

① Intro to BA : Role of Analytics for Data Driven Decision Making

⇒ Business Analytics means ⇒ "means utilizing tools, techniques, and procedure to analyze past business data in order to gain insight and help in decision making & problem solving"

① Data is Everything, Analytics is necessary for survival. Analytics can be used for process improvement (eg. using it to reduce time to deliver order), problem solving (predicting customer cancellation & frauds), decision making (switch to a new market etc.).

② Introduction to concept of Big data Analytics

Big Data ⇒ refers to extremely large/complex sets of data that cannot be easily managed, processed or analysed with traditional data processing tools.

⇒ The goal of big data analytics is to extract valuable insights, patterns and knowledge.

⇒ Big data is characterized by 3 Vs :-

- ① Volume (large amount of data)
- ② Velocity (High speed at which data is generated & processed)
- ③ Variety (diff. type of data like structured / ☒ unstructured)

Sources of big data

- ① Social media data
- ② Satellite data
- ③ Banking, financial, healthcare data
- ④ Entertainment data
- ⑤ Research data, data from sensors, meters etc.

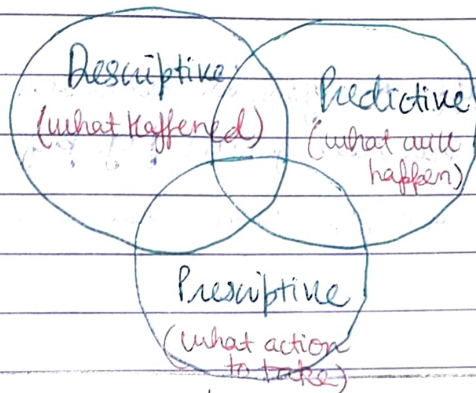
① Structured data :- Has clear structure and follows a regular sequence. Organized formatted

in a way which is easily readable by both humans & machines. It generally follows tabular structure. Relational Database entries :- Employee data in company database with columns Name, ID, Department, salary. Spreadsheet :- Data in Excel or similar software like financial data with date, expense etc.

(Relational Database is used to handle complex, large datasets while spreadsheets are simpler, used for basic data manipulation & smaller datasets and calculations)

② Unstructured data :- lacks formal structure. Doesn't conform to fixed schema and is often human generated.
Example word docs, pdfs, video clips, images, audio recordings etc

③ Types :- Descriptive Analytics, Predictive Analytics and Prescriptive Analytics



Descriptive is Numerical portion :- (covered in upcoming pages (one shot))

Descriptive	Predictive	Prescriptive
↓ Summarizes historical data to understand what happened in the past. [Tells you what happened in the past] Eg. You run a e-commerce business, it will take involve looking at the past data (sales) and answer ques. like :- ① what were our best selling products last year? ② which months had the highest sales?	↓ use historical data & statistical algorithms to make prediction about future outcomes. [Take educated guess of what will happen in the future] Eg. Continuing with e-commerce one, it would forecast future sales or specific trends like sale of certain clothes ↑ during certain seasons.	↓ goes a step further and recommends actions to optimize or improve outcomes. gives advice/ action to achieve desired result [Specific action to achieve best outcome] Eg. Predictive analytics told us sale would go up during this period ↓ so prescriptive analytics might recommend ↑ing inventory of those products, running targeted ads or discounts.

Digital Media Analytics

⇒ Both play crucial role in digital marketing and online business. Strategies of business and organizations by leveraging insights from these analytics, entities can make informed decisions to optimize their online performance & user engagement.

③ WEB Analytics

⇒ It provides insights into performance of website ⇒ That how users are interacting with their online content. Tools like Google Analytics is used for this.

Key Aspect of this are:- Traffic Analysis, User Behaviour (Time they spend), Conversion Tracking (Sale/Purchase/Lead generate), Page Performance (Load time), A/B Testing (making two pages & one page with more images and seeing which performs better).

③ Social Media Analytics

⇒ It involves collection and analysis of data from social media platforms.

Its Key Aspects include:- Audience insights, Engagement metrics, Sentiment analysis (positive, neutral, negative), Influencer impact, Campaign performance etc.

⇒ Tools used are Google Analytics, Facebook Insights/Twitter Analytics etc.

⑤ Overview Machine Learning Algorithms

And Statistical Software Packages

③ Machine Learning Algorithms

They are set of rules/statistical model that enable computer to learn and make predictions without being explicitly programmed to do so. They learn from data and improve over time.

Ex. A very commonly used example of Machine Learning Algo is "Spam Email filter", which often uses a technique called Naive Bayes classification.

• Example:- Spam Email filter

• Problem:- You receive lots of spam emails. You want to automatically filter out rather than check each manually.

• Machine Based Algorithm:-

③ Training:- Give dataset of email labeled as spam or not spam. It learns patterns/probabilities of certain words being used in spam email like offer, free, discount, urgent and calculate probability of it being spam.

→ This is how machine learning can automate tasks.

③ How Machine Learning Operates:-

⇒ By training model on a dataset, allowing algorithm to identify patterns and rule "Spam data". The trained model can then be used to make decision when presented with new/unseen data.

③ Statistical Software / Software Packages

⇒ They are the type of computer program designed to perform statistical analysis on data. Helps, researchers, business analyst to make sense of data by providing various tools.

• Benefits:-

① Helps in Data Analysis:- Analyse trend/pattern "not" in large data.

② Visualization:- Representing visual representation for more easy interpretation of large data.

③ Modelling:- Helps in the creation of statistical models.

Popular Software:- SPSS (Statistical Package for Social Sciences), R (open language, statistical analysis software), MATLAB (Matrix Laboratory)

Descriptive Statistics

①

Data

Based on Data Structure

Structured Data

(Data in a matrix labelled as rows & columns)



Unstructured Data

(Data not in matrix form, no fixed structure)



Time Series Data: Taking price of some thing at regular time intervals like every 15 days to see how it's value fluctuates. E.g. tracking stock price of a company for every 15 days.

Based on Scale of Measurement

Cross sectional data

Taking picture at one moment at time (snapshot)

E.g. survey 100 people on street today and ask them fav. colour. If giving you picture of people preference at that particular moment.

Panel Data (aka longitudinal)

Types of Data Measurement Scales

Nominal

In it data refers to names which can be given a numerical value to represent them (no measuring means rank order)

Example - 1

Gender (with)

In it data has clear order or ranking but intervals between categories may not be consistent.

Example - 2

Ranking

In it data has clean order and consistent intervals between them.

Example - 1

Temperature

In it data has clear order, consistent intervals and true zero point.

Example - 1

Weight

Ordinal

Interval

Ratio

Mean (Average)

It is the mathematical average value of data

$$\text{Mean} = \bar{X} = \frac{x_1 + x_2 + \dots + x_n}{n} = \frac{\sum_{i=1}^n x_i}{n}$$

If, with frequency,

$$\text{Mean} = \frac{\sum_{i=1}^n f_i x_i}{\sum_{i=1}^n f_i}$$

Imp. property of mean is that summation of deviation from observation from mean is zero.

$$\sum_{i=1}^n (x_i - \bar{X}) = 0$$

③ Population and Sample

Population: - The entire grp you are interested in studying.

Sample: - smaller grp, selected from population to represent & make inferences about the larger grp.

E.g. You want to find avg. height in college. You cannot ask 8000 students. You select 10 students from each tier and infer about them it about all.

④ Measures of Central Tendency

★ "Median" is the thing in descriptive stats which is not calculated based on the entire data.

★ Median

⇒ Divides data into two equal halves. Portion of observation below & above median will be 50%.

Individual Series

⇒ Arrange the data, find middlemost value.

Bivariate Series:-

Odd no. of observations,

$$\left(\frac{n+1}{2}\right)^{\text{th}} \text{ term} = \text{Median}$$

Even no. of observation,

$$\frac{\left(\frac{n}{2}\right)^{\text{th}} \text{ term} + \left(\frac{n}{2} + 1\right)^{\text{th}} \text{ term}}{2} = \text{value} = \text{Median}$$

Continuous Series



$$\text{Median} = l + \left(\frac{\frac{N}{2} - C.F.}{f} \right) \times h$$

① Calculate C.F.

★ l = lower limit of median class

★ f = frequency of median class

★ C.F = C.F of class above median class

★ h = class width

⇒ Median class is the class which has the highest frequency.

★ Mode

⇒ Mode is the frequently occurring value in dataset.

★ ⇒ Mode is the only measure of central tendency which is valid for Nominal (qualitative data). Since, the mean and median for nominal data is meaningless.

⇒ Eg. There is a data of marital status of customers namely (a) married (b) single (c) with children (d) without children. Mean, Median are meaningless here, but mode can tell what is the most occurring customer type.

Individual/Bivariate Series

→ Series with highest frequency.

Continuous Series



$$\text{Mode} = l + \left(\frac{f_1 - f_0}{2f_1 - f_0 - f_2} \right) \times h$$

★ "Modal class" - class interval with highest frequency.

l = lower limit modal class

f_0 = frequency of class above modal class

f_1 = frequency of modal class

f_2 = frequency of class below modal class.

h = class width

$$\# \text{ Mode} = 3 \text{ Median} - 2 \text{ Mean}$$

⑤ Percentile, Deviate and Decile

⇒ Percentile, Quartile and Decile are used to tell the position of the observation (and its value) in the dataset.

- Percentile Score can be used in identifying posⁿ of a student in class and in asset management.

$\Rightarrow P_x = 0$ ← gives value of data which $n\%$ of data is below that value.

$$p_{xi} = \frac{x_i(n+1)}{100} \text{th obs here, } n = \text{no. of observations in data}$$

Quartiles

- Special type of permeable which divide plate in 4 equal parts

eg. $Q_1 = 1.25$ = contains first 25% of data

$Q_2 = P_{50} = 11.50\% = \text{Median}$

Ex $P = \text{first-draw} \Rightarrow P_{10} = \text{contains } 10\% \text{ of data}$
 $P_{20} = 2^{\text{nd}} \Rightarrow P_{20} = \text{contains } 20\% \text{ of data}$

$P_{20} \approx 2^{nd} \Rightarrow P_{20} = \text{contains } 20\% \text{ of data}$

$Q_1 = P_{100}$ contains 100% of data

beides

③ special type of parentile which divide into 10 equal parts.

eg $P_{10} = \text{girls + deids} = P'_{10} = (\text{contains } 10\% \text{ of data})$

$P_{20} \approx 2^{nd}$ " $= P_{20} =$ contains 20% of data set.

99

22	39	56	76
24	44	66	77
26	45	67	78
31	46	75	78
21			

⇒ Suppose this is a data.

(1) company likes to know by what time 10% of the and 90% of the wire are

3

Ans: $n = 25$ mu, $p = 19$ (ast) = 0.6 $\uparrow$ "Observation"

80, value at approx (2.6) th will be :-

~~Subst~~ = value at 2nd + 0.6th obs

$$= 3\text{hr} + 0.6 (\text{value at 3rd pos} - \text{value at 2nd pos})$$

$$= 3 + 0.6 (5-3)$$

$$= 3 + 1.2 = 4.2\text{hr Ans.}$$

while,

$$r_{90} = \frac{90(25+1)}{100} = \frac{26 \times 9}{10} = \underline{23.4} \text{ } ^{\circ}\text{obs}$$

$$\text{value at 23.4th obs} = 23^{\text{rd}} \text{ obs} + 0.4^{\text{th}} \text{ obs} \\ = 78 + 0.4^{\text{th}} (\text{value at 24th obs} - 23^{\text{rd}} \text{ obs})$$

6. Measures of Variation

longe
ISD
Variance
Standard
Deviation

Range

→ Difference b/w highest and lowest value in the set. It tells how much data is spread.

29. [70, 85, 92, 60, 75] [Score of students]

- Highest = 92
- Lowest = 60 $\Rightarrow$ Range = $92 - 60 = 32$

This means the hole is 32 ft across a range of 32 ft

② Inter-Quartile Distance (IQD) [also k/a IQR (Inter-Quartile Range)]

⇒ It is the measure of distance b/w Q_1 and Q_3 .

Eg. Dataset = [10, 15, 20, 25, 30, 35, 40]
 $Q_2 = 25$ (median)

⇒ Lower half of data = [10, 15, 20] = $Q_1 = 15$

Upper half of data = [30, 35, 40] = $Q_3 = 35$

$$IQD = IQR = 35 - 15 = 20$$

↑ it means middle 50% of data lies in range of 20 units.

★ ⇒ IQD is used to identify outliers in data.

means, those values in dataset which differ significantly and thus forms an outlier in context.

⇒ Formula to identify outliers using IQR :-

• Lower Bound = $Q_1 - 1.5 \times IQR$ (here -)

• Upper Bound = $Q_3 + 1.5 \times IQR$ (here +)

Eg. [1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 90]

$$Q_1 = 3 \Rightarrow IQR = 9 - 3 = 6$$

outlier value = $Q_1 - 1.5 \times IQR$

$$= 3 - 1.5(6) = -6$$

$$= 9 + 1.5(6)$$

$$= 9 + 9.0 = 18$$

Here, in our dataset there is a value which is higher than than '18' (upper bound) outliers = 90, 90 is a outlier.

V. Imp. FORMULA Trick

$$\text{Variance} = \sigma^2 =$$

② Variance and Standard deviation

$$\text{Variance} = \sigma^2 = \sum_{i=1}^n \frac{(x_i - \bar{x})^2}{N}$$

(Here, $\bar{x}$ = Sample mean
 $\bar{x}$ = Population mean)

$$S.D. = \sigma = \sqrt{\sum_{i=1}^n \frac{(x_i - \bar{x})^2}{N}}$$

(In case of Population Variance)

$$\text{Example} \quad \text{Variance} = \frac{S^2}{n} = \frac{\sum_{i=1}^n (x_i - \bar{x})^2}{N-1}$$

(SD b/w $S = \sqrt{\sum_{i=1}^n \frac{(x_i - \bar{x})^2}{N-1}}$)

(If k/a sample variance, we divide by $n-1$ in case we are calculating variance for sample. This is k/a "Bessel's correction". It is denoted by S^2 .)

[If, we divide for calculating sample by "n" only then, it will give less value which is incorrect. ∴ this is also k/a downward bias. To solve from it we divide by $n-1$.]

③ Degree of freedom

Eg. You want to create a sample of 4 numbers with mean = 10

$$\begin{aligned} \text{1st choice (free choice)} &= 8 \\ \text{2nd choice (free choice)} &= 12 \\ \text{3rd choice (free choice)} &= 9 \end{aligned} \quad \left\{ \begin{array}{l} \text{This sum} = 29 \end{array} \right.$$

⇒ Now, for mean to be 10, sum of 4 no. should be 40. Then only $\frac{40}{4} = 10$. So, 4th no. should be 40 - 29 = 11.

⇒ So, we didn't have this freedom to choose this number. So, Degree of freedom = $N-1$

7. Chebyshev's Theorem

⇒ Chebyshev Theorem / Inequality tells that how much of your data lies b/w certain interval defined by mean & standard deviation.

⇒ Probability of finding a random selected value in an interval defined by $\mu \pm k\sigma$ is $1 - \frac{1}{k^2}$.

$$P(\mu - k\sigma \leq x \leq \mu + k\sigma) \geq 1 - \frac{1}{k^2}$$

$\left(\begin{array}{l} 1 = \text{probability} \\ \mu = \text{mean} \\ \sigma = \text{S.D.} \\ k = \text{no. of S.D. from } \mu \end{array} \right)$

Q. Amount spent per month by a segment of credit card users of a bank has a mean value of 12000 and standard deviation of 2000. Calculate proportion of customers who are spending b/w 8000 and 16000.

Ans. Given, mean = $\mu = 12000$
S.D. = 2000

$$\text{Now, } \mu \pm k\sigma \Rightarrow 12000 \pm k(2000) \quad \left\{ \begin{array}{l} = 8000 \rightarrow \text{join here } k=2 \\ = 16000 \rightarrow \text{join here } k=2 \end{array} \right.$$

So, Now,

$$P(\mu - 2\sigma \leq x \leq \mu + 2\sigma) = P(8000 \leq x \leq 16000) = 1 - \frac{1}{2^2} = 0.75$$

⇒ That is the proportion of customers spending b/w 8000 and 16000 is at least 0.75 or 75%.

8. Skewness and Kurtosis

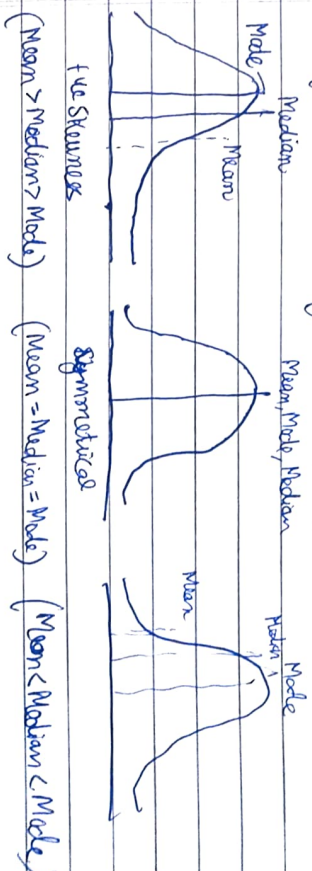
★ Skewness

⇒ Skewness is the measure of asymmetry of probability distribution of a real valued random variable. In other words skewness tells us, the proportion of data b/w μ and $\mu - k\sigma$ is same as μ and $\mu + k\sigma$ ($k = \text{some true constant}$).
It tells us the extent to which distribution leans to the left or right of mean.

$$g_1 = \frac{\sum_{i=1}^n (x_i - \mu)^3}{n \sigma^3} \quad \text{Here, } \begin{array}{l} x_i = \text{each value in dataset} \\ \mu = \text{mean of dataset} \\ n = \text{no. of items} \\ \sigma = \text{S.D.} \end{array}$$

if,

$g_1 = +ve = \text{the skewness is tail longer in right side}$
 $g_1 = -ve = \text{tail, left side}$
 $g_1 = \text{close to } 0 = \text{symmetrical distribution}$

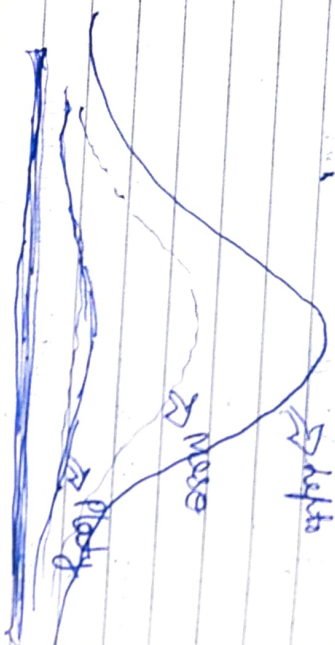


Kurtosis

Kurtosis tells us the peakness where a skewness tells where the tail is $\sqrt{\text{tail}} \neq \text{tail}$.

$$(K) \text{ Kurtosis} = \frac{\sum_{i=1}^n (x_i - \bar{x})^4}{n \cdot \sigma^4}$$

If, $k < 3 \Rightarrow \text{Platykurtic}$
 $k = 3 \Rightarrow \text{Mesokurtic}$
 $k > 3 \Rightarrow \text{Leptokurtic}$



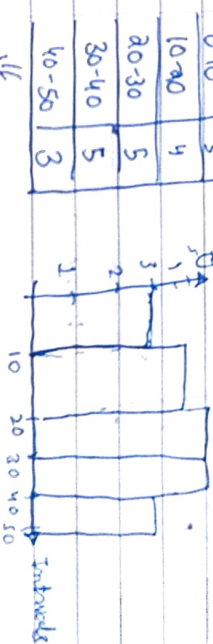
Graphs (Data Visualisation)

① Histogram ("Frequency polygon and frequency curve")

$\Rightarrow$ If you are not given data in frequency form, make in frequency.

Ex. 1, 3, 17, 32, 5, 63, 49, 50, 20, 36, 46, 50, 27, 29, 30, 32.

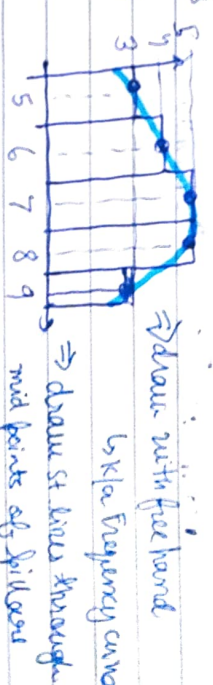
Interval / Bin Size



or

If use the above mid values

0-10	5	3
10-20	6	5
20-30	7	5
30-40	8	3
40-50	9	3



$\Rightarrow$ draw with free hand
 $\hookrightarrow$ Kipla "Frequency curve"
 $\Rightarrow$ draw st lines through mid points of pillars
 $\hookrightarrow$ Kipla "Frequency poly"
 [Make sure to extend freq. polygon to make it closed since polygons are closed]

Ogive Curve

$\Rightarrow$ curve made / represent cumulative frequency (C.F)
 $\Rightarrow$ It can be less than or more than type
 $\Rightarrow$ When they both meet, the intersection point is median. Hence we can calculate median from each one separately also.

3

0-10	3
10-20	27
20-30	15
30-40	3
40-50	5
	53

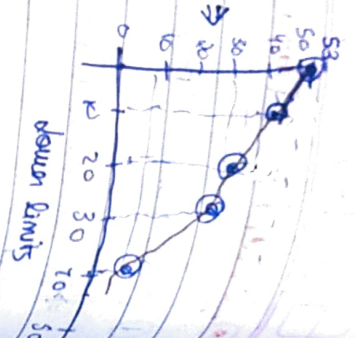
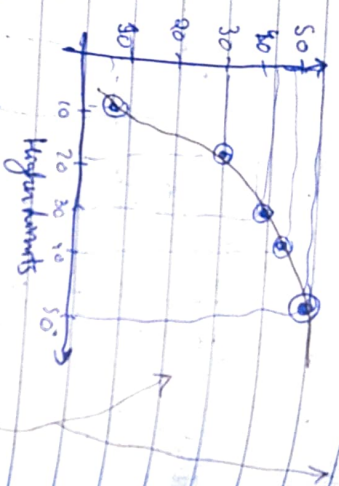
more than

more than 0	53
more than 10	48
" " 20	45
" " 30	30
" " 40	3

less than

CF

less than 10	3
" " 20	30
" " 30	45
" " 40	48
" " 50	53



⇒ New, $N = 53$

Median = $\frac{(53+1)}{2} = \frac{54}{2} = 27^{\text{th}}$ observation

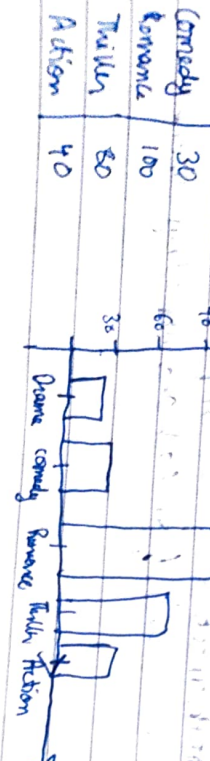
↓
37th value and
corresponding value of X will
be the value of 27th observation.

eg. category female



2. Bar graphs

used in case of Nominal data where names are associated with Numerical value.



3. Pie chart

eg. Subject frequency

Scale	6	72°
P.E.	4	48°
Maths	9	108°
English	7	84°
History	4	48°

Total = 30
Now, $360 = 12^\circ$

① draw a circle
and mark the angles using the center
bottom of P.



4. 5 point Summary & Box-Plot (or Box & whisker Plot)

concise depiction of dataset which gives 5 things:-

- ① Minimum Value
- ② Q_1
- ③ Q_2 or median
- ④ Q_3
- ⑤ Maximum Value

⇒ graphical representation of 5 point summary is Box Plot.

eg. Data = [55, 66, 71, 75, 80, 85, 90, 95, 99]

Min = 55
 $Q_1 = \frac{25(55+75)}{100} = \frac{25(130)}{100} = 32.5$
 $Q_2 = \frac{75(55+99)}{100} = \frac{75(154)}{100} = 115.5$
 $Q_3 = \frac{75(75+99)}{100} = \frac{75(174)}{100} = 130.5$
 Max = 99



Box Range from Q_1 to Q_3 shows from max & min and median is marked inside box. If in case of extreme, they are displayed as point beyond