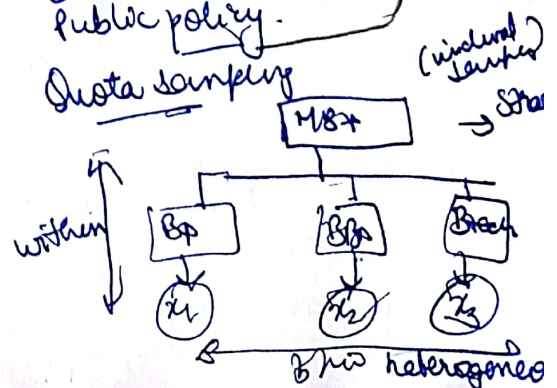
- Greater accuracy
- Greater geographical representation

Cluster

Judgmental

Small no. of sample
Business issues
Public policy

Quota sampling



Scaling techniques including likert, Thurston, Semantic Differential, Scaling techniques etc.

Measurement Scale :- It is the process of observing and recording the observations that are collected as the part of research effort.

Scale

A series of items which are arranged progressively according to value.

Types of Scale:-

1. Nominal Scale

3. Interval Scale

2. Ordinal Scale

4. Ratio Scale

Nominal Scale

- Simple Scale

- number or alphabets are written on it

eg no. on cricket shirt identity

Ordinal Scale

- According to value.

eg can say marks based

Interval Scale : simple mean

- not for objective & alternate

- based on difference make interval

make SD that

10-15, 15-20, 20-30, 30-40

Ratio Scale geometric mean

measure quality = absolute two relation between two

Criteria → Reliability & Validity

Likert Scale :-

respondent - agree or disagree make statement

Scaling Technique

(divide in groups marks on scale)

Rating Scale

→ Graphic Rating Scale (Good, bad, neutral)

→ Itemised Rating Scale (provide number for description)

→ Likert Scale

→ Semantic Differential Scale (upto 7 point bipolar scale)

→ Stapel's Scale (↑+ ↓-) (fill statement)

→ Thurstone Scale

Ranking Scale

→ Method of Rank order

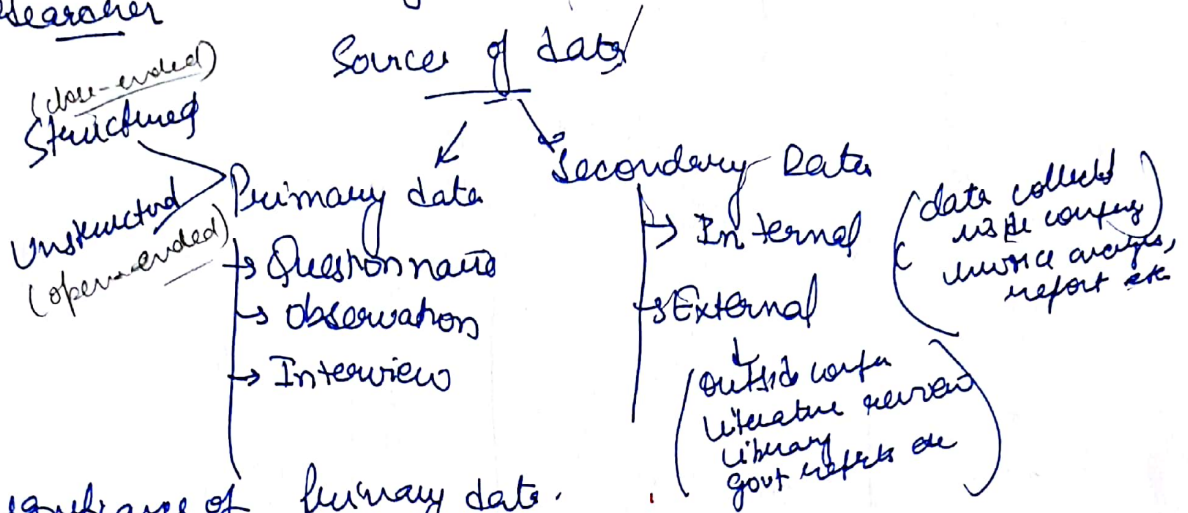
→ Method of paired comparison

(compare pairs ranks)

2 marks 1 mark 20 subject prefer

2 Techniques of Data collection: Primary and Secondary; Classification and Tabulation of Data

Data collection → It is the process to gather information about the relevant topic of research which is being done by researcher



Significance of Primary data.

1. Reliability
2. Availability of wide range of Techniques
3. Control.

Limitation

1. Cost
2. Time
3. Large Data

Significance of Secondary Data

1. Economical
2. Quickness
3. Availability
4. Useful in Exploratory research.

Classification of Data and Tabular Presentation

→ Qualitative (Attribute/Characteristic), Quantitative (Numerical data)

Temporal (time), Spatial (Geographical)

1) Qualitative

Location	Skilled	Unskilled
Rural.		
Urban		

Tabular form

2) Quantitative

Marks	No of students
0-10	3
10-20	7
20-30	12
30-40	15
40-50	22
	59

3. Temporal $\rightarrow$ Time

Where
time
is
main
characteristic =

Year	Sales (units)
2010	50,000
2011	70,000
2012	90,000
2013	1,00,000
2014	1,50,000

4. Spatial Classification + Place/location

Indian Students in different countries in year 2016

Country	No. of Students
USA	50,000
UK	15,000
Japan	17,000
Russia	2,000
Australia	25,000

Introduction to Statistics

Statistics is a mathematical science including methods of collecting, organizing and analyzing data in such a way that meaningful conclusions can be drawn from them. In general, its investigations and analyses fall into two broad categories called descriptive and inferential statistics.

Descriptive statistics deals with the processing of data without attempting to draw inferences from it. The data are presented in the form of tables and graphs. The characteristics of the data are described in simple terms. Events that are dealt with include everyday happenings such as accidents, price of goods, business incomes, epidemics, sports data and population data.

Inferential Statistics is a scientific discipline that uses mathematical tools to forecasts and projections by analyzing the given data. This is of use to people employed in such fields as engineering, economics, biology, the social sciences, business, agriculture and communication.

II Population and Sample

A population often consists of a large group of specifically defined elements. For example, the population of a specific country means all the people living within the boundaries of that country.

- Usually, it is not possible or practical to measure data for every element of the population under study. We randomly select a small group of elements from the population and call it a sample. Inferences about the population are then made on the basis of several samples.

Example:-

P#1

114

Quantitative and Qualitative Data

Data is quantitative if the observations or measurements made on a given variable of a sample or population have numerical values.

Example: height, weight, number of children, blood pressure, current, etc.

Data is qualitative if words, groups & categories represent the observations or measurements.

Example: Colors, yes-no answer, blood group.

Quantitative — discrete — if corresponding data values take discrete values.
— continuous — if the data values take continuous values.

Example of discrete data: number of children, number of cars.

Example of continuous data: Speed, distance, time, pressure.

⇒ Elements

- Collection of data
- Classification of data
- Presentation of data
- Comparison
- Analysis and interpretation

* Functions of Statistics

- Simplifies Complexity
- Measures periodic changes
- Proper presentation of facts
- Formulation of policies
- Study of relationship
- Comparison and forecasting
- Testing and hypothesis

Economics
Commerce (business)
Management (help in policy making)
Mathematics & Statistical
Natural science

* Scope of Statistics

- Nature of Statistics — Science or art
- Relation with others (subject)
- Limitation of Statistics.
 - deals with group only
 - only quantitative data
 - based on estimates and approximation
 - Applied for specific purpose
 - No cause and effect relation

* Relation with others

- Deals with only quantitative data.
- Doesn't study individual events — True on an average

Central tendency :- It is Central Value which represent the whole data

A measurement of data that indicates where the middle of the formation lies. ~~then Central~~

Central tendency is a way to describe what is typical for a set of data. Central tendency does not tell you specifics about the individual pieces of data, but it does give you an overall picture of what is going on in the entire data set. There are three major ways to show central tendency: mean, mode and median.

Mean:

Divide the sum of all given value by total number of values

Median :- Median is a value which divides the whole set into two

parts
steps

1. Set ascending or descending order
2. ~~Arrange the given data~~

(even) $M = \frac{N+1}{2}^{\text{th}} \text{ item}$

$N =$ No of item or observation.

Odd $\frac{\text{Sum of two closest centre item}}{2}$

$M = 4.5$

$M = \frac{4^{\text{th}} \text{ item} + 5^{\text{th}} \text{ item}}{2}$

Discrete series

- (i) Arrange data in ascending order
- (ii) Calculate CF $\rightarrow$ Cumulative frequency

$M = \frac{N+1}{2} \rightarrow \text{cf}$

$N = \Sigma f = \text{Total of frequency}$

Mode :- The mode is the most common number in a set of data.
e.g. the mode of 1, 2, 2, 3, 5, 6 is 2

* Measures of Dispersion

Dispersion :- Dispersion is the measure of the variation of the items.

- The degree to which numerical data tend to spread about an average value is called the variation or dispersion of the data.

⇒ Significance of measuring variation

1. To determine the reliability of an average
2. To serve as a basis for the control of the variability
3. To compare two or more series with regard to their variability
4. To facilitate the use of other statistical measures

⇒ Properties of a good measure of variation

- (i) It should be simple to understand
- (ii) It should be easy to compute
- (iii) It should be rigidly defined
- (iv) It should be based on each and every item of distribution
- (v) It should be amenable to further algebraic treatment.
- (vi) It should have sampling stability
- (vii) It should not be unduly affected by extreme items

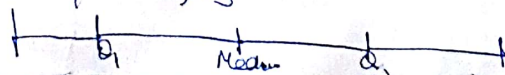
⇒ Measure or methods of dispersion

1. The range
2. The interquartile range and quartile deviation
3. The Mean average deviation
4. The Standard deviation
5. The Lorenz curve

⇒ Range = The difference between the greatest and lowest values in the data set.

• Merits and demerits of range.

→ Interquartile range = The difference between the lower (first) quartile, Q_1 and the upper (third) quartile, Q_3



Lower (First) Quartile, Q_1 - Value below which 25% (one-fourth) of the scores lie.

Upper (Third) Quartile, Q_3 - Value below which 75% (three-fourth) of the scores lie.

→ Box Plot :- Can be created from the lower quartile, median, upper quartile, IQR, and upper and lower Extremes

→ Upper Extrema → largest data value

→ Lower Extrema → smallest data value.

$$\text{Interquartile range} = Q_3 - Q_1$$

$$\text{Quartile deviation} = \frac{Q_3 - Q_1}{2}$$

↓
Gives average amount by which two quartiles differ from the median

$$\text{Coefficient of Q.D.} = \frac{(Q_3 - Q_1)/2}{(Q_3 + Q_1)/2} = \frac{Q_3 - Q_1}{Q_3 + Q_1}$$

→ Merits and Limitations of IQR

* Mean Absolute deviation :-

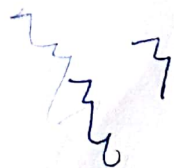
Make use of the absolute value to find the distance to find the distance each data point is away from the mean, median, founded by taking all absolute value differences then averaging them together.

⇒ Standard deviation :- Square root of variance.

→ steps to find standard deviation.

1. Find the mean
2. Subtract the mean from each number
3. Square the difference from step 2
4. Find the sum of square used in step 3
5. Divide by the number of values
6. ~~Divide by the number of value~~
7. Take the square root

Variance



P#13

The Coefficient of Variation (CV) is a measure of relative variability
the ratio of standard deviation to the mean (average).

$$\text{Coefficient of Variance} = \left(\frac{SD}{\bar{x}} \right) * 100$$

eg. "The SD is 15% of mean" is a CV

CV useful - compare results of two survey or tests that have different measures or values.

test A = 95%
B = 25%
more variation related to its mean

$$\text{CV for population} = \frac{\sigma}{\mu} * 100\%$$

$$\text{CV for sample} = \frac{s}{\bar{x}} * 100\%$$

Binomial distribution

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x}$$

$$SD = \sqrt{np(1-p)} = \sqrt{npq}$$

Poisson distribution

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

Normal distribution

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

$$z = \frac{x - \bar{x}}{s}$$

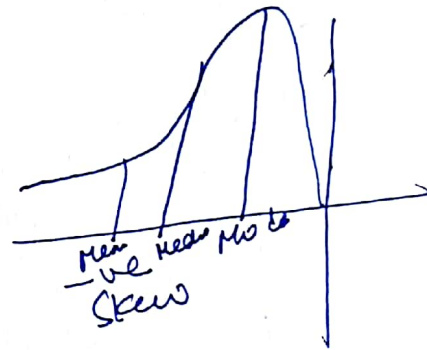
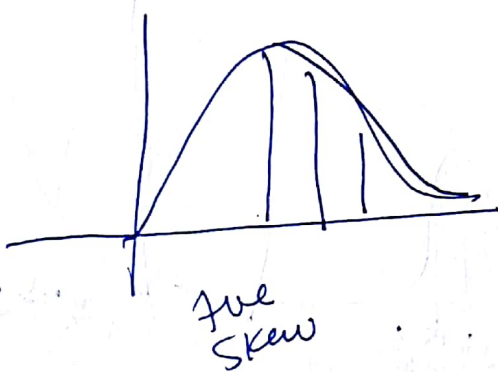
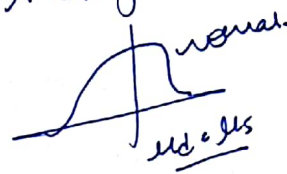
$$r = \sqrt{\frac{n-2}{1-r^2}}$$

$$\frac{(R_{xy} - R_{xz}) \cdot R_{yz}}{\sqrt{(1-R_{xz}^2)(1-R_{yz}^2)}}$$

Skewness and Kurtosis : Concept and Measures

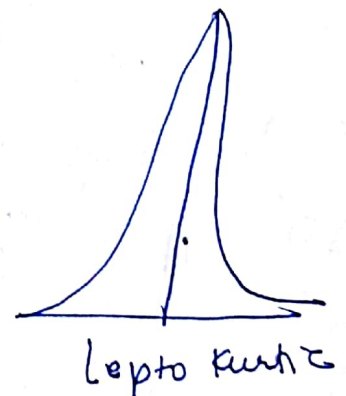
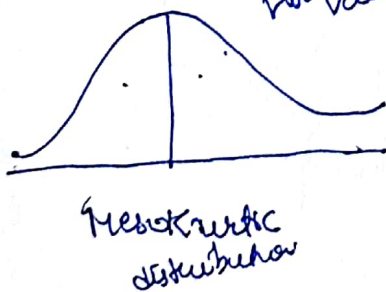
Skewness :- Skewness is a measure of symmetry, or more precisely, the lack of symmetry. A distribution, or data set, is symmetric if it looks the same to the left and right of the centre point.

Kurtosis is a measure of whether the data are heavy-tailed or light-tailed relative to normal distribution. That is, data sets with high kurtosis tend to have heavy tails, or outliers. Data sets with low kurtosis tend to have light tails, or lack of outliers. A uniform distribution would be the extreme case.



Kurtosis → peakedness of distribution

3 types - differing amount of peakedness
no. values bunched across the curve



Measures of central tendency

1. Moment coeff of Kurtosis

$$a_4 = \frac{\mu_4}{s^4}$$

SD → s

P#6

= 3 → distribution is Mesokurtic
 < 3 platykurtic
 > 3 leptokurtic

2. Percentile Coeff of Kurtosis (k)

$$= \frac{1}{2} (Q_3 - Q_1) \rightarrow \text{Semi-interquartile range.}$$

$$= \frac{\frac{1}{2}(Q_3 - Q_1)}{\frac{2}{3}(P_{90} - P_{10})}$$

$$y = 0.263 \rightarrow \text{Standard value} - \text{Mesokurtic}$$

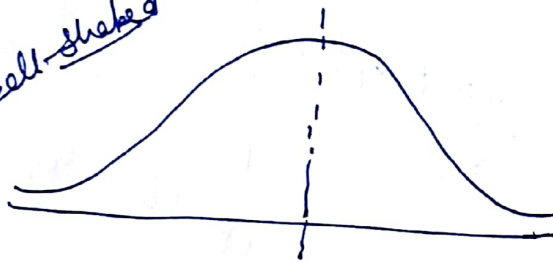
$$< 0.263 - \text{platykurtic}$$

$$> 0.263 - \text{leptokurtic}$$

Karl Pearson's Coefficient of Skewness

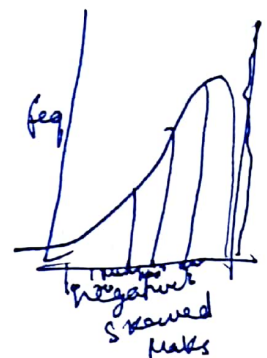
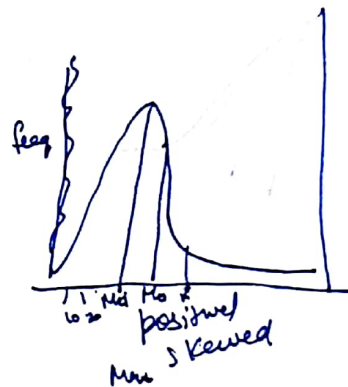
Skewness $\rightarrow$ ^{distribution} off-centre
Normal distribution

bell-shaped



$$\bar{X} = Md = Mode$$

$$\bar{X} = M_d = M_o$$



1. Karl Pearson's Pearsonian coefficient of Skewness

measure relative to dispersion

$$S_{kp} = \frac{\bar{X} - M_o}{\frac{S}{S.D.}}$$

$\rightarrow - \textcircled{S} \rightarrow +ve$

$$\bar{X} = M_o - 3(\bar{X} - M_d)$$

$$S_{kp} = \frac{3(\bar{X} - M_d)}{S}$$

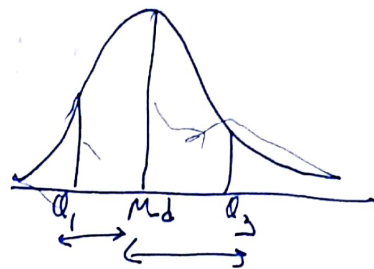
Coeff - ve (true skew)
= -ve (ve skew)

$$= 0 \text{ (i.e. } \bar{X} = M_o)$$

$0 \rightarrow$ Symmetric distribution

Quartile coefficient of Skewness

$$SKQ = \frac{(Q_3 - Md) - (Md - Q_1)}{Q_3 - Q_1}$$



3. Percentile coefficient of Skewness

Divide into 100 equal areas

Percentile coefficient of Skewness =
$$\frac{(P_{90} - Md) - (Md - P_{10})}{P_{90} - P_{10}}$$

4. Moment coefficient of Skewness

$$m_3 = \frac{\sum (X_i - \bar{X})^3}{n}$$

$$\mu_3 = \frac{\sum (X_i - \bar{X})^3}{n}$$

moment coefficient
$$\alpha_3 = \frac{\mu_3}{s^3}$$

0 → Symmetric
+ve → +ve skew
-ve → -ve skew

⇒ Test of Skewness

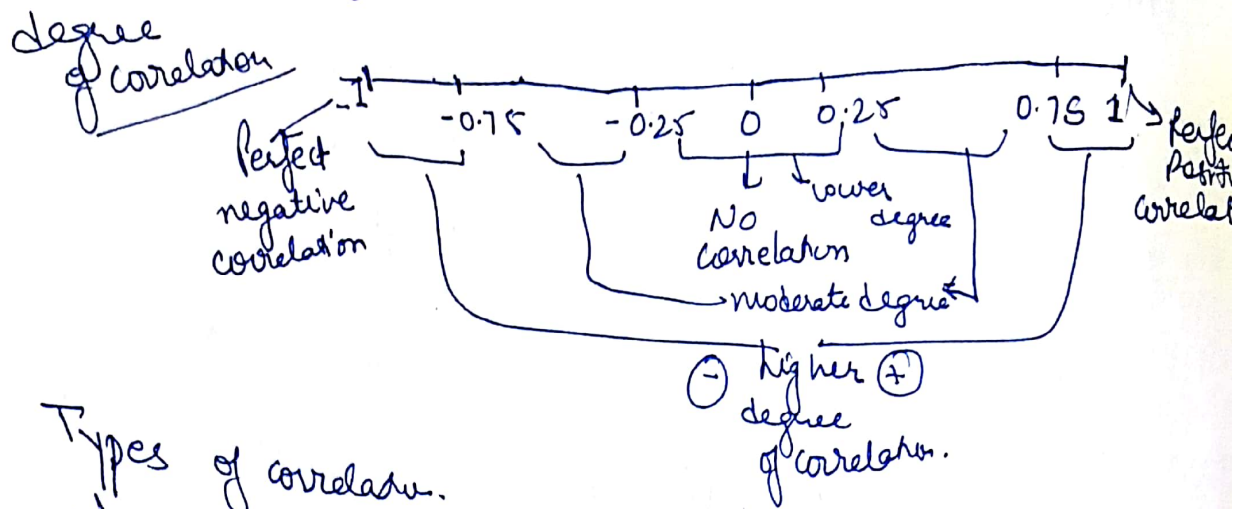
- 1) value of mean, median and mode (not equal)
- 2) Not bell shaped
- 3) Positive values of median are not equal to negative values of median
- 4) Quartile deviation is not different at the right position

Skewness - lack of symmetry
frequency distribution of series goes away from symmetry
- shows the direction and extent of variation
(limitation of dispersion) →

Tests

Correlation (r)

It is the relationship b/w two or more than two variables
the value of correlation lies between ~~to~~ -1 to +1



Types of correlation.

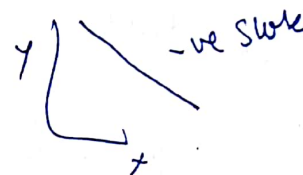
1.) Positive correlation.

$x \uparrow y \uparrow$
 $x \downarrow y \downarrow$



② Negative correlation

$x \uparrow y \downarrow$
 $x \downarrow y \uparrow$



③ Simple correlation

b/w two variable only

4.) Partial correlation

we have more than two variable, but we study only two variable at a time.

5.) Multiple correlation

we deal with atleast three variables simultaneously

Two variables
-1 to +1

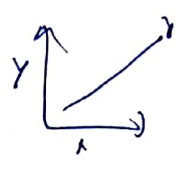
0.75 $\rightarrow$ Perfect Positive correlation

Example
mean of height of students after them

Linear Correlation

Situation in which any change in one variable leads to change in other variable in same ratio

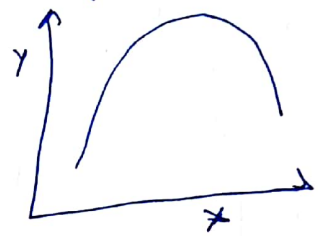
X	Y	X/Y
1	10	1/10
2	20	1/10
3	30	1/10
4	40	1/10
5	50	1/10



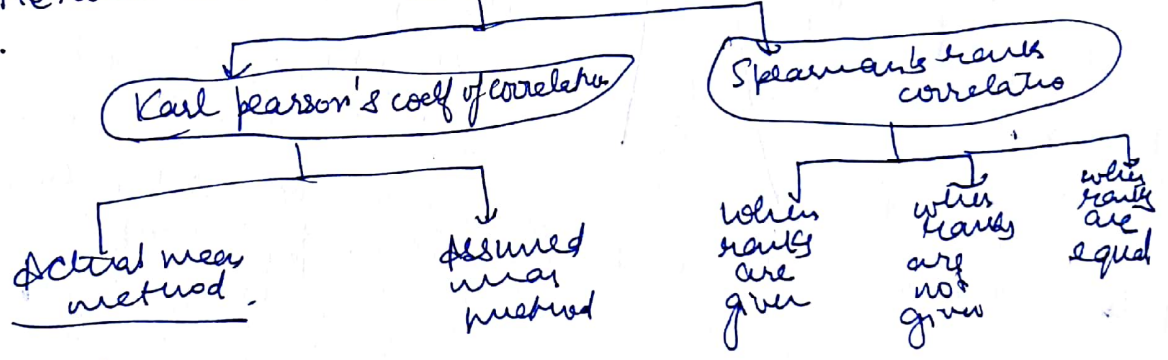
7. Non-linear correlation

change in variable, ratio tends to change

X	Y	X/Y
1	10	1/10
2	18	1/9
3	24	1/8
4	28	1/7
5	30	1/6



8. Methods to measure correlation



$$r = \frac{\sum xy}{\sqrt{\sum x^2 \cdot \sum y^2}}$$

$$x(x - \bar{x}) \text{ \& } y(y - \bar{y})$$

Assumed mean method

$$r = \frac{N \sum dxdy - \sum dx \cdot \sum dy}{\sqrt{N \sum dx^2 - (\sum dx)^2} \sqrt{N \sum dy^2 - (\sum dy)^2}}$$

$$dx = (x - A)$$

$$dy = (y - A)$$

A is assumed mean

Regression Analysis :-

Regression is a statistical measure used in finance, economics and other disciplines that attempts to determine the strength of relationship between one dependent variable (Y) & series of another change variable (independent variable).

two types

Linear Regression

Use one independent variable to explain or predict the outcome of the dependent variable Y

Linear Regression: $Y = a + bx + u$

Multiple Regression.

uses two or more independent variables to predict the outcome

Multiple Regression =

$$Y = a + b_1x_1 + b_2x_2 + b_3x_3 + \dots + b_kx_k + u$$

a = intercept

b = slope

u = regression residual

Y = the variable that you are trying to predict (dependent variable)

X = independent variable

(the variables that you are using to predict Y).

main objective of linear regression

Establish if there is a relationship between two variables
+ve relation - move together
-ve relation - one ↑ other ↓

- More specifically, establish if there is a statistically significant relationship b/w two
- Example: Income & Spending, wage and gender, student height
Exam Score

• Forecast new observations

- Can we use what we know about the relationship to forecast unobserved values?
- Examples: what will sales be over the next quarter?
what will the ROI of a new store opening be contingent on store attributes

→ Variable Role

- dependent variable (Y)
- Independent variable (X)

The Magic! A linear Equation

$$y = a + b \cdot x$$

y (m) x + b

state word.

$$y = \beta_0 + \beta_1 x$$

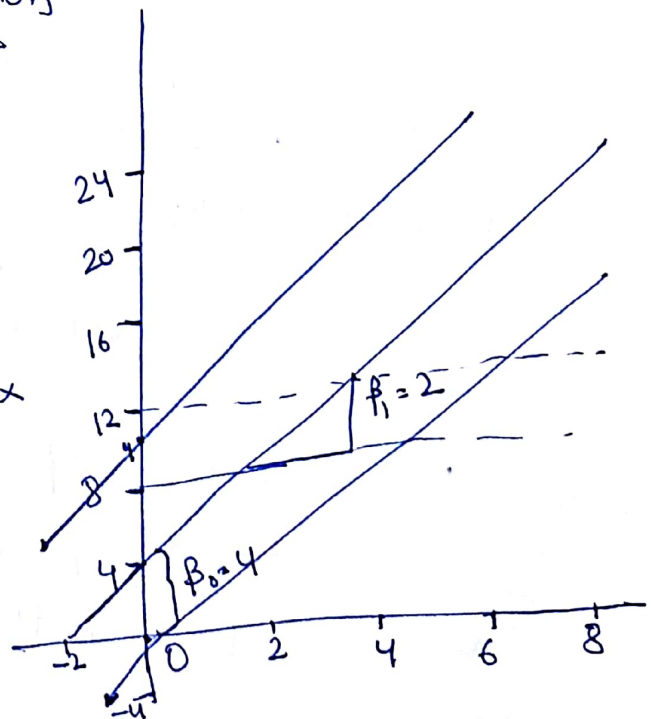
Intercept β_0 β_1 x
coeff. or slope of x

$$y = \beta_0 + \beta_1 x$$

$$y = 4 + 2x$$

$$y = 9 + 2x$$

$$y = -2 + 2x$$



Simple linear regression

$$y = \beta_0 + \beta_1 x + \epsilon$$

β_0 → constant or intercept
 β_1 → x's slope or coeff.
 ϵ → Error term.
 y → dependent variable
 x → independent variable

Slope (b)

$$b = r \frac{S_y}{S_x}$$

$$a = \bar{y} - b\bar{x}$$

$$r = \frac{\sum (x - \bar{x})(y - \bar{y})}{\sqrt{\sum (x - \bar{x})^2 \sum (y - \bar{y})^2}}$$

Elementary Probability Theory: Concepts, Definitions and Examples

Probability Theory

Sample Space (S) = Set of all possible outcomes. In an experiment.

Event = one particular outcome. $\in S$ or $\subseteq S$

Basic probability

How do we find probability of one (or more) things happen?
event

ex) experiment: flip a coin
 event: Get head.

Note: $P(H) = \frac{1}{2}$ or 50%

ex) roll a dice exp.
Get a 3 event
 $P(3) = \frac{1}{6}$

ex)

experiment: Roll 2 dice
event: sum of 7.

$$P(\text{Sum of 7}) = \frac{6}{36} = \frac{1}{6}.$$

General

$$P(E) = \frac{\text{no. of fav. outcomes}}{\text{Total no. of outcomes}}$$

symbolically

$$P(E) = \frac{n(E)}{n(S)}$$

$$0 < P(E) < 1$$

$P(E) = 0$	$E = \text{impossible}$
$P(E) = 1$	$E = \text{guaranteed}$

Calculate Probability of Compound Events
↳ OR, NOT, AND

Mutually Exclusive events

dependent on each other
if one happens, other will not happen
vice-versa

non-mutual
have common things

if $E_1 \cap E_2 = \emptyset$ = mutually exclusive
 $E_1 \cap E_2 \neq \emptyset$ = non-mutually exclusive



$$P(H) = \frac{FC}{\text{Total Cases}} \\ (FC + OC) \rightarrow P(H)$$

sure event, the probability = 1.

* Probability Distributions: Binomial, Poisson and Normal Distributions

A probability distribution is a statistical function that describe all possible value and likelihoods that random variable can take within a given range.

Discrete - countable/finite.

Random variable - ~~out~~ Outcome based on ~~change~~ chance.

A variable, X , that has value for each outcome for a procedure, that is determined by chance.

Probability distribution - A table that gives the probability for each value of a random variable.

~~Ex.~~ EX. DIE : roll a die

X	$P(X)$
1	$\frac{1}{6}$
2	$\frac{1}{6}$
3	$\frac{1}{6}$
4	$\frac{1}{6}$
5	$\frac{1}{6}$
6	$\frac{1}{6}$

discrete { Discrete random variable - A variable with a countable or finite no. of variable value.

eg. Student in a class

Continuous random variable - A variable with infinite number of possible value

Histon

STOCHASTIC FROM PROB. DIST.

HORIZONTAL: Value of random variable

VERTICAL: Probability

$E \pm$: WEIGHTED DIE

X	P(X)
1	.05
2	.15
3	.35
4	.30
5	.10
6	.05
	<u>1.00</u>

NOTE.

$$0 \leq P(X) \leq 1$$

AND

$$\sum P(X) = 1.00$$

$\Rightarrow$ Mean, variance & SD

Mean

$$\mu = \frac{\sum (X \cdot P)}{N}$$

$$\approx \left[\frac{X \cdot F}{N} \right]$$

Variance

$$\sigma^2 = \sum [X^2 \cdot P(X)] - \mu^2$$

$$\text{S.D. } \sigma = \sqrt{\sum [X^2 \cdot P(X)] - \mu^2}$$

$$\mu = \frac{\sum [X \cdot P(X)]}{N} \rightarrow \text{Mean expected value}$$

X	P(X)	X · P(X)	X ²	P · (X) ²
1	.05	.05		
2	.15	.30		
3	.35	1.20		
4	.30	.50		
5	.10	.30		
6	.05	.30		
	<u>1.00</u>	<u>$\sum (X \cdot P(X)) = 3.4$</u>		

$$\mu = 3.4$$

$$\sigma^2 = 1.34$$

$$\sigma = \sqrt{1.34}$$

P#11

Binomial distribution

$X =$ ~~no. of 11 year fluffy cat 5 trials~~
two possible outcomes
 Success/failure
 $p =$ probability of success

$3 = n =$ no. of trials

$2 = X =$ no. of times Successes (0 - n)

probability of success $\rightarrow p = .90$

prob. of not getting success $\rightarrow q = 1 - p = .10$

$$P(X) = \frac{n!}{(n-x)! x!} \cdot p^x q^{n-x}$$

$$P(X=8) = \frac{13}{(13-8)! 8!} (.5)^8 (.5)^5$$

$$p = .5$$

$$q = .5 = 1 - .5 \quad P(X=8) = .157$$

$$P(\text{at least } 8) = P(X \geq 8)$$

$$= P(8) + P(9) + P(10) + P(11) + P(12) + P(13)$$

$$P(\text{More than } 8) = P(X > 8) =$$

$$P(9) + P(10) + P(11) + P(12) + P(13)$$

$$P(\text{at most } 8) = P(X \leq 8)$$

value between 0 & 8

$$P(\text{fewer than } 8) = P(X < 8)$$

Binomial evaluates the probability of an event occurring several times over a given number of ~~trials~~ trials and given the event's probability in each trial.

Poisson distribution : helps in describing the chance of occurrence of a number of events in some given interval or given space conditionally that the value of average number of occurrence of even is known. $\rightarrow$ major & only condition of PPD

$$P(x, \mu) = \frac{(e^{-\mu}) (\mu^x)}{x!}$$

μ = avg no. of success

x = no. of occurring success

$e = 2.71828$

$\rightarrow$ Normal PD :

68% - one SD of mean

95% - two SD of mean

99.7% - 3 SD of mean

$M = Md = Mo$: equal
total curve area = 1

Hypothesis test

Types of Hypothesis :

Null Hypothesis H_0 :

Alternative Hypothesis H_a :

Errors in hypothesis test

Type I error - reject null Hypothesis

Type II error (β) - ~~the~~ failure to reject a false null hypothesis

P#12

Procedure in the hypothesis

- 1.) State the null hypothesis and alternate hypothesis
- 2.) Select the appropriate test statistics and level of significance.

testing hypothesis of mean = z-statistics

z-Statistic. σ (S.D) known.

$n > 30$

normal data

t-test σ unknown

normality or $n > 30$

2) Level of significance

0.10 - political polling

0.05 - consumer research project

0.01 - quality assurance

' α ' = probability of rejecting the null hypothesis when it is true

eg. $\alpha = 0.05 \rightarrow 5\%$ risk of concluding that a difference exists when there is no actual difference

3) State the decision rule

null hypothesis H_0 - accept or reject.

C.V of test - determined by α

4) Compute the appropriate test statistics & make decision.

z-Statistic, $z = \frac{\bar{x} - \mu}{\frac{\sigma}{\sqrt{n}}}$, t-Statistic $t = \frac{\bar{x} - \mu}{\frac{s}{\sqrt{n}}}$

5) Interpret the decision

chi-square test

don't depend on normal distribution to interpret findings

purpose

- i) to test hypothesis of no association between two or more groups, population or categories (ie check independence b/w 2 variables)
- ii) to test how likely the observed distribution of data fits with the distribution that is expected.

$$\chi^2 = \sum_{i=1}^n \frac{(O_i - E_i)^2}{E_i}$$

Expected
total = 29+24+22+19+21
+ 19+20+23+18
+ 20+23

O - observed freq.
E - Exp. freq.

Category	observed	Expected	Residual (Obs - Exp)	(Obs - Exp) ² /E	Component = (Obs - Exp) ² /E
Aries	29	21.333	7.667	58.78889	2.75549
Taurus	24	21.333	2.667	7.112889	0.333421
Gemini	22	21.333	0.667	0.442889	0.02142
Cancer	19	21.333	-2.333	5.442889	0.255139
leo	21	21.333	-0.333	0.110889	0.005198
Virgo	18	21.333	3.333	11.108889	0.52037
Libra	19	21.333	-2.333	5.442889	0.255139
Scorpio	20	21.333	-1.333	1.776889	0.08322973
Sagittarius	23	21.333	1.667	2.778889	0.13026451
Capricorn	18	21.333	-3.333	11.108889	0.520373
Aquarius	20	21.333	-1.333	1.776889	0.08332973
Pisces	23	21.333	1.667	2.778889	0.13026451
					5.094017203

→ A chi-square test for independence shows how Categorical Variables are related.

- Variation depends upon -
- how data collected
- how hypothesis ~~is~~ recorded

→ Chi-Square hypothesis test is appropriate if you have:

- Discrete Outcomes (Categorical)
- Dichotomous variable.
- Ordinal variable.

→ Test a chi-square hypothesis : steps

- Sample 1
1. 11 degrees of freedom
 2. Chi-Square test statistics of 5.094.

Step 1: Take the Chi-Square statistic
Find p-value in chi-square tables

- closest value for $df = 11$ & 5.094 is b/w 0.002 & 0.01
p-value 0.9256

Step 2: Use p-value.

Small p-value (1% - 5%)

large p-value (92.56%)

5% reject
0.05 - level

- reject null hypothesis

- not reject null hypothesis

→ The Chi-Square distribution:

- Special case of gamma distribution

Let's have random sample taken from ND.

1st χ^2 -distribution is the distribution of the sum of these random samples squared

$df(x) =$ no. of samples being summed.

eg. 10 sample from ND
Then, $df = 10$

(also known as χ^2)

- always right skewed.

high degree of freedom looks like N.D.

P#15

(continuation on next P#14)

Introduction to Analysis of Variance (ANOVA)

Introduction

- Different Types of ANOVA
- Assumptions of ANOVA
- Hypotheses in ANOVA
- Main Effect / Interaction Effects
- Post Hoc Test
- The F Distribution

ANOVA is a statistical Method used to compare the means of two or more groups

- Factors (Variables)
- Levels

0mg	50mg	100mg
9	7	4
8	6	3
7	6	2
8	8	4
8	8	4
9	7	3
8	6	2

Factor = "Dosage"

Levels - "0mg", "50mg", "100mg"

Types of ANOVA

- One way ANOVA - one factor with at least two levels, levels are independent
- Repeated-Measures ANOVA - one factor with at least two levels, levels are dependent.

Day 1	Day 2	Day 3
9	7	4
8	6	3
7	6	2
8	8	4
8	8	4
9	7	3
8	6	2

Ans.

Factorial ANOVA — Two or more factors (each of which with least two levels), levels can be either independent, dependent or both (mixed).

	Day 1	Day 2	Day 3
Men →	9	7	4
	8	6	3
	7	6	2
Women →	8	7	3
	8	8	4
	8	7	3

Assumptions in ANOVA

1. Normality of Sampling Distribution of Means
The distribution of sample means is normally distributed.

2. Independence of Errors

Errors between cases are independent of one another.

3. Absence of Outliers.

Outlying scores have been removed from the data set.

4. Homogeneity of Variance

Population variances in different levels of each independent variable are equal.

Hypothesis:

ANOVA with one factor ("A", three levels):

$$H_0: \mu_1 = \mu_2 = \mu_3$$

← Main effect

H_1 : not all means are equal.

ANOVA with two factors ("A" & "B", each with 3 levels):
with $H_0 = \mu_{A1} = \mu_{A2} = \mu_{A3}$
then $H_1 =$ not all means are equal.

$$H_0 = \mu_{B1} = \mu_{B2} = \mu_{B3}$$

$H_1 =$ not all means are equal

H_0 : interaction absent $\leftarrow$ Interaction Effect
 $H_1 =$ interaction present. $(A \times B)$

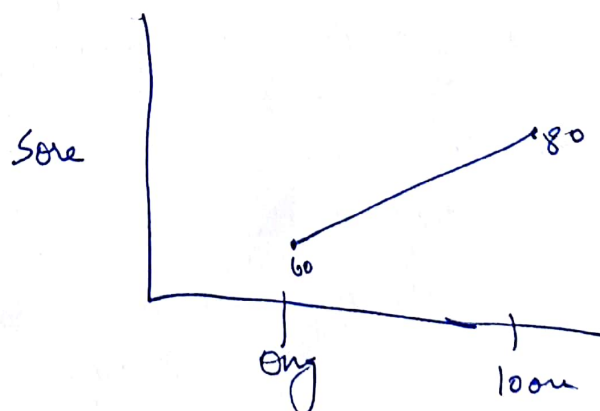
Main Effects in ANOVA

$$H_0 = \mu_{A1} = \mu_{A2} = \mu_{A3}$$

H_0 : not all means are equal

Pretend we are comparing the test score of people who have received a medication (100mg dosage group) and people who have not received a medication (0mg dosage group). The 0mg condition has a mean of 60, while the 100mg condition has a mean of 80. This could be represented in a graph like this:

0mg - 60
100mg - 80



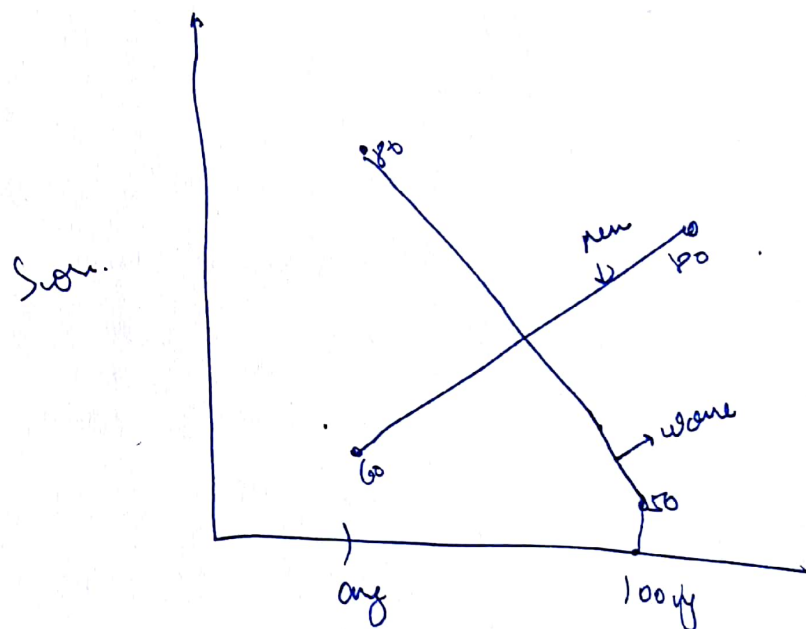
#. Interaction Effects in ANOVA

$$H_0: \mu_1 = \mu_2 = \mu_3 \quad \left| \quad H_0: \mu_B = \mu_2 = \mu_3 \quad \left| \quad H_0: \text{interaction absent} \right. \right.$$

$$H_1: \text{not all means are equal} \quad \left| \quad H_1: \text{not all means are equal} \quad \left| \quad H_1: \text{interaction absent} \right. \right.$$

Here we have a 2-way ANOVA with two factors: dosage (0mg & 100mg) and gender (men & women). In the 0mg dosage condition, men have a mean of 60 while women have a mean of 80. In the 100 mg dosage condition, men have a mean of 80 while women have a mean of 50. This could be represented in a graph like this.

	M	W
0 mg	60	80
100 mg	80	50



⇒ Post-Hoc Analysis in ANOVA

$$H_0: \mu_1 = \mu_2 = \mu_3$$

$H_1: \text{not all means are equal.}$

If we reject the null hypothesis, all we know is that there is a difference somewhere among the groups.

Additional test called post hoc test can be done to determine where difference lies.

ie F distribution in ANOVA.

When doing an ANOVA, we calculate an "F" statistic. It is similar to other statistics such as "Z" & "t".

$$F = \frac{\text{Treatment Differences} + \text{Random Differences}}{\text{Random Differences}}$$

If there are no treatment differences (that is, if there is no actual effect), we expect F to be 1. If there are treatment differences, we expect F to be greater than 1.

The F statistic has its own one-tailed distribution, much like how the 'z' and 't' statistics have their own separate distributions.

#. Interaction Effects in ANOVA

Chi-square test (contingency) (a statistical method assessing the goodness of fit between a set of observed values and those expected theoretically).

= Uses

- Confidence interval estimation for a population standard deviation of a normal distribution from a sample standard deviation
- Independence of two criteria of classification of qualitative variables
- Relationship between categorical variable (contingency table)
- Sample variance study when the underlying distribution is normal
- Tests of deviations of differences between expected and observed frequencies (one-way tables).
- The Chi-square test (a goodness of fit test).

⇒ Chi distribution

It distributes square root of a variable distributed according to Chi-square distribution with $df = n - 1$ degrees of freedom near probability density function of:

$$f(x) = \frac{1}{2^{(1-n/2)} \Gamma(n/2)} x^{(n-1)/2} e^{-(x^2)/2}$$

⇒ Chi-Square p-values

- tells us the result is significant or not.

Two things needed:

1. Df (no. of categories - 1)

2. alpha level (α).

usual $\alpha = 0.05$ (5%)

could be 0.01 or 0.10

tell us deviation b/w O & E values is by chance or because of some factor