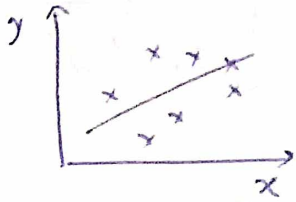


Linear Regression

Best fit Line



$$\hat{y} = mx + c$$

$$\Rightarrow h_{\theta}(x_i) = \theta_0 + \theta_1 x_i$$

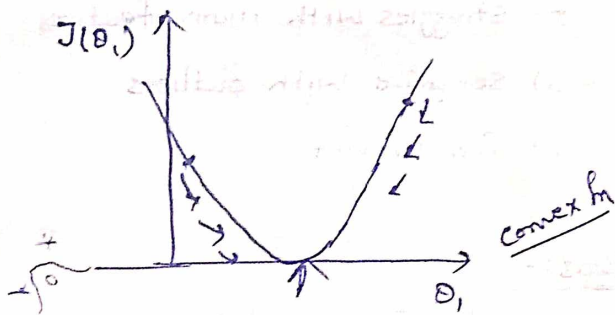
$$h_{\theta}(x) = \theta_0 + \theta_1 x_1 + \theta_2 x_2 + \dots + \theta_n x_n$$

Cost fn :-

$$J(\theta) = \frac{1}{2n} \sum_{i=1}^m [h_{\theta}(x_i) - y_i]^2$$

Residual err

$$= J(\theta_0, \theta_1)$$



$$\theta_j = \theta_j - \alpha \frac{\partial}{\partial \theta_j} [J(\theta)]$$

• mean squared err -

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2$$

⊖: 1) Not Robust to outliers
2) It changes it's unit.

• mean Absolute Err -

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i|$$

⊕: 1) Robust to outliers.
2) Same unit

⊖: 1) Convergence usually takes more time

⊖: 2) optimization is a complex process

• Rmse:-

$$RMSE = \sqrt{MSE}$$

⊕:

⊖:

* Ridge and Lasso Regression

overfitting

Train Accuracy - 90%

Test Accuracy - 70%

Low Bias

High variance

under fitting

Train Accuracy = 60%

Test Accuracy = 62%

High bias

High variance

Generalized model → Train Acc - 90%
Test Acc - 89%

Low Bias & Low var.

• Ridge Regression (L2 Regularization)

$$Loss = MSE + \lambda (\text{slope})^2$$

penalizing ✓

→ This prevent overfitting by Adding penalty.

• Lasso Regression (L1 Regularization)

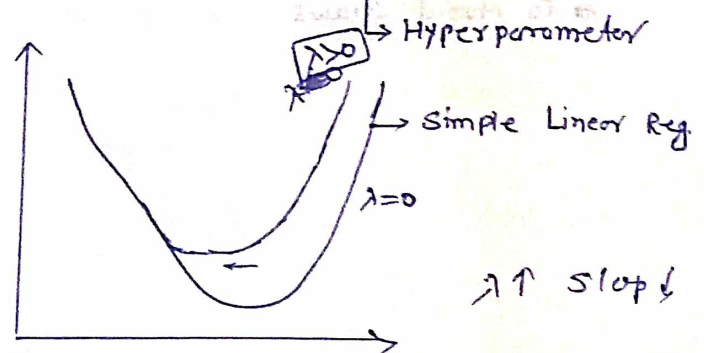
$$Loss = MSE + \lambda |\text{slope}|$$

→ prevent overfitting and helps in feature selection.

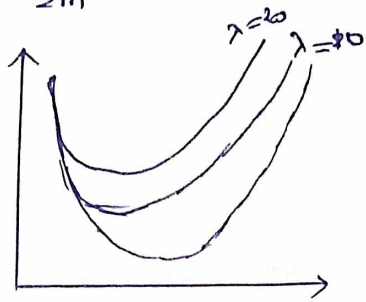
$$\Rightarrow Loss = MSE + \lambda \sum |\text{slope}_i|$$

if Cost fn = $\frac{1}{2m} \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \lambda \sum_{i=1}^n (\text{slope})^2$

Ridge



if Cost fn = $\frac{1}{2m} \sum (y_i - \hat{y}_i)^2 + \lambda \sum |\text{slope}|$



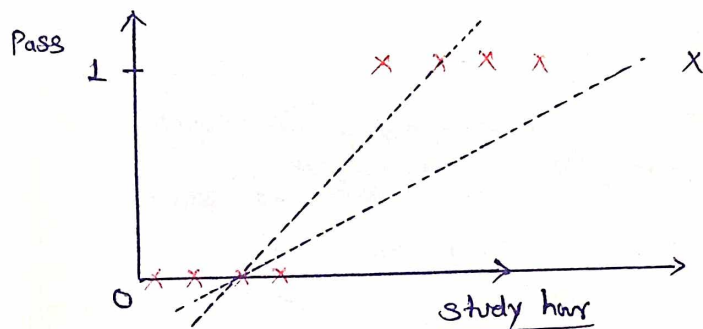
• Elastic-Net Reg.

1) Reduce overfitting

2) feature selection

$$\text{Cost fn} = \frac{1}{2m} \text{MSE} + \lambda_1 \sum (\text{slope})^2 + \lambda_2 \sum |\text{slope}|$$

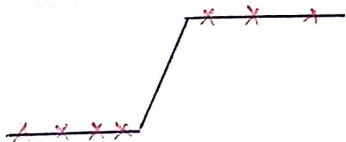
• Logistic Regression → Classification



Problem with Linear Regression

- = { i) outliers
- ii) Sometimes > 1 or < 0

→ so for best fit line we have squash our line with sigmoid fn.



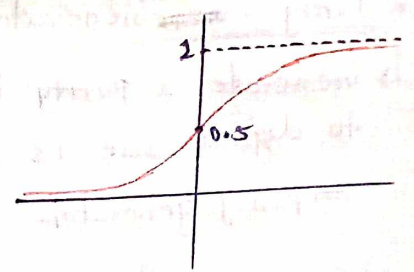
$$h_{\theta}(x) = g(\theta_0 + \theta_1 x)$$

$$\therefore g(z) = \frac{1}{1 + e^{-z}}$$

$$h_{\theta}(x) = \frac{1}{1 + e^{-(\theta_0 + \theta_1 x)}}$$

g : Sigmoid fn

$$g = \frac{1}{1 + e^{-z}}$$



$$z = \theta_0 + \theta_1 x$$

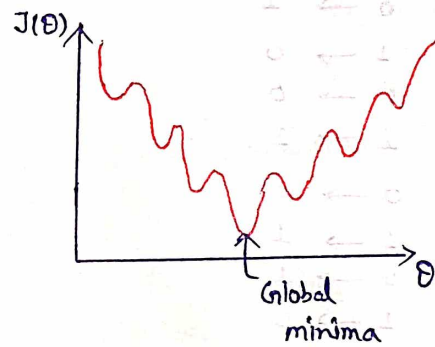
$$z = \theta^T x$$

$$h_{\theta}(x) = \frac{1}{1 + e^{-(\theta^T x)}}$$

Cost fn:

$$J(\theta_0, \theta_1) = \frac{1}{2m} \sum_{i=1}^m [h_{\theta}(x_i) - y_i]^2$$

↓
Non-Convex fn



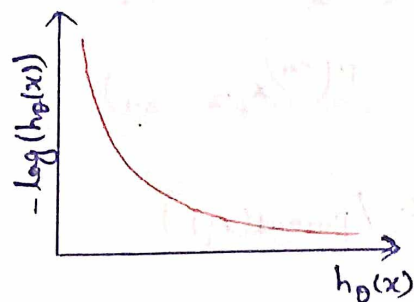
$$J(\theta_0, \theta_1) = \frac{1}{2m} \sum [h_{\theta}(x_i) - y_i]^2$$

↓
Cost ($h_{\theta}(x_i), y_i$)

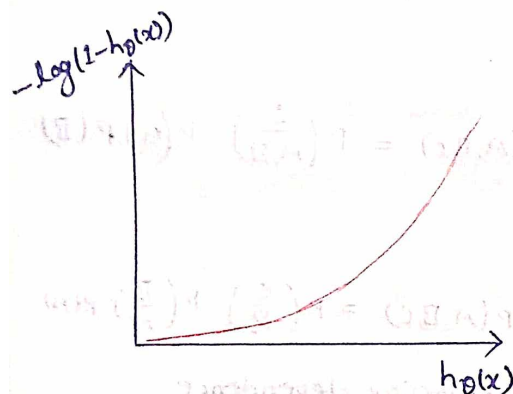
$$\text{Cost}(h_{\theta}(x), y) = \begin{cases} -\log(h_{\theta}(x)) & \text{if } y=1 \\ -\log(1-h_{\theta}(x)) & \text{if } y=0 \end{cases}$$

$$\text{Cost}(h_{\theta}(x), y) = -y \log(h_{\theta}(x)) - (1-y) \log(1-h_{\theta}(x))$$

if $y=1$



if $h_0(x) = 1 \Rightarrow \text{Cost} = 0$



Day	outlook	Temp	Humidity	Play Tennis
D1	Sunny	Hot	High	No
D2				
D3				
D4				

* KNN (K-Nearest Neighbor):

→ used in both Classification and Regression Problem.

→ It Works by finding the 'K' closest data points (Neighbors) to a given input and makes a predictions based on the majority class (for classification) or the Average value (for Regression).

$$J(\theta_0, \theta_1) = \frac{1}{2m} \sum (y_i \log(h_0(x_i)) + (1-y_i) \log(1-h_0(x_i)))$$

$$\theta_j = \theta_j - \alpha \frac{\partial}{\partial \theta_j} J(\theta_0, \theta_1)$$

→ Lazy Learner Algorithm because it does not learn from the training set immediately, despite ~~of~~, it stores entire dataset and perform Computation at the time of classification.

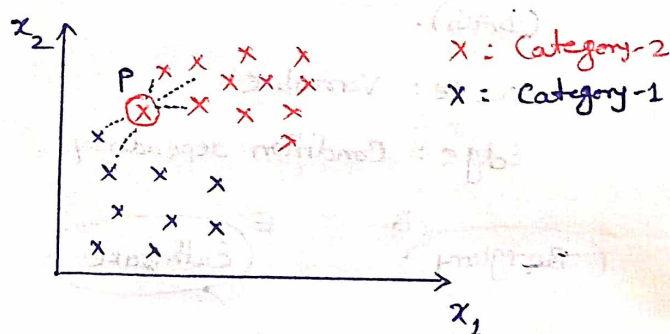
* Naive Bayes Algorithm:-

$$P(B/A) = \frac{P(B) \cdot P(A/B)}{P(A)}$$

If x_1, x_2, x_3 are Events
and y is an output

$$P\left(\frac{y}{x_1, x_2, x_3}\right) = \frac{P(y) \cdot P\left(\frac{x_1, x_2, x_3}{y}\right)}{P(x_1, x_2, x_3)}$$

$$= P(y) \cdot P(x_1/y) \cdot P(x_2/y) \cdot P(x_3/y)$$



→ Point P Classified to ~~Category-1~~ Category-2 because for $K=5$, 3 points of Category 2 are Nearest

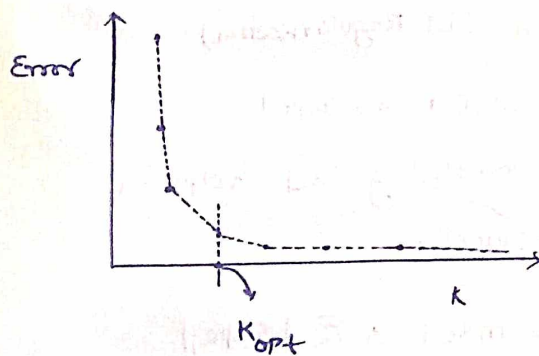
Statistical method for selecting 'K'

Cross-validation:-

This method divides the dataset into K parts. The model is trained on some of these parts and tested on remaining ones. This process is repeated for each part. The K value that gives the highest avg accuracy during these tests is usually the best one to use.

Elbow method:-

We draw a graph showing the error rate or accuracy for different K values. As K increases the error usually drops at first. But after a certain point error stop decreasing quickly. The point where the curve changes the direction and looks like an 'elbow' is usually best choice for K .



• odd values of K

▲ To Avoid chaws

* Distance ~~and~~ metrics:

1) Euclidean distance:

$$\text{dist}(x, x_i) = \sqrt{\sum_{j=1}^d (x_j - x_{ij})^2}$$

2) Manhattan distance

$$d(x, y) = \sum_{i=1}^n |x_i - y_i|$$

3) MinKowski distance

$$d(x, y) = \left[\sum_{i=1}^n (x_i - y_i)^p \right]^{1/p}$$

- ⊖:-
- 1) Slow with large data
 - 2) Struggles with many features
 - 3) Sensitive with outliers
 - 4) Can-overfit

* Bias:-

- Related to Training data
- difference b/w actual or expected values and predicted values are known as Error or bias Error.

→ Let y : true value
 $\hat{y}$: Estimate value

$$\text{Bias}(\hat{y}) = E(\hat{y}) - y$$

Low Bias: Training dataset error ≈ 0

High bias: Training dataset error is high.

↳ A model will not match the training dataset closely.

* Variance :

- Related to Test data.
- This is the measure of spread in data from its mean position.
- Variance is the variability of the model that how much it is sensitive to another subset of training dataset.

→ Let y : Actual value
 $\hat{y}$: predicted value

$$\text{Variance} = E[(\hat{y} - E(\hat{y}))^2]$$

$$\frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2$$

* Performance metrics:-

Dataset

x_1	x_2	y	$\hat{y}$
-	-	0	1
-	-	1	1
-	-	0	0
-	-	1	1
-	-	1	1
-	-	0	1
-	-	1	0

A) Confusion matrix:-

→ Compares the prediction made by the model with the actual results.

	Predicted Positive	Predicted Negative
Actual Positive	True Positive (TP)	False Negative
Actual Negative	False positive (Fp)	True Negative

According to Exa-

		Predicted	
		1	0
Actual	1	3	1
	0	2	1

$$1) \text{ Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

$$\text{Acc} = \frac{4}{7} = 57\%$$

→ This can be misleading when one class is more dominant over the another.

2) Precision:- Used to minimize false Positive

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

3) Recall:- used to minimize FN.

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

$$4) \text{ F-Beta score} = \frac{\text{Precision} * \text{Recall}}{(1 + \beta^2) \beta * (\text{Precision} + \text{Recall})}$$

if FP & FN both are Important

$$\hookrightarrow \beta = 1$$

$$\text{F1 score} = \frac{2PR}{P+R}$$

if $FP \gg FN$, Fp is more important

$$\beta = 0.5$$

$$= \frac{[1.25] PR}{0.25(P+R)}$$

if $FN \gg FP$

$$\beta = 2$$

$$= \frac{5PR}{4(P+R)}$$

5) specificity :-

→ It measures the ability of a model to correctly identify negative instances.

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

6) Type 1 Error = $\frac{FP}{FP + TN}$

Type 2 Error = $\frac{FN}{TP + FN}$

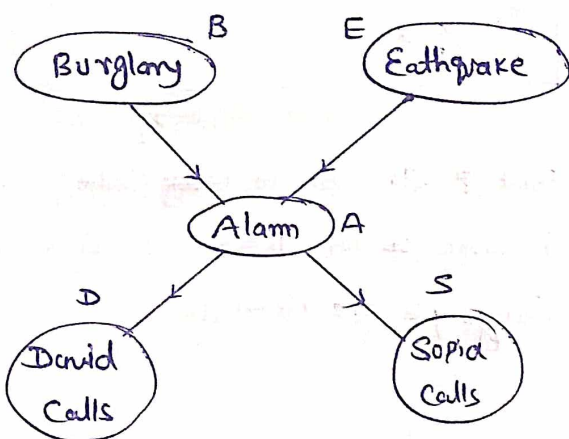
* Bayesian Belief Network (BBN)

→ Bayesian Network or Belief Network

→ It is a graphical model that represents a set of variables and their conditional dependencies using a directed acyclic graph (DAG).

Node: Variable

Edge: Condition dependency.

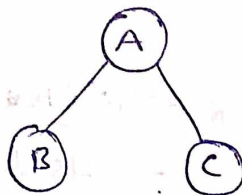


→ Each Node has Conditional Probability Table

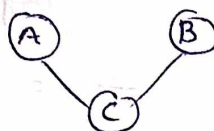
$$P(x_1, x_2, x_3 \dots x_n) = P(x_1) \cdot P(x_2/x_1) \cdot P(x_3/x_1, x_2) \dots$$

$$P(x_n/x_1, x_2 \dots x_{n-1})$$

$$= \prod P(x_i / \text{parent}(x_i))$$



$$\Rightarrow P(A, B, C) = P\left(\frac{B}{A}\right) P\left(\frac{C}{A}\right) P(A)$$



$$\Rightarrow P(A, B, C) = P\left(\frac{C}{A, B}\right) \cdot P(A) P(B)$$



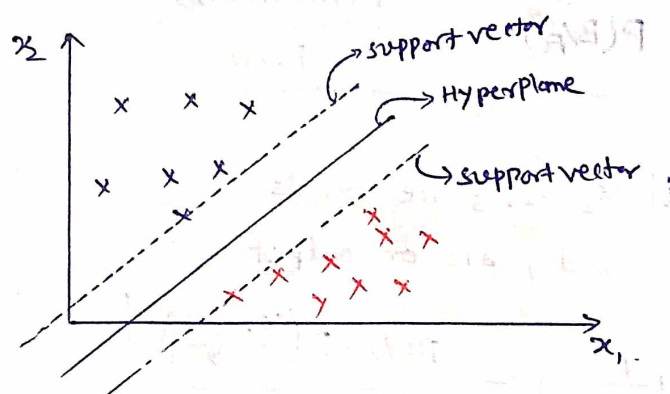
$$\Rightarrow P(A, B, C) = P\left(\frac{C}{B}\right) \cdot P\left(\frac{B}{A}\right) \cdot P(A)$$

↳ Markov dependence

* Support vector machine (SVM) :-

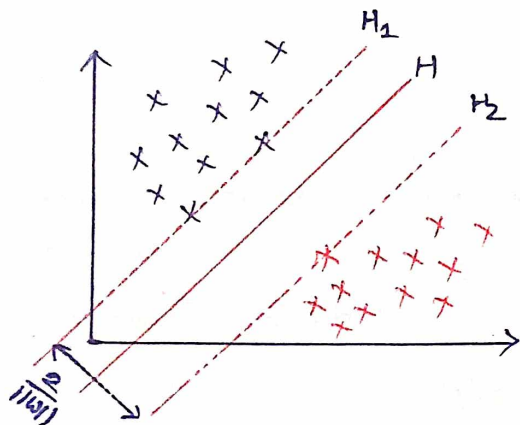
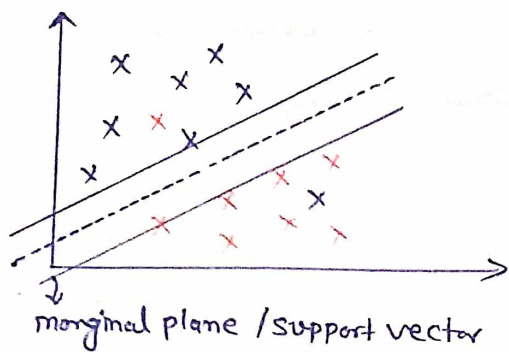
A) Support vector classifier :-

→ The key idea behind the SVM Algorithm is to find the hyperplane that best separates two classes by maximizing the margin between them.



Hard margin :- A maximum margin hyperplane that perfectly separates the data without misclassification.

Soft-margin: Allows some misclassification.



Eqn of hyperplane -

$$H: w^T x + b = 0$$

$$H_1: w^T x + b = +1$$

$$H_2: w^T x + b = -1$$

$$\text{dist}(H_1, H_2) = \frac{2}{\|w\|}$$

We have to maximize marginal plane distance $\Rightarrow$ minimize $\|w\|$

$$y_i = \begin{cases} +1, & w^T x + b \geq 1 \\ -1, & w^T x + b \leq -1 \end{cases}$$

for all accurate datapoint

$$y_i (w^T x + b) \geq 1$$

$$\text{Hinge loss} = \max(0, 1 - y_i (w^T x_i + b))$$

~~hinge~~

$$\text{Cost fn} = \frac{1}{2} \|w\|^2 + C \left[\frac{1}{n} \sum_{i=1}^n \max(0, 1 - y_i (w^T x_i + b)) \right]$$

$$\Rightarrow \text{Minimize}_{w, b, \zeta} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^n \zeta_i$$

No. of misclassified point

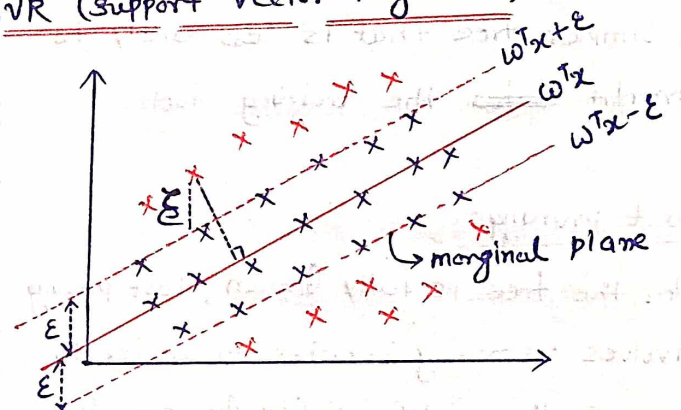
$$\Rightarrow y_i (w^T x_i - b) \geq 1 - \zeta_i$$

$$\Rightarrow \zeta_i \geq 0$$

C : This is a regularization parameter that controls the trade-off b/w margin maximization and penalty for misclassification.

ζ : It is a slack variable that represent the degree of violation of the margin by each data point.

* SVR (Support vector Regression):



$\therefore$ We have to minimize - $\|w\|$

$$\text{And} - |y_i - w^T x_i| \leq \epsilon$$

$$|y_i - \hat{y}| \leq \epsilon$$

for All points -

$$|y_i - w^T x_i| \leq \epsilon + |\zeta_i|$$

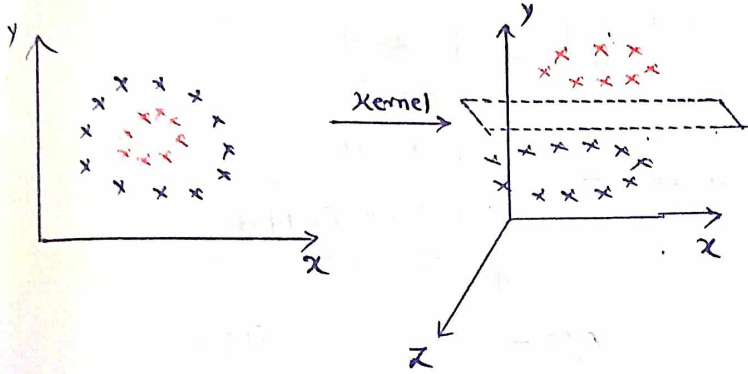
$$\text{Cost fn} = \text{Minimize}_{(w, b)} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^n |\zeta_i|$$

Hinge loss

Q. Draw an ER diagram of Railway management system

SVM Kernels:-

→ A Kernel is a function that transforms the data into a higher dimensional space where it becomes easier to separate using hyperplane.



Why Use ~~K~~ Kernels :-

- 1) Non-Linearity handling
- 2) flexibility: different kernels can be used depending upon the nature of the data and the problem.

3) feature Extraction

→ for Transformed feature $\phi(x)$

$$x_i \longrightarrow \phi(x_i)$$

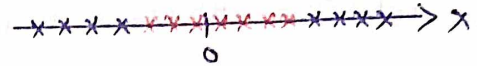
$$K(x_i, x_j) = \phi(x_i) \cdot \phi(x_j)$$

Common SVM Kernels-

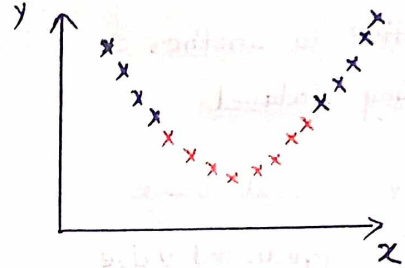
- 1) Linear: $K(x_i, x_j) = x_i \cdot x_j$
- 2) Polynomial: $K(x_i, x_j) = (\gamma x_i \cdot x_j + c)^d$
- 3) Gaussian or Radia Basis-
$$K(x_i, x_j) = e^{(-\gamma \|x_i - x_j\|^2)}$$
- 4) Sigmoid: $K(x_i, x_j) = \tanh(\gamma x_i \cdot x_j + c)$

Exq -

(1d)



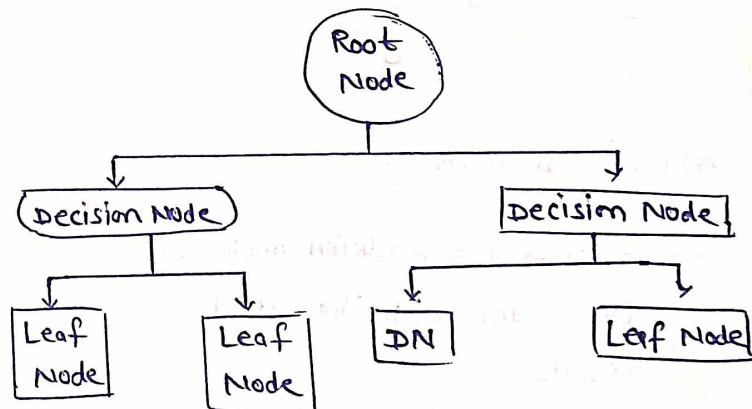
$$\downarrow y = x^2$$



* Decision Tree :-

→ A decision Tree is a supervised learning algorithm used for both classification and Regression tasks. It has a hierarchical tree structure which consists of root node, branches, internal nodes and ~~leaf~~ leaf nodes.

- Internal nodes represent attribute test.
- Branches represent attribute values.
- Leaf nodes represent final decisions or predictions.



• Entropy: This is the measure of Uncertainty of a random Variable.

→ It characterizes the impurity of an arbitrary Collection of Examples.

→ The higher the Entropy more the information Content.

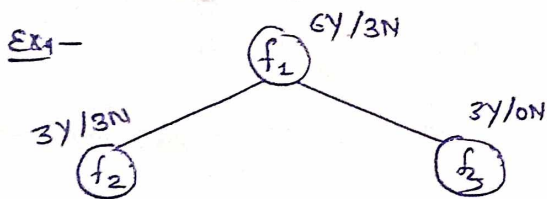
$$H(S) = -P_+ \log_2(P_+) - P_- \log_2(P_-)$$

• Gini Index:-

→ This is a metric to measure how often a randomly chosen element would be incorrectly identified.

→ Attribute With Lower Gini Index should be preferred.

$$G-I = 1 - \sum_{i=1}^n P_i^2$$

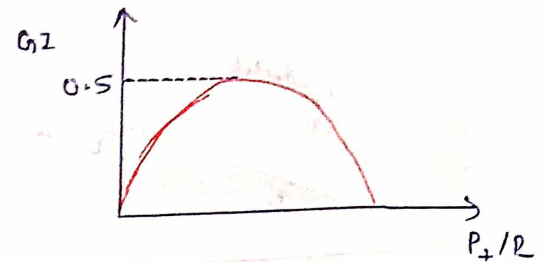
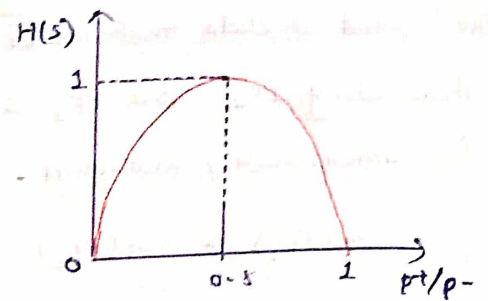


$$\begin{aligned}
 H(f_2) &= -\frac{3}{3} \log_2\left(\frac{3}{3}\right) - \frac{0}{3} \log_2\left(\frac{0}{3}\right) \\
 &= -1 \log_2(1) = \underline{\underline{0}} \text{ (Pure Split)}
 \end{aligned}$$

$$\begin{aligned}
 H(f_3) &= -\frac{3}{6} \log_2\left(\frac{3}{6}\right) - \frac{3}{6} \log_2\left(\frac{3}{6}\right) \\
 &= -\log_2\left(\frac{1}{2}\right) = \underline{\underline{1}} \text{ (Impure Split)}
 \end{aligned}$$

$$\begin{aligned}
 GI(f_2) &= 1 - [P_+^2 + P_-^2] \\
 &= 1 - (0.5^2 + 0.5^2) = 0.5
 \end{aligned}$$

$$GI(f_3) = 1$$



• Information Gain:-

→ Information Gain tells us how useful a question (or feature) is for splitting data into the groups.

→ measures how much uncertainty decreases after split.

→ if S is a set of Instances, A is an attribute.

$$\begin{aligned}
 \text{Gain}(S, A) &= H(S) - \sum \frac{|S_u|}{|S|} \cdot H(S_u) \\
 &\quad \downarrow \\
 &\quad \text{Entropy of Root Node}
 \end{aligned}$$

• Pruning in Decision Tree:

→ decision Pruning is a technique used to prevent decision trees from overfitting the training data.

→ We remove unwanted Nodes.

1) Pre-Pruning:-

- We set max depth to ~~limit~~ limit the decision Tree depth.
- Set a minimum threshold for the Number of Samples in each leaf node. (minimum Samples per leaf)
- specify the minimal Number of Samples needed to break up a node. (minimum Samples per Split)
- Restrict the quantity of features Considered for splitting.
- Maximum depth.
- By Pruning Early, we come to be with a simpler tree that is less likely to overfit ~~the~~ the training facts.

2) Post pruning:-

- After the tree is fully Grown, Post pruning Involves removing branches or nodes to Improve the model's ability to generalize.
- Cost Complexity Pruning
- Reduced Error pruning
- minimum Impurity Decrease
- minimum Leaf size
- This simplifies the tree while preserving its Accuracy.

* PCA (Principle Component Analysis)

→ This is a dimensionality reduction technique used in data Analysis and machine learning.

→ It helps you to reduce the Number of features in a dataset while keeping the most Important Information.

because more Number of features after a limit degrades the performance of model.

3 feature → Accuracy = 60%.

4 feature → Accuracy = 80%.

5 feature → Accuracy = 90%.

7 feature → Accuracy = 85%.

9 feature → Accuracy = 70%.

→ 2 Ways to remove curse of dimensionality -

1) Feature selection

2) Dimensionality Reduction (PCA)

→ We select only Important features.

→ Why dimensionality Reduction -

1) Prevent → curse of dimensionality

2) Improve the performance of model

3) visualize the data and understand the data.

→ To see How features relate to each other wheather they increase or decrease together we use Covariance -

$$\text{Cov}(x, y) = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$$

if $\text{Cov}(x, y) = (+)\text{ive} \Rightarrow x \uparrow y \uparrow$

$\text{Cov}(x, y) = (-)\text{ive} \Rightarrow x \uparrow y \downarrow$

if $\text{Cov}(x, y) = 0 \Rightarrow$ No relation b/w x, y

$$\text{Pearson Co-relation} = \frac{\text{Cov}(x, y)}{\sigma_x \sigma_y}$$

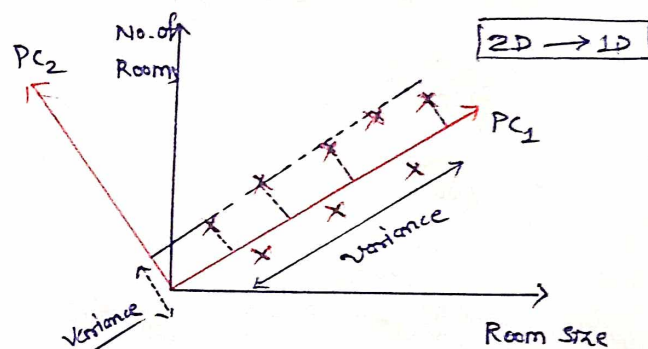
$$-1 \leq \text{Co-relation} \leq 1$$

if $\frac{\text{Cov}(x, y)}{\sigma_x \sigma_y} \rightarrow 1 \Rightarrow$ more (+)ive correlated

$\frac{\text{Cov}(x, y)}{\sigma_x \sigma_y} \rightarrow -1 \Rightarrow$ more (-)ive correlated

$\frac{\text{Cov}(x, y)}{\sigma_x \sigma_y} \rightarrow 0 \Rightarrow$ No correlation

→ PCA Uses linear Algebra to transform data into new features called principal Components. It finds these by calculating Eigen vectors (direction) and Eigen values (importance) from the Covariance matrix.



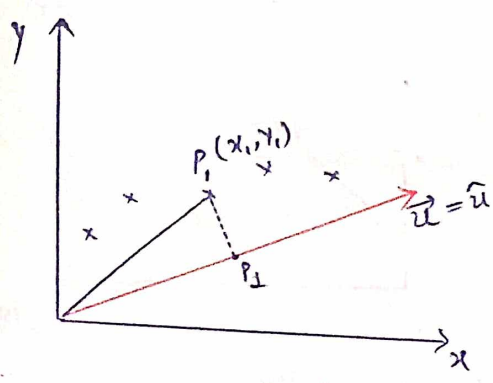
→ PCA identifies the new directions: PC_1, PC_2 .

→ PC_1 Captures the maximum variance meaning it holds the most information.

PC_2 Captures the remaining variance and is PC_1 perpendicular to PC_1 .

→ The spread of data much wider along PC_1 than along PC_2 . So PC_1 is chosen for dimensionality reduction.

$$\text{Var}(PC_1) > \text{Var}(PC_2)$$



$$\text{Proj}_{P_1}(u) = \frac{P_1 \cdot u}{\|u\|}$$

$$\therefore \|u\| = 1$$

$$\text{Proj}_{P_1}(u) = \frac{P_1 \cdot u}{1} = P_1'$$

We get $\{P_0', P_1', P_2', \dots, P_n'\}$

$$\{x_1', x_2', x_3', \dots, x_n'\}$$

$$\text{Max variance} = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$$

↳ Cost fn

Eigen vector or Eigen values -

- 1) find Covariance matrix between features
- 2) Eigen vector and Eigen value will be found out from Covariance matrix

$$AV = \lambda V$$

- 3) Eigen vector with high Eigen value should be chosen,

→ for 2 features - (x, y)

Covariance matrix -

$$A = \begin{matrix} & \begin{matrix} x & y \end{matrix} \\ \begin{matrix} x \\ y \end{matrix} & \begin{bmatrix} \text{Var}(x) & \text{Cov}(x, y) \\ \text{Cov}(x, y) & \text{Var}(y) \end{bmatrix} \end{matrix}$$

for 3 features -

$$A = \begin{bmatrix} \text{Var}(x) & \text{Cov}(x, y) & \text{Cov}(x, z) \\ \text{Cov}(x, y) & \text{Var}(y) & \text{Cov}(y, z) \\ \text{Cov}(x, z) & \text{Cov}(y, z) & \text{Var}(z) \end{bmatrix}$$

$$AV = \lambda V$$

↳ find λ

$$\begin{matrix} \lambda_1 & \lambda_2 & \lambda_3 \\ \downarrow & \downarrow & \downarrow \\ PC_1 & PC_2 & PC_3 \end{matrix}$$