

Part - A

[10 Marks]

1. a) Define ridge regression and mention when it is used. [1M]

Ans: Ridge regression is a regularization technique used in linear regression to address problems like multicollinearity and overfitting. It is also called as L2-regularization. It penalizes the sum of the squared coefficients.

$$\text{Min}(\|Y - X\beta\|^2 + \lambda\|\beta\|^2)$$

b) Define Logistic Regression. [1M]

Ans: Logistic Regression is a supervised learning algorithm used for classification problems. It is a type of classification algorithm that predicts a discrete or categorical outcome.

c) What is Bias-Variance trade-off? [1M]

Ans: Bias is the error from overly simple assumptions in a model, leading to underfitting, while variance is the error from excessive sensitivity to training data leading to overfitting. Bias-variance tradeoff is the fundamental challenge of balancing these two sources of error to create a model that generalizes well to new, unseen data.

d) What is the purpose of Bootstrapping? [1M]

Ans: It is a procedure for estimating the distribution of an estimator by resampling technique that allows you to estimate properties of an estimator by repeatedly drawing samples from the original data.

e) What is Gradient Boosting [1M]

Ans: Gradient Boosting constructs models in a sequential manner where each weak learner minimizes the residual error of the previous one using Gradient descent.

f) Define Ada Boosting

[1M]

Ans: Ada boost is a boosting technique that combines several weak classifiers in sequence to build a strong one.

g) Define the Concept of Backpropagation.

[1M]

Ans: It is a core algorithm in neural networks that works by calculating the error between a network's output and the desired output, then using chain rule to propagate that error backward through the network.

h) Define K-Nearest Neighbor Classifier.

[1M]

Ans: It is non parametric, supervised machine learning algorithm that categorizes the new data points by finding the "K" closest data points on the training set and assigning the new point to the most common class among its neighbors.

i) Define Unsupervised Learning.

[1M]

Ans: It is a type of machine learning algorithms analyze unlabeled data to discover hidden patterns, structures, and relationships without any prior human guidance.

j) List the applications of Association rule Mining. [1M]

Ans:
* Credit card transactions

* Telecom uses for call waiting, call forwarding

* Banking Services

* Market Basket Analysis

* Insurance Claim Frauds

* Medical patient histories.

Part - B

[50 marks]

2. a) How does LDA differ from logistic regression for classification? [5M]

Ans LDA assumes features follow a multivariate normal distribution and have equal covariance matrices across classes, while logistic regression is more flexible and does not require these assumptions, making it more robust to violations. — 1M

LDA estimates the joint probability of the data and class, whereas logistic regression estimates the conditional probability of the class given the data. — 2M.

Key differences

Feature	LDA	Logistic Regression
Assumptions	Assumes features are multivariate Normal forms and have equal covariance matrices across classes.	makes fewer assumptions; does not require normality or equal covariance matrices.
Approach	Models the joint Probability $P(x, y)$ or assumes a specific distribution for x and y .	Models the conditional probability $P(y)$. — 3M
Robustness	less robust to outliers and non normal data	More robust to non-normal data and violations of assumptions. — 4M
Performance	performing well when its assumptions are met	Often performs better in practice, especially with non-normal data, due to its flexibility
Decision boundary	Creates a linear decision boundary under the assumption of equal covariance matrices.	Creates a linear boundary, but can be extended to non-linear boundaries through feature engineering. — 5[M]

2b) Analyze the effect of regularization parameters in Lasso and Ridge [5 M]

Ans: Regularization parameter, often denoted as α (alpha), controls the strength of the penalty applied to the model's coefficients in both Lasso and Ridge regression. — 1 M

A higher α increases the penalty, leading to smaller coefficients for Ridge and the possibility of some coefficients becoming exactly zero for Lasso. This effect shrinks coefficients to combat overfitting, but Ridge retains all features while Lasso performs automatic feature selection. — 2 M

Effect of the regularization parameter (α)

* Low α (closer to zero):

- * The penalty is weak
- * Coefficients are shrunk less, closer to the original least square values.
- * A model with a very low α can still overfit if it's too complex for the data. — 3 M

* High α :

- * The penalty is strong
- * Coefficients are shrunk more aggressively towards zero
- * In Lasso, a higher α increases the likelihood of some coefficients being set to exactly zero, effectively removing features from the model. — 4 M
- * In Ridge, coefficients are shrunk, but none are set to zero exactly zero.
- * If α is too high, model can become too simple, leading to underfitting and loss of predictive power. — 5 M.

3@ What is Lasso regression and Ridge regression? Explain. [5 M]

Ans: A regularization technique L1 Regularization in a regression model. It is also called as Lasso (Least Absolute Shrinkage and selection operator) regression. — 1 M

It adds absolute value of magnitude of the coefficient as a penalty term to the loss function (L). This penalty can shrink some coefficients to zero which helps in selecting only the important features and ignoring the less important ones. — 2 M

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

Where:

m - no. of features:

n - no. of examples

y_i - Actual Target value

$\hat{y}_i$ - Predicted Target value

Ridge Regression

It is a model that uses the L2 regularization technique is called Ridge regression. It adds the squared magnitude of the coefficient as a penalty term to the loss function (L). It handles multicollinearity by shrinking the coefficients of correlated features instead of eliminating them. — 4 M

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

Where

n - No. of examples or datapoints

m - No. of features i.e; predictor variables.

y_i - Actual target value for the i^{th} example

$\hat{y}_i$ - Predicted target value for the i^{th} example.

w_i - Coefficients of the features.

λ - Regularization Parameter that controls the strength of regularization. — 5 M.

3.b) Derive the cost function in Logistic Regression.

[5 M]

Ans Logistic Regression is a supervised learning algorithm used for classification problems. To measure how well the model is predicting or performing, we use a cost function, which tells us how far the predicted values are from the actual ones. In Logistic Regression, the cost function is based on log loss (cross-entropy loss) instead of mean squared error.

— 2 M

- * It measures the error between the predicted probability and the actual class label (0 or 1)
- * Instead of a straight line (like linear regression), logistic regression works with probabilities between 0 and 1 using sigmoid function.
- * The cost penalizes wrong predictions more heavily when the model is confident but wrong.
- * The cost function is defined as

$$\text{cost}(h_0(x), y) = -y \cdot \log(h_0(x)) - (1-y) \cdot \log(1-h_0(x)) \quad \text{— 4 M}$$

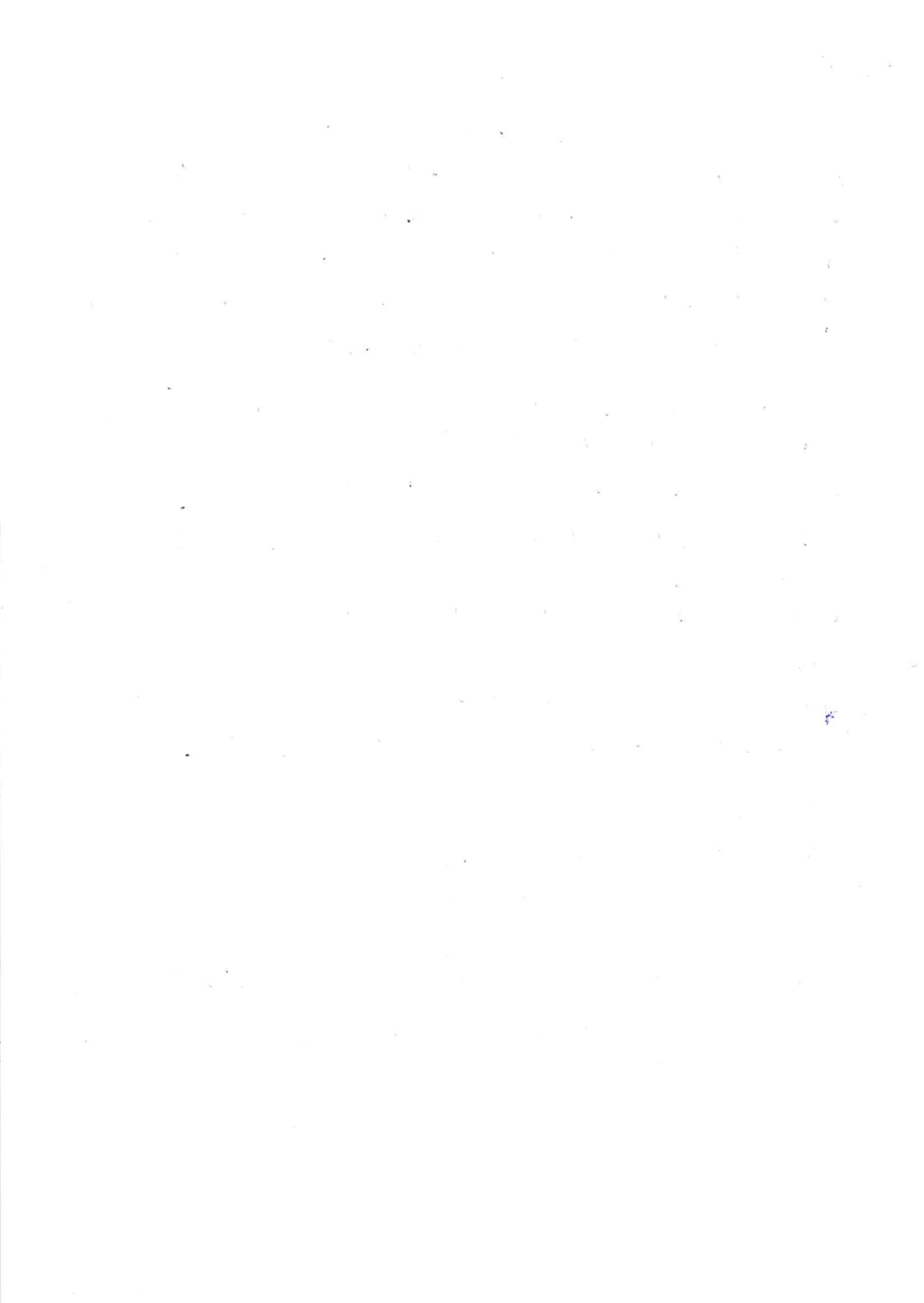
Where:

- * $h_0(x)$: Predicted probability using sigmoid
- * y : Actual value (0 or 1)

For all training examples, the cost function (log loss) becomes:

$$J(\theta) = -\frac{1}{m} \sum_{i=1}^m \left[y^{(i)} \log(h_0(x^{(i)})) + (1-y^{(i)}) \log(1-h_0(x^{(i)})) \right]$$

— 5 M.



Q4. Demonstrate K-fold cross validation with an example. [10]

Ans: K-fold cross validation is a technique that divides a dataset into K-subsets or K-folds, to reliably evaluate a machine learning models performance. In each iteration, one fold is used as a test-set while the remaining K-1 folds are used for training, a process that is repeated K-times with each fold acting as the test-set once. — 2M

The models final performance is then calculated by averaging the results from all K iterations. This method ~~will~~ helps estimate how well the model will generalize to unseen data by training and testing it on different data partitions. — 4M

How it works:-

Divide the data into K equal sized folds. A common value for K is 10.

- * Iterate K-times this process repeats K times.
 - * In each iteration, one fold is designated as the validation set.
 - * The model is trained on the remaining K-folds.
 - * The models performance is evaluated on the validation set.
- 6M

Calculating the final score:

After all K-iterations the performance matrices from each fold are averaged to get a final, more robust estimate of the models performance.

Purpose of K-fold or When use it:

— 8M

- * Reliable performance estimation
- * Model comparison
- * Generalization.

Considerations

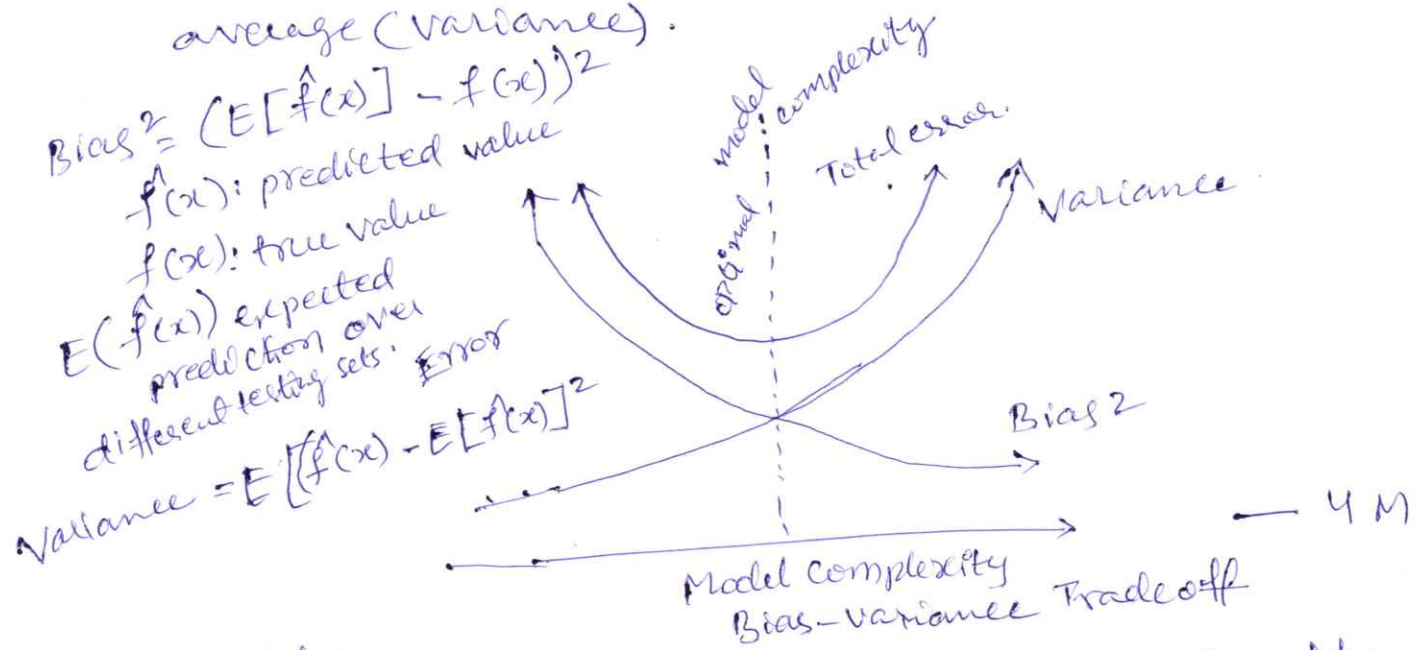
- * Choosing K.
- * Data leakage
- * Stratified K-folds.

10M.

5.A) Derive the relation ship between bias, Variance and MSE. [5M]

Ans The Relation ship between bias, Variance and MSE is $MSE = Bias^2 + Variance$.

Here MSE (Mean squared error) is the sum of the squared bias and the Variance of an estimator. This means an estimators avg error is composed of two parts. How far off the avg prediction is from the true value (bias) and how spreadent the predictions are from that average (Variance).



Model type

underfitting

optimal

overfitting

Bias

High

Moderate

low

Variance

low

Moderate

High

Result

poor training & test performance

Best generalization

poor test-performance

— 5 M.

5B. How can effective no. of parameters be used in model selection. [5 M]

Ans:- The effective no. of parameters is used in model selection to quantify the trade-off between a model's complexity and its goodness of fit, guiding the choice towards the model that balances simplicity with accurate prediction. — 1 M

It measures the actual degrees of freedom a model uses, which can be less than the nominal no. of parameters due to regularization or other factors like non-linearity. — 2 M

By comparing models based on this effective number, techniques like Akaike Information Criterion (AIC) or Bayesian Information Criterion can help selecting a more parsimonious model that generalizes better to new data and avoids overfitting. — 3 M.

How it works

- * Balancing fit & complexity
- * Quantifying complexity.
- * Model selection Criteria
 - AIC
 - BIC.

— 4 M

Applications: This concept is used for tasks like

- * choosing the order of a polynomial.
- * selecting the no. of features or clusters.
- * comparing different models, with non-linearities

— 5 M

6.A) Discuss bagging and boosting techniques. [5 M]

Ans Bagging & Boosting are two types of Ensemble Learning. These two decrease the variance of a single estimate as they combine several estimates from different models. — 1 M

Bagging is a homogeneous weak learners model that learns from each other independently in parallel and combines them for determining the model average.

Boosting also homogeneous weak learners model but works ~~effectively~~ differently from bagging. — 3 M.

In this model, learners learn sequentially and adaptively to improve model predictions of a learning algorithm.

Bagging - Bootstrapping Aggregating also known as Bagging. Example of Bagging: Random Forest

for applications — 4 M.

Differences and similarities — [5 M]

6B. Illustrate the use of additive models of nonlinear relationships.

[5 M]

Ans: Additive models illustrate non linear relationships by extending linear models to include a sum of non-linear functions, each representing the relationship of one predictor variable to the outcome.

Instead of a single line, these models use flexible curves or "smooths" of the predictor variables, rather than just a linear combination. — 1 M

Instead of a single line, relationship is represented by a flexible curve for each predictor, which is built using techniques like splines. This approach provides a balance between the flexibility of complex models and the interpretability of linear models. — 2 M.

How it works

* Extending linear models: An additive model is like a linear model ($Y = b_0 + b_1X_1 + b_2X_2$) but replaces the linear terms with smooth, non-linear functions.

* Non linear functions: The model is expressed as:

$$Y = b_0 + f_1(X_1) + f_2(X_2) + \dots + \epsilon$$

* Smooth functions (f_i): Each function (f_i) captures the non linear relationship between its corresponding predictor X_i and outcome Y .

* Splines: These functions are often built using splines, which are flexible curves that can take on various shapes to fit the data structures. — 4 M

Benefits & limitations:

- * Interpretability
- * flexibility
- * Balance
- * Additive assumption
- * Complexity

— 5 M

⑦ compare additive models, regression trees, and boosting in predicting New Zealand fish size data [10M]

Ans: For predicting New Zealand fish size data, boosting generally offers superior predictive performance compared to generalized additive models (GAMs) and single regression trees, particularly when dealing with complex ecological data involving non-linear relationships and interactions.

— 2 [2M]

Model Comparison

<u>Feature</u>	<u>Additive Models (GAMs)</u>	<u>Regression Trees</u>	<u>Boosting (BRTs)</u>
<u>Predictive performance</u>	Generally good, but often outperformed by ensemble methods like boosting. can overfit if not carefully managed	can have limited predictive power & are prone to overfitting with noisy data or a lack of strong patterns.	High Predictive accuracy; often the best among the 3. Avoids overfitting more flexible, effectively than GAMs or single trees. 3M
<u>Model complexity / Interpretability</u>	High interpretability: clearly shows the smooth, non-linear relationships of each predictor variable with the response	Smooth & highly interpretable.	A "black-box" approach: models are complex ensembles of many trees, making 4M
<u>Handling Non-linearity</u>	Handles non-linear relationships well using smoothing smooth functions	Excels at modeling non linear dependencies and interaction through recursive binary splits.	Combines many simple trees to effectively model complex non-linear responses and interactions. 5M

Data Requirements

Requires careful handling of interaction terms.

less data preparation needed; can handle mixed datatypes.

Requires careful tuning of parameters and is computationally intensive.

— 6 M

Application in Fisheries

widely used for species distribution & catch standardization studies

used but less frequently than GAMs or BRTs due to performance limitations.

Increasingly used in ecology and fisheries for its strong performance in modeling complex environmental data.

— 7 M

Conclusion for New Zealand Fish Size Data

For the specific task of predicting New Zealand fish size data, which ^{involves} ~~involves~~ complex, non-linear relationships between fish size and various environmental factors (e.g., water temperature, depth, food availability, location), boosting is the strongest choice.

— 8 M

Research in similar ecological contexts has shown that Boosting Regression Trees (BRTs) provide substantially superior predictive performance and better handle complex interactions compared to generalized additive models, and they inherently outperform single ~~and~~ regression trees. While BRTs are less ^{9 M} interpretable than the other methods, this accuracy makes them a powerful tool for robust prediction in fisheries management.

— 10 M

⑧ Illustrate the working of K-Nearest Neighbor with a numerical example.

A: KNN is a supervised learning method that classifies a new data point based on the majority class of its 'k' nearest neighbors in the training data.

How KNN works: A numerical example. 1M.

Let's use a simple dataset to predict if a new fruit is an Apple (Class A) or a Banana (Class B), based on its "weight" & "diameter".

<u>Fruit</u>	<u>Weight (gm)</u>	<u>Diameter (cm)</u>	<u>Class</u>
		9	A
F1	70		A
	90	10	A
F2		8	
F3	100		B
	110	11	
F4		9	B
F5	120		B
	130	10	
F6			

New data point (P): we want to classify a new fruit with
Weight = 95 gm and Diameter = 9 cm. - 2M

Choose k: Let's set $k=3$ (we will look at the 3 nearest neighbors).

Step 1: calculate the distance
We use Euclidean distance formula to find the distance between the new point P (95, 9) and all other points in the dataset.

$$\text{Distance}(P, F_i) = \sqrt{(\text{Weight}_P - \text{Weight}_{F_i})^2 + (\text{Diameter}_P - \text{Diameter}_{F_i})^2}$$

- 2M

<u>Fruit</u>	<u>Coordinates</u>	<u>calculation</u>	<u>Distance</u>
F1	(70, 9)	$\sqrt{(95-70)^2 + (9-9)^2} = \sqrt{25^2 + 0^2}$	25
F2	(90, 10)	$\sqrt{(95-90)^2 + (9-10)^2} = \sqrt{5^2 + (-1)^2}$	5.09
F3	(100, 8)	$\sqrt{(95-100)^2 + (9-8)^2} = \sqrt{(-5)^2 + 1^2}$	5.09
F4	(110, 11)	$\sqrt{(95-110)^2 + (9-11)^2} = \sqrt{(-15)^2 + (-2)^2}$	15.13
F5	(120, 9)	$\sqrt{(95-120)^2 + (9-9)^2} = \sqrt{(-25)^2 + 0^2}$	25
F6	(130, 10)	$\sqrt{(95-130)^2 + (9-10)^2} = \sqrt{(-35)^2 + (-1)^2}$	35.01

Step 2: Sort & Find ~~k~~ nearest neighbors

Sort the distances in ascending order & select the top k (3) entries.

<u>Distance</u>	<u>Fruit</u>	<u>class</u>
5.09	F2	A
5.09	F3	A
15.13	F4	B
25	F1	A
25	F5	B
35.01	F6	B

Step 3: Predict the class

We look at the classes of the 3 nearest neighbors we selected in the previous step.

- F2 : class A
- F3 : class A
- F4 : class B

By majority vote, class A appears twice, while class B appears once. Therefore, the new fruit is classified as an Apple (class A).

9a). Compare SVM and neural networks for classification tasks.

A: SVMs are better for smaller, high-dimensional datasets and provide a simpler, more interpretable model with a guaranteed optimal boundary, while neural networks excel on large, complex datasets, learn hierarchical features, and are better suited for tasks like image & speech recognition.

SVM:

Strengths:

- effective in high-dimensional space
- works well for smaller datasets.
- good at finding a clear, optimal decision boundary (hyperplane) by maximizing the margin between classes.
- less prone to overfitting
- can handle non-linear classification using kernels.

Weaknesses:

- computationally expensive
- scales poorly with very large datasets.
- Requires more careful data pre-processing (missing values)
- less effective with large no. of classes.

Interpretability: higher, as the decision boundary is based on support vectors.

Best for: Text, image classification.

Neural Nws:

Strengths:

- Excellent for large datasets & complex, non-linear problems.
- can automatically learn hierarchical features, which is crucial for unstructured data like images & audio.
- Extremely versatile, with deep learning variants.

Weaknesses:

- Requires significant computational resources
- Hyperparameter tuning is more complex
- less interpretable than SVMs, as the decision boundary can be complex & difficult to visualize.
- More prone to overfitting

Interpretability

- lower, as the internal workings of deep nws are often considered a "black box"

Best for: NLP, speech recognition, complex image recognition.

Qb. Write the issues related to overfitting in neural networks.

A: Overfitting in NN occurs when a model learns the training data too well, including noise and ~~ind~~ idiosyncrasies, leading to poor performance on new, unseen data. The key issues are: $(\frac{1M}{2})$

- Poor Generalization:

The most significant issue is the model's inability to generalize to new data. $(\frac{1M}{2})$

- High variance: Overfit models exhibit high variance, meaning they are highly sensitive to the specific training data. $(1M)$

- Memorization Vs. Learning: Instead of learning the underlying patterns and relationships in the data, an overfit model essentially memorizes the training examples. $(1M)$

- Wasted Computational Resources: Training an overly complex model for too long can lead to overfitting and consume significant computational resources. $(1M)$

- Misleading Performance Metrics: High training accuracy can be deceptive, as it doesn't reflect the model's true performance on real-world data. $(1M)$

10). Compare & Discuss how random forests handle missing data and feature correlation with suitable examples.

A: Random forests can handle missing values without

A: explicit imputation in two primary ways:

- Surrogate Splits:

During tree construction, if a primary splitting feature has missing values for certain instances, the algorithm can identify a "Surrogate split" based on a correlated feature.

- Proximity-Based Imputation:

Some implementations of random forests leverage the concept of "proximity" between data points to impute missing values. After building the forest, a proximity matrix is generated, indicating how often two instances end up in the same leaf node.

Handling Feature Correlation:

Random forests are inherently less sensitive to multicollinearity compared to linear models for following reasons:

- Random feature subsampling:

At each split in a decision tree, only a random subset of features is considered for finding the best split. This reduces the chance that highly correlated features will consistently be chosen together, mitigating their combined influence.

- Ensemble Averaging: The final prediction in a random forest is an aggregation of the predictions from multiple independent decision trees. Even if some individual trees are influenced by correlated features, the averaging process tends to smooth out these individual biases.

— 2 M (2 marks)

Handling Missing values with Random Forest using python.

Here we focus on predicting missing 'Age' values in Titanic dataset which is a classic dataset used in Machine Learning and data analysis.

* Implementing program.

```
# Import Libraries.  
import pandas as pd  
import numpy as np
```

— 2 M

* Loading Dataset

```
# Importing dataset & setting 'PassengerId' as Index  
Data = pd.read_csv('Data.csv', index_col='PassengerId')  
Data.head()
```

output will display dataset

* Data preprocessing

Handling Missing values.

code Data.isnull().sum() is used to check for missing values in a Data frame called Data.

Missing values

```
Data.isnull().sum()
```

Dropping 'Cabin' column due to missing values.

```
Data = Data.drop(columns=['Cabin'], axis=1)  
Data.head()
```

— 2 M

Data.Embedded.value_counts():

* Model Building

1. Importing Random Forest Regressor
2. Creating Random Forest Regressor
3. Training the model

* Prediction

Predicted_Ages = rf_age.predict(Independent_Feature)

11a) Write the Hierarchical clustering and its linkage methods.
SM.

A: Hierarchical clustering groups data points into a tree-like hierarchy, either by progressively merging clusters or splitting a single cluster (divisive). The key to this process is the linkage method, which determines the distance between clusters. Common linkage methods include single linkage (minimum distance), complete linkage (maximum distance), average linkage (average distance), and Ward's method (minimizing the increase in squared error).
LM.

Hierarchical Clustering Methods:
QM.

→ Agglomerative clustering:

This is a bottom-up approach where each data point starts as its own cluster. Clusters are then merged based on similarity until only one cluster remains.

→ Divisive clustering:

This is a top-down approach where all data points are in their own cluster.

→ Dendrogram:

The final output is a tree diagram called a dendrogram, which visualizes the hierarchy of clusters of clusters and the relationships between them. (2 Marks)

Linkage Methods:

(2 Marks)

→ Single linkage: The distance between two clusters is the minimum distance between any two points in the two clusters.

$\left(\frac{1}{2} M\right)$.

→ Complete Linkage:

The distance between two clusters is the maximum distance between any two points in the two clusters.

$\left(\frac{1}{2} M\right)$.

11b). Compare the Random forests with Decision Trees?

A: Random Forests and Decision Trees are both widely used machine learning algorithms, but they differ significantly in their structure and performance characteristics.

Decision Trees:

- Single Model: A decision tree is a single, tree-like model that makes predictions by following a series of rules based on input features.
- Interpretability: They are highly interpretable and easy to visualize, making it straightforward to understand the decision-making process.
- Overfitting: Decision trees are prone to overfitting, especially when allowed to grow deep without pruning, as they can become too specific to the training data and fail to generalize to new data.
- Computational Efficiency:
They are computationally efficient and fast to train and predict. (2½ marks)

Random Forests:

- Ensemble Model: A random forest is an ensemble learning method that combines multiple decision trees. It builds a "forest" of trees, each trained on a random subset of the data.

- Reduced Overfitting: ϕ

By averaging the predictions of multiple trees, random forests significantly reduce the risk of overfitting and improve generalization.

- Interpretability: Less interpretable than single decision trees due to the ensemble nature, making it harder to visualize the entire decision-making process.

- Computational cost:

More computationally intensive than single decision trees due to the training of multiple trees.

(2½ Marks)