

Unit-Wise Concise Theory BA

NIT-1

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"

• 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)

Examples of Big data:
Social media data
E-commerce data
Banking, financial, healthcare data
Entertainment data
Research data, Data from sensors, etc.

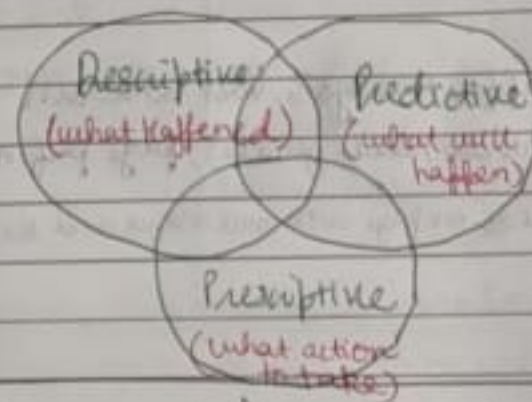
④ Structured data:- Has clear structure and follows a regular sequence. Organized formatted in a way which is easily readable by both humans & machines. It generally follows tabular structure.

Relational Database entries:- Employee data in company database with columns Name, ID, Department, Salary.
⑤ 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)

⑥ 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 lose 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 action to optimize or improve outcomes. gives advice/action to achieve desired result. [Suggests action to achieve best outcome] Eg. Predictive analytics told us sales would go up during this period & so prescriptive analytics might recommend using inventory of those products, running targeted ads or discounts.

independent variable.

S. Jini 23/03/21

VIF = measures how much the variance of a regression coeff. is inflated due to multicollinearity in model. $VIF=1$ indicates no correlation among 'size' and other variable like bedroom, age etc. hence, no multicollinearity. As a rule of thumb VIF greater than 5-10 is a problematic amount of multicollinearity.

Model Evaluation

- $mse = \text{mean_squared_errors}(y_pred, y_test)$
- $r^2_score(y_train, model.predict(x_train))$
- $r^2_score(y_test, y_pred)$
- computes MSE b/w model predictions on test data and the actual value of the test data. [Actual value of y , and the value of y which model predicted.] If MSE is less meaning our model is a better fit.
- y_train :- actual value of y used in training
- $model.predict(x_train)$:- predict values of y based on the training data of x . (NOTE: x is x_train & y is y_train)
- This line basically calculate R^2 for "training set". (x_train used)
- Similar to previous line it calculates R^2 but for "testing set". It compares the actual values (y_test) to the predicted values (y_pred). It measures how well unseen data (test data) is predicted by the model.

Checking Multicollinearity with VIF

- ~~Suppose~~ x is a dataframe, with the columns :- Size, Bedroom and age. These features are used to predict y = price of house.
 $x = df.drop(\text{column} = 'House price')$ is actually this only.

⇒ Now,

import lib to
calculate VIF

Import Statsmodel. api as sm

from Statsmodel.stats.outliers_influence import variance_inflation_factor

Splitting the data

$x_{train}, x_{test}, y_{train}, y_{test} = \text{sk.train_test_split}$

$(x, y, \text{test_size} = 0.2, \text{random_state} = 42)$

output of the fnx

It returns four values:-

x_{train} :- 80% of x chosen randomly to train

x_{test} :- 20% of remaining x is stored here.

(Similarly with y)

y_{train} :- corresponding to randomly chosen 80% x (x_{train}) data the rows of y . [it is also 80% hence]

y_{test} :- corresponding to randomly chosen 20% of x (x_{test}) data the rows of y . [it is also 20% hence].

fnx

$\Rightarrow$ train_test_split a fnx of sklearn lib. used to split arrays/matrices into random train & test subsets.

$\Rightarrow x$ = dataset with house price, y = house price data or columns

$\Rightarrow$ test_size = 0.2 $\Rightarrow$ specifies that for testing 20% of data will be used (of both x & y) and for training 80% of x & y data used.

$\Rightarrow$ random_state = 42 ensures data is selected i.e. 80% & 20% is selected randomly not like first 80% then 20% it should be random.

Creating and Training the model

$\text{model} = \text{LinearRegression}()$

$\text{model.fit}(x_{train}, y_{train})$

$\Rightarrow$ A ~~new~~ linear regression model by name "model" is created and trained (i.e. in it fit data) on data x_{train} & y_{train} .

Making prediction

$y_{pred} = \text{model.predict}(x_{test})$

$\Rightarrow$ The model is used to predict the house price (y) based on the training set (x_{test}). The predictions are stored in y_{pred} .

• Creating a dataframe for VIF values

```
vif = pd.DataFrame()
vif['Features'] = X.columns
```

→ a new dataframe is created by the name vif

→ In that dataframe we have a column named Features which is same / taken from columns in X.

• Calculating VIF for Each Feature

```
vif['VIF'] = [variance_inflation_factor(X.values, i) for i in range(X.shape[1])]
```

⇒ This ~~code~~ line calculate VIF for every feature one by one. at the end, after running this code we get :-

Features	VIF
Size	4.5
Bedrooms	7.2
Bathrooms	3.8
Age	1.2

⇒ VIF = 1 means no correlation b/w size and bedrooms

⇒ VIF greater than 5-10 indicates there is a problematic amount of multicollinearity.

Multicollinearity

↳ refers to a situation in statistical model (usually in multiple regression) where two or more independent variable (X) are highly correlated with each other.

⇒ When multicollinearity is present in ~~model~~, it can cause problem in regression analysis.

Eg: y = price of house, X = size, no. of bedrooms etc. Now, larger houses (size ↑) tend to have more bedrooms so, these two variables are highly correlated. This is k/a multicollinearity. If we include both in regression model, it's difficult to determine their individual effect on house prices also, model may favour one or unfavour another as they are highly correlated. ⇒ This is a problem.

⇒ In order to detect this ^{and its power} we use VIF (Variance inflation factor) for each

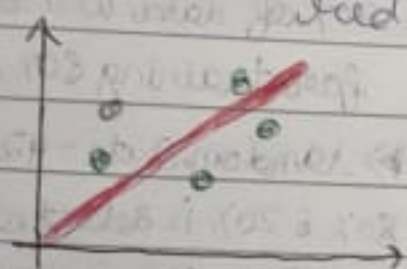
⇒ $mse = \text{mean squared error (y, y-predicted)}$.

The line which the model draws actually calculate mse , b/w the actual y values and predicted one. It is a measure of how close a regression line is to a set of points.

The lower the MSE , the better the model performance.

Plotting the result

* for st. line ⇒ `plt.scatter(x, y, colour='green')`
`plt.plot(x, y-predicted, colour='red')`
`plt.grid`
`plt.show`



Multiple

(we have multiple independent variables i.e. 'x').

⇒ Here there are multiple x (independent variable) predicting a single y (= dependent variable). Eg. Predicting house price ($y = \text{house price}$) based on square area, no. of bedrooms, age of house etc. (multiple x).

⇒ consider, `df = pd.read_csv('House.csv')`

`df`

Suppose we get :-

	House price	no. of bedroom	age of house	sq. area
0	10,000	2	2	100
1	20,000	4	4	121
2	30,000	6	5	111

Importing the data

Import sklearn model selection as `sk`

`x = df.drop(column='House price')`

`y = df['House price']`

→ this means x contains all columns except house price

Calculate model parameters (intercept & coefficients)

`model.intercept` —

`model.coef[0]` or `model.coef`

(indexing operation)

used to access the first element of the array. But, if we use `array[value]` if we use `[0]` forward to it it will directly give `[value]`.

this it can give result

Predicting and Evaluating the data

`y_predicted = model.predict(x)`

`r_square = r2_score(y, y_predicted)`

`mse = mean_squared_error(y, y_predicted)`

⇒ `model.predict(x)` ⇒ uses the model to predict 'y' values for given 'x' values. This line basically say ⇒ Based on the x you learned so far, tell me what 'y' would be for each x . The result is stored in `y_predicted`.

⇒ `r2_score` ⇒ calculates coeff. of determination, k/a R^2 . It gives you an idea of how well your independent variable (x , which you can change) explain the variance in your dependent variable (one which you are trying to predict).

$R^2 = 1$ → model perfectly fit. You could correctly predict y based on x .

$R^2 = 0.8$ → means 80% of the variance in y can be explained by x .

as R^2 ↓ means your model is less & less accurate and at

$R^2 = 0$ means model does not explain / predict variance in y .

requirement of such models).

-1 in reshape = tells Numpy to automatically calculate the no. of rows based on the length of array & no. of column

1 → specifies that there should be one column (since, you have one feature or one independent variable)

$X = \text{np.array(adult).reshape(-1,1)}$
 $y = \text{np.array(sale)}$

Creating and training the linear regression model

$\text{model} = \text{LinearRegression}()$
 $\text{model.fit}(x, y)$

⇒ linearRegression is the command that creates a linear regression model (which will predict value of a dependent variable^(y) based on an independent variable^(x)). It is stored in a variable named "model".

⇒ model.fit(x, y) It is used to fit or train the model based on the data we provided.

⇒ x, y are the data we provided to the model. where, x (independent variable) ^{used} to predict y (dependent variable).

⇒ we are trying to predict " y " if we gave the model only " x " is the purpose of this model.

⇒ when we run, $\text{Model.fit}(x, y)$ we tell the model to look at the data ' x ' & ' y ' and learn the rel^n b/w them and try to figure out the changes in ' x ' associated with changes in y .

⇒ After this the model finds the best fit line to represent rel^n b/w x & y .

All codes Unit-2

SIMPLE

Independent variable $\rightarrow$ I can change this variable, I can control over it.
Dependent variable $\rightarrow$ It is the output/response to the independent variable.
Eg. Sale = dependent variable, $[y]$ adut = independent var. $[x]$ (I can change its amount)

Import numpy as np

Import pandas as pd

$\rightarrow$ Import matplotlib.pyplot as plt

★ $\rightarrow$ From sklearn.linear_model import LinearRegression

★ $\rightarrow$ From sklearn.metrics import r2_score, mean_squared_error

Data for regression

adut = [10, 12, 10, 15, 16, 18, 19, 20]

Sale = [24, 35, 30, 28, 25, 30, 32, 35]

$\Rightarrow$ Out of this one is independent (I can change = adut) and one is dependent (= Sale).

Making the data compatible for linear regression.

$\rightarrow$ We convert the data into array like the models we use expect the data to be in array. As, it enable faster computation. Also it expect input (i.e. x) to be 2D array so that's why reshape $(-1, 1)$ for n variable. It transforms 1D array to 2D array (which is the

• To calculate stuff, on a particular column/row

→ `stats.zscore(df['column name'])`

→ `df['column name'].mean()`

`df[' '].median()`

• `var()`

• `std()`

• `quantile(0.25)` or `.quantile([0.25, 0.75])`

• `skew/kurtosis()`

[in case of row use
`df.loc[0]`
↑ index of row
(minus row number)]

• For graph of a particular column or row

`plt.hist(df[' '])`

(you can include density, bins etc. too to the code.)



★ Shapiro Wilk Test for Normality

import scipy.stats as stats

exam_scores = [data, ...]

→ `shapiro_stat, shapiro_p_value = stats.shapiro`

`print("shapiro-wilk test - statistic:", shapiro_stat * (exam_scores))`

`print("shapiro-wilk test - p-value:", shapiro_p_value)`

formule for both shapiro-wilk test statistic and p-value is same.

We get the output as

shapiro-wilk test statistic: -
" " " " p-value: +

Output.

◎ Find correlation

→ `df.corr()`

→ It will give output like

var1	1	0.6	0.1
var2	0.9	1	-0.2
var3	-0.2	0.14	1

Correlation measures the strength and direction of association b/w two continuous variables.

→ value of correlation (-1 to 1)

→ 1 = perfect +ve correlation
0 = no correlation

-1 = perfect -ve correlation

You should have a df meaning
since `df = pd.DataFrame(data)`
or `df = pd.read_excel(" ")`

apne south son (var1-1, 2-2, 3-3 den correlation 1 he nahega)

- No. of rows & columns

df.shape

gives (Total no. of rows, Total no. of columns)

- Info on data types in the df

↳ gives the type of data in every column

df.dtypes

- All columns names

↳ gives the name of every column.

df.columns

- Dropping rows with missing value

df2 = df.dropna()

drop/removes any row where there is a missing value.

df2

↳ on printing now we will get only that row which has all data points.

Eg.

	Var 1	Var 2	Var 3
0	10	20	60.0
1	20	30	40.0
2	10.1	nan	69.1
3	100	20	70.0
4	400	10.1	nan

after dropna() and printing

new table

2 rows missing as they had missing values.

	Var 1	Var 2	Var 3
0	10	20	60.0
1	20	30	40.0
3	100	20	70.0

- To access a particular row or column.

★ iloc → In pandas, stand for integer location. It gives us Select rows and columns values.

df.iloc[0] = gives 1st row of df

df.iloc[:,0] = gives 1st column of df

df.iloc[0,0] = gives 1st element of 1st row & column

df.iloc[:,3:4] = gives 4th and 5th column

(Indexing in Python starts from 0th & 1st means → 1st & 2nd etc. Rows & Columns)

[Index 3 & 4 means 4th posⁿ and 5th posⁿ]

→ NaN or nan = not a number

- To select a particular column or Row

Selected _ column = df . column _ name
OR

Selected _ column = df . ["column 1", "column 2", "column 3"]

(for rows we use integer location iloc, alternatively iloc can also be used for columns.)

- To get a summary

Summary _ stats = df . describe ()

↳ It will give out 8 things :-

- count :- It means the total no. of entries without nan. in a column
- mean
- std

And 5 point summary :- min, 25%, 50%, 75% and max

- Replacing missing values with True

↳ Tells us if any missing value

df . isnull ()

⇒

	Var 1	Var 2	Var 3
0	False	False	False
1	False	True	False
2	False	True	True

• True means there is a nan or missing value here in this point.

- Sum of the missing values

↳ Tell the total missing value in each column

df . isnull () . sum ()

like here,

Var 1	0
Var 2	1
Var 3	2

output

Var 3 has two datapoints at which there is no value.

★ Dataframes (df)

⇒ In the pandas library, we have a powerful tool for data analysis which provides data in structured form (like the one arranged in excel) k/a 'Dataframe'.

⇒ We always import excel in form of dataframes.

All commands of Dataframe ★

① Import Excel

import pandas as pd
df = pd.read_excel("path")
df

② Display rows

df.head() ⇒ gives first five rows of dataframe.

df.head(10) ⇒ gives 10 rows

df.tail(10) ⇒ gives the last 10 rows.

③ Get info

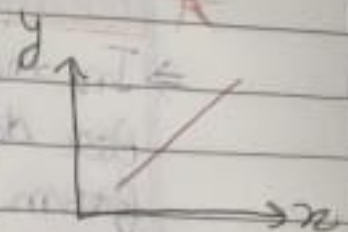
df.info ⇒ gives the info on non-null entries, memory usage, data type (like
int = integers, float = real no. (2.34, 5.05
whole no. (3, 4 etc.) etc.)

string = text/characters, boolean = True or False
(str) (bool) data

etc.

For Simple line graph

`plt.plot(x, y, colour='red')`

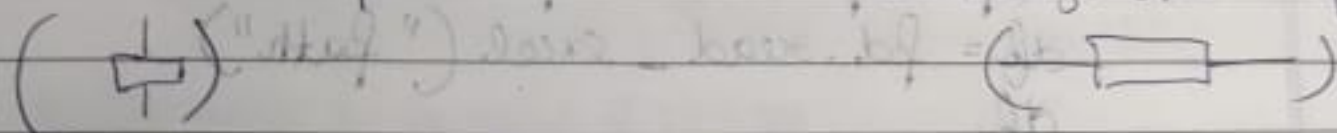


NOTE :- adding this code just after histogram code will give a frequency curve.

Boxplot

→ `plt.boxplot(data1)` OR `plt.boxplot(data1, vert=False)`

For vertical boxplot For Horizontal boxplot



→ `plt.title('title')`
→ `plt.show`

★ In case we are drawing scatter plot of two different columns (independent and dependent)

then, `plt.scatter(df['columnname1'], df['columnname2'])`
`plt.title('title')`
`plt.xlabel('column 1')`
`plt.ylabel('column 2')`
`plt.show`

$q_1 = np.percentile(data, 25)$

$q_2 = np.percentile(data, 50)$

$q_3 = np.percentile(data, 75)$

• Graphs.

For Histogram

Import matplotlib.pyplot as plt
`plt.hist(data1)` OR `plt.hist(data1, bins=10)`

Remains same
`plt.xlabel('x')`
`plt.ylabel('y')`
`plt.title('title')`
`plt.show`

OR
`plt.hist(data1, bins=10, density=True, alpha=0.5)`



For transparency of graph

For frequency curve

Frequency curve is drawn above histogram so, add the following line at the end of histogram code.
⇒ Histogram code → then

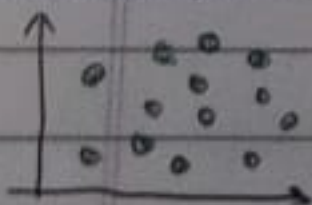
`plt.plot(x, y, 'r-', linewidth=2)`

r- means red colour line of width = 2

or you can simply type `colour='red'`



For Scatter plot

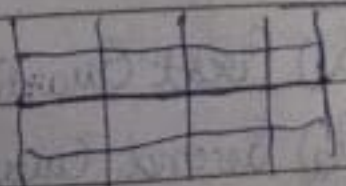


`plt.scatter(x, y, colour='green')`

green dots of (x,y) points will be shown



`plt.grid` → can be used to make grids on graphs. Just add this line at end or anywhere.



- For Sorting data $\rightarrow$ Using $\Rightarrow$ Using order

Ascending / Using $\Rightarrow$ `sorted(data)`

Descending / Using $\Rightarrow$ `sorted(data, reverse=True)`

Just directly these commands no 's' sign.

- For Range

$\text{Range} = \text{max}(\text{data}) - \text{min}(\text{data})$

[compute value b/w min and max value]

- For Standard deviation, Variance, Skewness, Kurtosis, Z-score

★ `std = np.std(data)`

★ `var = np.var(data)`

} Both, std and var are numpy lib

★ from scipy.stats import skew
`skew_value = skew(data)`

★ from scipy.stats import kurtosis
`kurtosis_value = kurtosis(data)`

} we have to import these two from scipy.stats hence no np. $\Rightarrow$ in these formula.

★ Import scipy.stats as st
`Z_score = st.zscore(data)`

$\Rightarrow$ we import scipy lib so `(st.zscore(data))`, tells they are i.e. zscore is from scipy lib/module.

- For Percentile

$\Rightarrow$ we use 'numpy' lib to calculate percentiles.

NOTE

(Q₁) First Quartile corresponds to 25th percentile

(Q₂) Second Quartile " " 50th percentile aka Median

(Q₃) Third Quartile " " 75th percentile

★ List of Import Statement

- ① `from statistics import mode`
- ② `from scipy.stats import skew`
- ③ `from scipy.stats import kurtosis`
- ④ `from sklearn.metrics import r2_score, mean_squared_error`
- ⑤ `from statsmodels.stats.outliers_influence import variance_inflation_factor`

★ Basic fnx codes

• For Mean, Median, Mode → define, `data = [9, 10, 12, 13, 8, 7, 5, 9, 9, 10]`

① `mean_value = np.mean(data)`

To print just write → `mean_value` or `print(mean_value)` (`np.mean` is the main formula. The variable in which it is stored i.e. `mean_value` name can be changed to any.

② `median_value = np.median(data)`

③ `from statistics import mode`

then, `mode_value = mode(data)` (* In formula of mode there is no `np.`) [As, it is not called from numpy library]