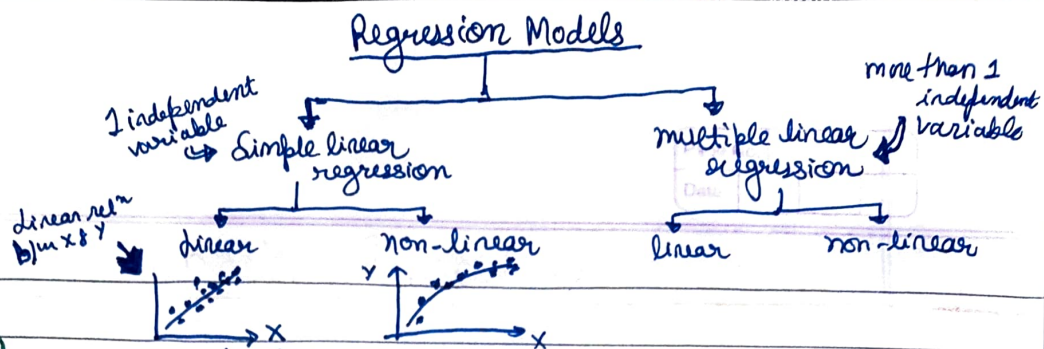


Unit-2

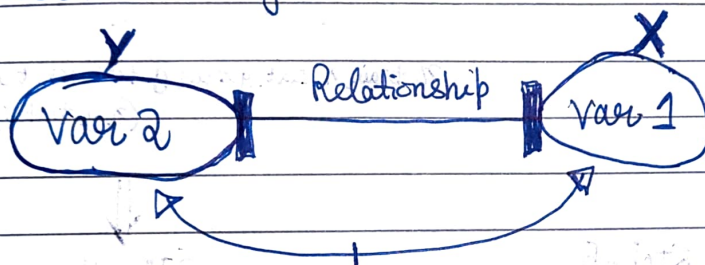


Simple Linear Regression (SLR)

→ It is a statistical technique to find existence of an association relationship b/w a 'dependent variable' and an 'independent variable'.

⇒ Here, there is only one independent variable in the model.

★ Regression can only tell



does not mean that value of independent variable (Var 2 or Y) depends on value of independent variable (Var 1 or X)

• Defⁿ ⇒ It is a statistical model in which there is only one independent variable and the functional relⁿ b/w the dependent variable and regression coefficient is linear.

$$Y_i = \beta_0 + \beta_1 X_i + E_i$$

Here,

Y_i = dependent variable in sample / Response variable / outcome variable

X_i = independent variable in sample / predictor variable / input or explanatory variable

β_0 and β_1 = Regression parameter or Regression coefficient

β_0 = intercept of regression line (y-intercept)

E_i = random error (or residuals)

β_1 = slope of the regression line

Simple Linear Regression Model Building

Step 1 Collect/Extract

data

(on the dependent and independent variable. Time & resource consuming & expensive even with organisation having well-designed resource planning system) enterprise (ERP)

Eg. $adult = [\dots]$
 $sale = [\dots]$
OR

$df = pd.read_excel('---')$
 $df.head()$

Step 2

Pre-process data

- Ensure Quality of data
- Check for issues such as scalability, completeness, usefulness, accuracy etc.
- All the steps required to give clean data

Step 3

Dividing data into training and Validation datasets

• Data divided into two sets - training data & Test (validation) data

• % of training data is usually 70-90% rest is testing data

• Training data used for developing model
Testing data to check & select the model

$x_{train}, x_{test}, y_{train}, y_{test} = ck.train_test_split(x, y, test_size=0.2, random_state=0)$

Step-5

Build the Model

• Model is build using training data.

• OLS (Ordinary least square) method

need to estimate parameters in linear regression model.

• The goal of OLS is to find the best fit line through a set of data points by minimizing the sum of squares of vertical distances (residual) b/w the observed and value predicted by the model.

[Use SSR (sum of squared residuals) = $\sum_{i=1}^n (y_i - \hat{y}_i)^2$ where $\hat{y}_i$ are predicted values]

• Using this we find values of β_0 and β_1 and once these regression parameters are set $\Rightarrow$ for a given x predicted value of y will be $\hat{y} = \beta_0 + \beta_1 x + \epsilon$. i.e. our model is made ✓

Eg. $model = \text{LinearRegression}()$
 $model.fit(x, y)$
 $model.intercept_$
 $model.coef_[0]$

Step-4

perform descriptive analysis

of the data like box plot, scatter plot (to identify any outliers)

• to reveal real b/w the two variables
• useful to determine functional relationship between dependent and independent variable.

Eg. $df.boxplot('column name')$
 $df.hist(df['I'])$
 $df['column name'].median$ etc.

OR
 $mean_value = df.mean(adult)$
 $plt.boxplot(adult)$

Step-6 $\Rightarrow$ perform model Diagnostic

- Important to run diagnostic test to see if the model runs correctly
- Important to ensure model runs consistently and somewhat accurately on the ~~training~~ testing data as it did on training data.

Eg. $y_predicted = model.predict(x)$

Step-7 :- Validate the model and Measure model accuracy / plot the model

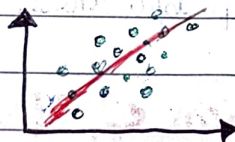
⇒ Ensure model runs accurately on the testing data and perform various test to see if the models fit perfectly.

$$r^2\text{-square} = r^2\text{-score} (y, y\text{-predicted})$$

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

Plot the model :-

```
plt.scatter(x, y, colour='green')
plt.plot(x, y_predicted, colour='red')
plt.grid
plt.show
```



Step-8 :- Decide on model Deployment

⇒ Final step involves figuring out how to make strategic decisions like adjusting budgets, setting prices etc. based on the model output. It basically help in decision making.

Regression Eqⁿ and its Analysis

$$Y_i = \beta_0 + \beta_1 X_i + \epsilon_i$$

y-intercept slope

So, if :- $\hat{Y}_i = 61655.3553 + 3076.1774 X_i$ (MBA stud vs % of marks in salary this class 10th)

Interpretation :-

↓
If a student have 0% in class 10th still his/her salary = ₹ 61,555 approx

↘ It means for every 1% increase in class 10th score the salary increases by ₹ 3076.1774...

(X=0, then, y=? is y-intercept)

★ Estimation of Parameters (using OLS)

⇒ A crucial step in regression model building is estimation of regression parameters.

⇒ Given an dependent variable (Y_i) and the corresponding independent variable values (X_i) and each subject to random errors (ϵ_i), we have to find the best eqⁿ to represent the relationship b/w these two variables.

● Assumption of OLS-based Linear Regression model

① Linearity :- Regression model must be linear in its parameters meaning, β_0 and β_1 must be linear. [variables X_i etc. can or cannot be linear]

② Variation in independent variable (X_i) :- There should be sufficient variation in X_i values then only model can be correct. Eg. Y = consumption level and X = income level. And the X is income taken of every person is ₹10,000 ⇒ meaning no variation ⇒ So, model will not able to predict a new $\hat{Y}$ based on a new X .

③ The explanatory variable X is assumed to be non-stochastic (ie. X is deterministic). [Fixed value of independent variable]

⇒ This means, X is controlled and not subjected to random variation or influences ⇒ which will help predicting Y more straightforwardly.

Eg. X = Study hour

Y = marks in exam

} If, we take into account the variation of student's feeling like sometime they wanna study so before 1 day they study 18 hrs but before they did not due to their mood ⇒ So it makes determining Y difficult. So we make ' X ' more controlled and consistent so its easy to determine Y .

④ Residuals (ϵ_i) follow a normal distribution and have zero mean

⇒ This means, if we plot all residuals from our model, they would form the classic bell shaped curve of normal distribution.

⇒ $\text{mean} = 0$, means on average the errors don't systematically overestimate or underestimate the dependent variable (Y).

★ Now, Based on above 4 assumption :- In OLS, our objective is to find optimal value of β_0 and β_1 such that we minimise the Sum of Squared Errors (SSE) i.e.

$$SSE = \sum_{i=1}^n \epsilon_i^2 = \sum_{i=1}^n (Y_i - \beta_0 - \beta_1 X_i)^2$$

on partial derivating it w.r.t. β_0 and β_1 , we get :-

$$\beta_0 = Y_i - \beta_1 X_i, \quad \beta_1 = \frac{\text{cov}(X, Y)}{\text{var}(X)} \quad \text{OR} \quad r_{xy} \frac{\sigma_x}{\sigma_y}$$

calculate β_1 and put value here. (σ_x, σ_y are SD).

★ Imp

⑤ The variance (meaning how spread out the residuals are) of the residuals $\Rightarrow \text{Var}(\epsilon_i | X_i)$ is constant for all values of X_i .

⇒ When the variance of residual is constant for different values of X_i its k/a $\Rightarrow$ Homoscedasticity.

⇒ But, when the variance of residual is non-constant its k/a $\Rightarrow$ Heteroscedasticity

(We want Homoscedasticity ✓ and we don't want heteroscedasticity ✗)

Eg. Whether its a cold day (less ice cream sold) or a hot day (more ice cream sold), the amount to which our prediction is wrong should have same spread. If the variance changed with temp $\Rightarrow$ then Heteroscedasticity which will violate our OLS assumption.

[Here, X_i = temp, Y_i = ice cream sold]

⑥ In case of time series data, residuals are uncorrelated that i. $\text{cov}(\epsilon_i, \epsilon_j) = 0$ for all $i \neq j$.

⇒ In time series data ⇒ data is collected over time ⇒ residual are uncorrelated means ⇒ residual/error at one point of time (e.g. error in predicting electricity consumption on Monday) does not tell anything about the residual (error in " " on Tuesday). They are independent of each other.

• Validation of Simple Linear Regression Model •

⇒ It's important to validate the regression model to ensure its validity and goodness of fit before it can be used for practical applications.

⇒ The following measures are used :- (+ Their interpretation on calculating them)

* Coeff. of Determination (R^2) [R-square]

↳ Objective of regression is to explain the variation in Y (or predict) based on knowledge of X .

⇒ R^2 ⇒ measures the % tge of variation in Y explained by the model ($\beta_0 + \beta_1 X_i$)

Simple Linear Regression model

↓
Broken into -

$$\begin{array}{lcl} Y & = & \beta_0 + \beta_1 X_i + \epsilon_i \\ \text{(variation in } Y) & & \text{(variation in } Y \text{ explained by the model)} \quad + \quad \text{(variation in } Y \text{ not explained by the model)} \end{array}$$

So, Total variation = Variation in Y + Variation in Y
in Y explained by model not explained by model.

$$Y_i - \bar{Y} = \hat{Y}_i - \bar{Y} + Y_i - \hat{Y}_i$$

↓ Applying summation to this squares :-

$$\sum_{i=1}^n (Y_i - \bar{Y})^2 = \sum_{i=1}^n (\hat{Y}_i - \bar{Y})^2 + \sum_{i=1}^n (Y_i - \hat{Y}_i)^2$$

SST

SSR

SSE

(Sum of square of total variation)

(Sum of square of variation explained by regression model)

(Sum of square of variation)

(Sum of square of error or unexplained variation)

Coeff. of Determination = $R^2 = \frac{\text{Explained variation}}{\text{Total variation}}$

$$= \frac{SSR}{SST} \quad \text{OR} \quad 1 - \frac{SSE}{SST}$$

(Since, $SSR = SST - SSE$)

★ Interpretation

↓
 ◎ R^2 lies b/w 0 and 1

◎ Higher the value of R^2 (meaning closer to 1) the better fit the model.

◎ R^2 less than 0.5 (or 50%) $\Rightarrow$ not good fit
 (Independent var.)

Eg: $R^2 = 0.1563 \Rightarrow$ means, grade 10 marks explain 15.63% of variation in the starting salary of MBA students

($X =$ grade 10 marks, $Y =$ salary) $\Downarrow$ (dependent var.)

(not good fit) $\rightarrow$ it should had ideally been able to explain 100% variability.

Analysis of (ANOVA) (F-static)

same constant variance

★ Significance Test (p-value)

→ In SLR we are trying to find a relⁿ b/w X and Y.

⇒ p-value is used to determine whether X has statistically significant relationship with Y (dependent var).

How it works in SLR :-

- ① H_0 (null hypothesis) :- In SLR, null hypothesis states that there is "no" relⁿ b/w X and Y. $\Rightarrow$ means, $\beta_1 = 0$ (β_1 = coeff. of X).
- ② H_A (alternative hypothesis) $\Rightarrow$ There is relⁿ b/w two $\Rightarrow \beta_1 \neq 0$.
- ③ Significance level (α) $\Rightarrow$ we choose a significance level ($\alpha = 0.05$ mostly) before running regression analysis.
- ④ p-value $\Rightarrow$ Fit SLR model into data, we get a p-value for β_1 (coeff. of X).
- ⑤ Comparison $\Rightarrow$

If, p-value $< \alpha \Rightarrow$ Reject H_0 and hence there is relⁿ b/w two.
Accept H_A

If, p-value $> \alpha \Rightarrow$ Accept H_0 , there is no relⁿ b/w two.
Reject H_A

Eg. data set of House prices (Y) [trying to predict] based on size of house (X).

$\Rightarrow$ perform SLR and find p-value of the coeff β_1 (coeff. of X_1), suppose it comes to be $= 0.03$.

$\Rightarrow$ our significance level $= 0.05$

Since, p-value < 0.05 (α) $\Rightarrow$ we reject H_0 and say there is relⁿ b/w house price and size of house.

p-value for either test $\begin{cases} \rightarrow t\text{-test} \\ \rightarrow f\text{-test} \end{cases}$ } same for In simple linear regression.

so in other words,

p-value for t-test or f-test

$< 0.05 \Rightarrow$ means ~~me~~ μ of X and Y and model is statistically significant.

★ Interpretation

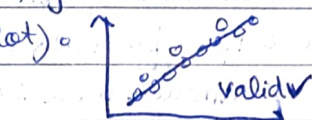
If, p-value less than 0.05 (α = level of significance) we reject H_0 meaning there is relationship b/w X (independent var) and Y (dependent var).

★ Residual Analysis

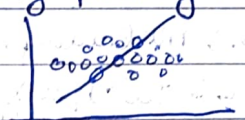
It is important to check if the assumptions of the SLR is being satisfied or not. It consist of checking the following four points:-

① The residual $(y_i - \hat{y}_i)$ is normally distributed or not.

↳ Easiest way to determine this is by plotting P-P plot (Probability-Probability plot).



P-P plot of a regression model
of $E_y = y_i = \beta_0 + \beta_1 x_i$

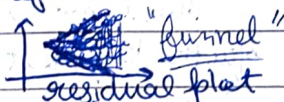


P-P plot of regression model of $E_y = y_i = \beta_0 + \beta_1 x_i$

dots close to diagonal $\checkmark \Rightarrow$ so residual follow normal distribution. And model is valid.

② Test of Homoscedasticity

If there is heteroscedasticity (meaning variance of residual, non-constant) there is funnel shape in residual plot.

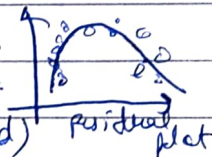


③ Testing functional form of Regression model.

$\Rightarrow$ Residual plot shouldn't have any pattern (or parabola).

If it is there then $\rightarrow$ model is not valid.

(Incorrect functional form used)



④ Outlier Analysis

↳ Outlier are those values whose value show a large deviation from mean value. Presence of outlier can have influence on regression coeff $\rightarrow$ changing the best fit line



★ The following distance measures are used to identify influential observation / our outliers:-

① Z-score $\Rightarrow$ It is the distance of observation from its mean value.

$$Z = \frac{Y_i - \bar{Y}}{\sigma_Y} \quad @ \sigma_Y \text{ and } \bar{Y} \text{ are SD and mean of dependent variable.}$$

★ Interpretation ★

$\Rightarrow$ Any observation with Z-score greater than 3 will be classified as outliers/influential obs. whose presence can alter regression coefficients.

② Mahalanobis distance $\Rightarrow$ Its value greater than Chi-square critical value for an obs. mean that that obs is outlier.

③ Cooks distance :- Cooks distance value more than 1 indicates highly influential obs. / outlier.

Extra point in R^2

$\Rightarrow$ Not all times, high R^2 value mean better fit. It can sometimes be misleading. $\Rightarrow$ This is k/a spurious Regression :- meaning it shows there is relⁿ b/w X and Y even though there is not.

that the model is

Eg. No. of Facebook users | No. of people left Finland.

2004	1 lakh	1.2 lakh
2009	2 lakh	2.9 lakh
		etc.

$\Rightarrow$ on calculating $R^2 \approx 0.8$ it means they both are related / model is correct. Our model eqⁿ is $Y = 1.996 + 2.1X$

$\swarrow$ means for every 1 person leaving Finland 2.1 users ↑

There is no relation b/w the two, but R^2 is

high $\Rightarrow$ hence mislead us so, R^2 is not always a correct measure to tell better fit.

Multiple Linear Regression

⇒ It is statistical technique used to find existence of relationship b/w a dependent variable (Y) and several independent variable.

⇒ Functional form of MLR is given by:-

$$Y_i = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \beta_4 X_4 + \beta_5 X_5 + \dots + \beta_n X_n + \epsilon$$

here, $\beta_1, \beta_2, \beta_3, \beta_4, \beta_5$ etc. are k/a 'Partial regression coeff'.

β_0 = y-intercept or constant term

ϵ = error term

★ Assumptions of MLR

↳ same as of that of SLR actually!!

MLR eqⁿ and its Analysis :-

⇒ consider, two independent variable: Promotion and starting point
dependent variable: - advt. revenue

⇒ consider, dependent variable ⇒ House Price,
independent variable ⇒ size and Age of house

lets say,

$$Y = 41258 + 5931.85 \times \text{Size} + -1000 \times \text{Age of house}$$

↓

○ for every one unit ↑ in size ⇒ price of house
↑ by ₹ 5931.85 (keeping age const.)

○ for every one unit ↑ in age of House ⇒ price of house
↓ by ₹ 1000 (keeping size const.)

⇒ In MLR, it may not be possible to control a variable (i.e. keeping it constant) in many situations.

→ Those which represent characteristic like person gender, marital status, exam status. It can take numerical value. Eg. 1 for male, 2 for female etc.

Working with categorical variables

↳ Eg. $Y = \text{Salary}$, $X_1 = \text{High school}$, $X_2 = \text{Bachelor}$, $X_3 = \text{Master}$, $X_4 = \text{Phd}$.

We take X_1 as reference and calculate others in reference to it.

Multi-collinearity and VIF

→ It refers to presence of high correlation b/w two or more independent variables in a multiple linear regression model.

→ It can cause issue with the estimation of coeff and effect the model.

→ It refers to ~~per~~ presence of perfect linear relationship b/w two or all independent variables in a regression model.

Eg. There are n independent variables then

$$\alpha_1 x_1 + \alpha_2 x_2 + \dots + \alpha_n x_n = 0 \quad (\text{where } \alpha_n \text{ are constant})$$

↳ perfect relⁿ with each other.

→ One of the Assumptions of CLRM (classical linear regression model) is that there should be no multicollinearity among the independent variable.

→ As the no. of independent variable $\uparrow$ $\Rightarrow$ problem of multicollinearity can also $\uparrow$.

→ Therefore we need to detect such multicollinearity in data and one such way is by calculating VIF.

⊙ Variance Inflation Factor ⊙

→ It is used to measure identifying existence of multi-collinearity. For eg. consider a regression model,

$$Y_i = \beta_0 + \beta_1 X_1 + \beta_2 X_2 \quad \text{and their } R^2 \text{ value be } R_{12}^2 \text{ then}$$

$$\text{VIF value} = \frac{1}{1 - R_{12}^2} \quad \text{Formula}$$

⇒ VIF greater than 5 requires investigation.

Interpretation ★

$VIF = 1 \Rightarrow$ no correlation b/w independent variables

$1 < VIF < 5 \Rightarrow$ moderate multicollinearity

$VIF > 5 \Rightarrow$ High degree multicollinearity, problematic

→ To fix high multicollinearity we can do :-

① Increase Sample Size.

② Drop non-essential ^{independent} variables.

Outliers Analysis ⇒ Discussed in SLR.

★ Auto correlation

↳ also k/a serial correlation. It happens when there is a relⁿ b/w current obs. and one or more past obs. It occurs when the residual of a regression model exhibit a pattern.

Eg. Measuring temp. every day → If since past 2-3 days temp is high. ⇒ then this has chances that temp. tomorrow will also be high. In other words, autocorrelation occurs when past values in a data series influence future values.

⇒ Durbin-Watson test is a statistical test used to detect the presence of auto correlation. It is based on estimated residuals.

★ Interpretation : Its Range 0 to 4

① A value of 2.0 ⇒ no auto correlation

② Value less than 2.0 ⇒ +ve auto correlation

③ Value greater than 2.0 ⇒ -ve " " (rare, Eg. value = 3.5 means sales this quarter can go down if last past quarter were up. sale High (Past) → Low now.)

High stock
yest or
High stock
today

High temp
yesterday → High temp
tomorrow

Validation of MLR model

→ To validate a MLR model we can follow the given measures and perform below test to check it.

① Coeff. of multiple Determination (R-square) and adjusted R-square

It explain how well
your model explain variation
in the dependant variable.

$$R^2 = 1 - \frac{SSE}{SST}$$

(R^2 = portion of Variance in Y
explained by the model)

→ we use adjusted R^2 more bcoz in MLR, normal R^2 can be manipulated as it can be increased artificially by adding more variables, but adjusted R^2 takes care of that aspect.

→ Adjusted R^2 considered more reliable than R^2 .

★ Interpretation

① $R^2 \Rightarrow R^2$ range 0 to 1. $R\text{-square} = 0 \Rightarrow$ means model doesn't explain any variability in

$\Rightarrow$ more, its near 1, the better fit the model with data is.

E.g. $R\text{-squared} = 70\%$ or 0.70 means 70% of variation in dependant variable (Y) is explained by the independent variable.

② Adjusted R-square $\Rightarrow$ ③ Also, b/w 0 to 1
④ less than or equal to R-square
⑤ More near to 1 the better fit the model is.

② T-test, F-test

$H_0 = \text{no rel. b/w } X \text{ \& } Y$
 $H_A = \text{rel. b/w } X \text{ \& } Y$

calculate p-value $\Rightarrow$ If $p\text{value} < 0.05 \Rightarrow$ H_0 reject, relⁿ b/w $X \text{ \& } Y$
 H_A accept

$p\text{value} > 0.05 \Rightarrow$ opposite of above

③ Check for multicollinearity, if present take app. steps.

④ And, check for autocorrelation (in case of timeseries data)

⑤ Residual Analysis (same as that of SLR $\Rightarrow$ check if Homoscedasticity is followed or not).

Unit-3

★ Logistic and Multinomial Regression

★ Logistic Regression

(Y)

↳ It is used when the dependent variable has two possible outcomes i.e. binary

⇒ Consider it answering a Yes or No question based on some info.

Eg. A doctor says a person has a ^{simple viral} disease (Yes or No) based on their symptoms (like cough and cold).

⇒ Logistic Regression will help in calculating probability the patient has disease based on system. It outputs value 0 and 1, which can be interpreted as probabilities.

★ Multinomial Regression

↳ It is extension of Logistic Regression, but it is used when dependent variable has more than ~~one~~ two categories.

Eg: In above case based on symptoms (like cold, cough, fever) a newly appointed doctor said that the patient ⇒ Yes (Have this normal fever), No (he doesn't have ^{its simply due to weather}), Maybe.

⇒ So, your dependent variable is no longer binary and has more categories.

Eg based on customer age, region, income level ⇒ The ice cream company has three flavours (= Y (dependent variable)) :- Choclate, Vanilla, Strawberry.

⇓

Instead of binary outcome, the outcome/dependent variable has multiple categories.

★ Logistic Fnx (Sigmoid Fnx)

↳ The logistic regression model (or binary regression model) is given by this sigmoid fnx :-

$$f(x) = \frac{e^z}{1 + e^z}$$

$$P(Y=1) = \frac{e^z}{1 + e^z}$$

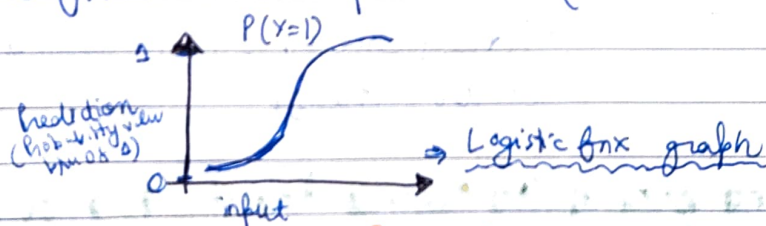
where,

$P(Y=1)$ - probability of obs labelled as 1

$$z = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots$$

↳ X_1, X_2, X_3 are independent var.

⇒ The logistic fnx is a S-shaped curve (as it is a probability fnx).



Estimation of Parameters in Logistic Regression

↳ One of the major assumption of MLR was, residuals follow normal distribution, but here, in logistic regression residuals don't follow normal distribution, so it is difficult to follow a normal distribution. ⇒ Thus we cannot use OLS to calculate parameters here. They are calculated using maximum likelihood estimator (MLE). ↴

using which β_0 & β_1 are calculated.

Interpretation of Parameters of Logistic Regr.

→ An easy interpretation of the regression coeff, β_i , is that:-

① +ve coeff ($\beta_i = +ve$) → means probability of ~~happ~~ event occurring ($P(Y=1)$) also rises

② -ve coeff ($\beta_i = -ve$) → means, probability of event occurring also ↓s.

Eg. Predict student pass ($Y=1$) or fail ($Y=0$) based on no. of hours they study (X)

→ If, $\beta_1 = 0.8$ → means, each additional hour of him/her studying will ↑ the probability of them passing.

$\beta_2 = -0.5$ → means, each additional hour of them studying will ↓ the probability of them passing. (counter-intuitive (consider $X = stress$)).

Estimating probability using logistic Regression

→ It involves fitting data into the model, and then using this model to predict the probability of an event occurring.

Example:- Estimating Probability of getting a loan Afforded.

→ Consider you are working for a bank and want to predict whether a person gets loan afforded ✓ or not X based on their age & income.

Step-1:- Collect Data

① Gather past loan applicants data including their age, income & whether they got their loan afforded or not.

Step-2:- Fit into Regression model.

- ① Use collected data to fit a regression model. The regression model would have an eqn:-

$$\log\left(\frac{P(Y=1)}{1-P(Y=1)}\right) = \beta_0 + \beta_1 \times \text{Income} + \beta_2 \times \text{Credit Score} \quad \text{--- ①}$$

⇒ here, $P(Y=1)$:- probability of loan approval, β_0 is intercept, β_1 and β_2 are coeffs of income and credit score.

Step-3:- Use model for prediction

⇒ once model is fitted, we can predict whether applicant loan will be approved or not.

Eg:- If age = 25 and income = ₹ 50,000 ^{new data}
we put into the above eqn ① to get log odds

↓
then convert that log odd to get the probability using logistic fnn:-

$$P(Y=1) = \frac{e^{\log\text{-odds}}}{1 + e^{\log\text{-odds}}}$$

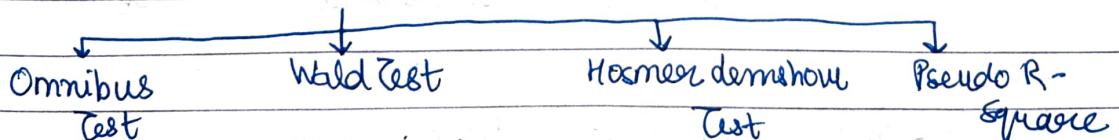
Step-4:- Interpret the results

↳ The result will be b/w 0 to 1 Eg:- If it comes 0.7 or 70%. meaning there is 70% chance that higher loan will be approved.

Logistic Model Diagnostic

⇒ We perform diagnostic test before a binary / logistic regression model is accepted for deployment.

4 test :-



Omnibus :-

- ⊙ To check overall significance of a logistic regression
- ⊙ Checks whether models with predictors (i.e. X or independent var) is significantly better at explaining the variance than a model with no predictors (null model)

Formula $\Rightarrow$ omnibus = $\frac{\text{deviance reduced} - \text{deviance full}}{\text{deviance reduced}}$ > we obtain these value by SPSS (Statistical soft.)

Interpretation :-

$\Rightarrow$ After performing this test $\rightarrow$ then doing hypothesis test $\rightarrow$ If we get $p\text{-value} < 0.05 \Rightarrow$ ^{then it} means, our predictors (X_1, X_2 etc.) were able to explain variance and have a significant impact on Y .

Eg. Model :- predict admission $\checkmark$ or X based on GPA and test-score. This test will show whether they both together significantly predict admission chances than a model that predicts admission without these X_1 & X_2 or predictors.

Wald Test :-

- ⊙ Check significance of individual coeff in the model (eg. checking significance of β_1, β_2 etc.)

Formula $\Rightarrow$ Wald test = $W = \frac{(\text{estimated coeff})^2}{(\text{standard error of coeff})^2}$

Interpretation :-

$\rightarrow$ If the $p\text{-value}$ calculated after Wald test < 0.05 then, the Wald test statistic for a coeff is a significant contributor to model.

$\Rightarrow$ $p\text{-value}$ for β_1 of $X_1 = 0.205 > 0.05 \Rightarrow$ this means X_1 may not have a significant impact on Y .

Hosmer Lemeshow Test

- ⊙ It is a goodness-of-fit test for logistic regression model.
- ⊙ It evaluate the model's performance in predicting outcome based on observed or predicted probabilities.

Interpretation :-

$\rightarrow$ If $p\text{-value}$ high $\Rightarrow$ model a good fit

$p\text{-value}$ low $\Rightarrow$ model not a good fit

Pseudo R-square :-

- Similar to what R^2 measure in LSR, but not directly comparable to it since it's 'Pseudo R-square' [since we don't use OLS here in logistic]

It provides indication of how well predictor variable (X) explain the variance in dependent variable (Y) in logistic regression.

There are 3 types of pseudo R-square

Interpretation

→ In general

→ McFadden
→ Cox and Snell
→ Nagelkerke

⇒ Range 0 to 1 closer to 1 the better

Higher values indicate a better fit.

E.g. Cox and Snell R-square = 0.5961 ⇒ 59% ⇒ meaning that around 59% of variance in the dependent/outcome variable (Y) is explained by the model which is moderate here.

Model Performance → The following things help in determining the model's performance.

★ Classification Table (or confusion matrix)

↓

helps to determine the performance of a model.

Actual (Reality mai hai)

+ve -ve

Predicted (हम अनुमान लगा रहे हैं)	+ve	TP	FP (Type I error)
	-ve	FN (Type II error)	TN

(True = correctly
False = incorrectly)

① Sensitivity / True Positive Rate / Recall

	Actual +ve	Actual -ve
Predicted +ve	TP	FP
Predicted -ve	FN	TN

$$\text{Sensitivity} = \frac{\text{Predicted +ve}}{\text{Actual +ve}}$$

$$= \frac{TP}{TP + FN}$$

② Specificity / True Negative Rate

	Actual +ve	Actual -ve
Predicted +ve	TP	FP
Predicted -ve	FN	TN

$$\text{Specificity} = \frac{\text{Actual -ve}}{\text{Predicted -ve}}$$

$$= \frac{TN}{FP + TN}$$

③ Accuracy :- Accurately diagnosed to the total no. of patients

$$\text{Accuracy} = \frac{TP + TN}{TP + FN + FP + TN}$$

	Actual +ve	Actual -ve
Predicted +ve	TP	FP
Predicted -ve	FN	TN

→ correctly predicted

④ Precision :- ~~correctly predicted~~

$$\text{Precision} = \frac{TP}{TP + FP}$$

	Actual +ve	Actual -ve
Predicted +ve	TP	FP
Predicted -ve	FN	TN

⑤ F-score (or F1 score) :- It is the harmonic mean of Recall and Precision

$$F\text{-score} = \frac{1}{\frac{1}{\text{Recall}} + \frac{1}{\text{Precision}}}$$

$$F\text{-score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

Q.

		Observed	
		Correct	Incorrect
Prediction	Correct	17 ^{TP}	3 ^{FP}
	Incorrect	0 ^{FN}	4 ^{TN}

① Precision = $\frac{17}{20}$, ② Recall = $\frac{TP}{TP+FN} = \frac{17}{17+0} = 1$ ③ Fscore = $\frac{2 \times \frac{17}{20} \times 1}{\frac{17}{20} + 1}$

④ Specificity = $\frac{4}{4+3} = \frac{4}{7}$

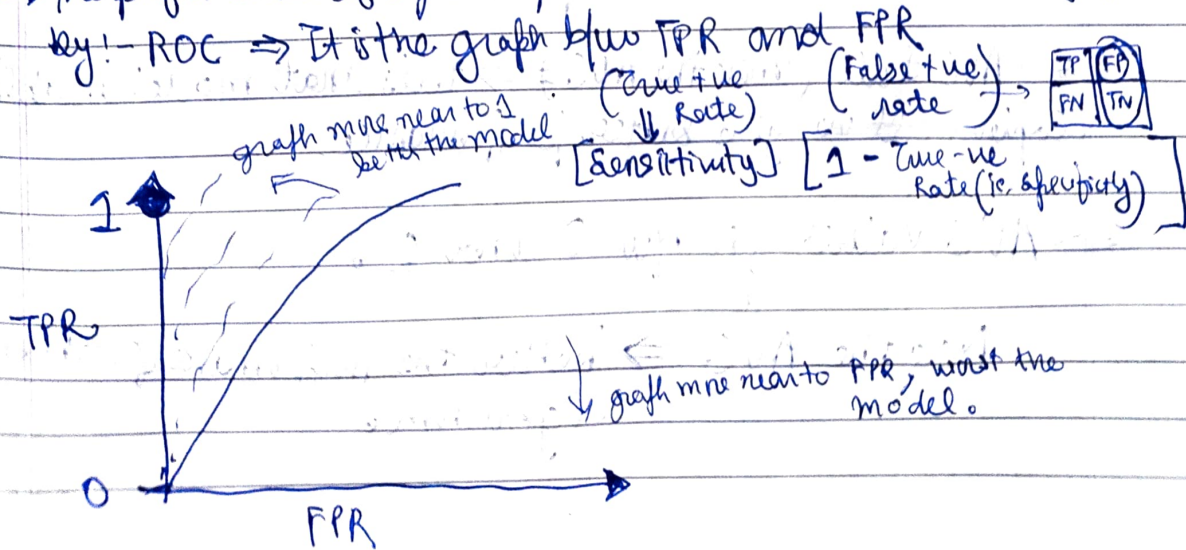
Ginni Coeff / Ginni Index

- It has a value b/w 0 and 1.
- performance of the model improves with ↑ ginni coeff. (more near to 1 → better the model)
- Ginni coeff can be calculated from AUC of ROC.

$$\text{Ginni coefficient} = 2 \times \text{AUC} - 1$$

ROC (Receiver operating characteristic)

- The performance of logistic regression model can be assessed graphically by! - ROC → It is the graph b/w TPR and FPR

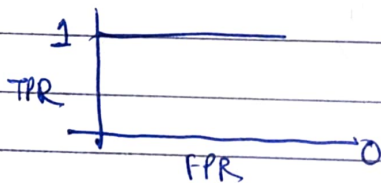


⇒ Basically we plot TPR and FPR at various threshold points then ⇒ ROC curve is formed.

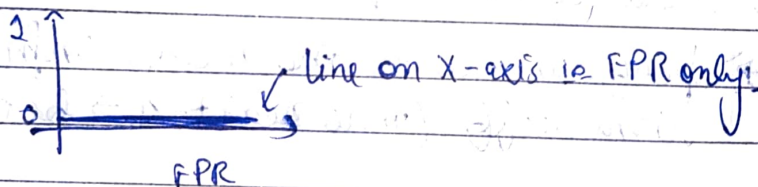
$$TPR = \frac{TP}{TP + FN}$$

$$FPR = \frac{FP}{FP + TN}$$

⇓
⊙ With a TPR = 1 and FPR = 0, a perfect classifier, so, ROC curve would look like :-



⊙ With TPR = 0 and FPR = 1, worst model, ROC curve look like :-



★ Area Under Curve (AUC)

↳ When making a ROC ⇒ AUC is used to access its effectiveness.

⇒ Calculating AUC means calculating the area under ROC curve.
⇒ no formula ⇒ done using software.

⇒ AUC tells (more the AUC) ⇒ tells that model is able to ~~explain all the~~ give better performance.

⇒ AUC values lie b/w 0 and 1.

↓
More the AUC ⇒ better the performance of the algorithm.