

Machine Learning (R22CS602PG)

R22

III Year II Sem Supply Exam - Dec - 2025 (19/12/2025)

Key Question Paper.

PART-A

1) a) Interpret handwriting recognition as well Posed learning Problem.

Sol: Handwriting recognition is a well Posed learning Problem because the task is to identify handwritten characters, the experience consists of labeled handwriting samples, and Performance is measured by recognition accuracy.

2) b) Consider the hypothesis $h_1 = \langle \text{Sunny}, ?, ?, ?, \text{Change} \rangle$. Justify the consistency of h_1 over the instance (or example) $d_1 = \langle \text{sunny}, \text{warm}, \text{Normal}, \text{strong}, \text{warm}, \text{same} \rangle$ whose target label is "Yes."

Ans: The hypothesis $h_1 = \langle \text{sunny}, ?, ?, ?, \rangle$ is consistent because the wildcard $?$ can match any attribute value, so it correctly classifies all possible positive training instances as "Yes";

3) Differentiate between a Perceptron and multilayer Perceptron.

Ans: c) A Perceptron is a single layer neural network that solves only linear Problems, whereas multilayer Perceptron has hidden layers and can solve non-linear Problems.

4) d) State SVM Principles.

Ans: Support vector machine works on the Principle of finding an optimal hyperplane that maximizes the margin between different classes.

e) Define Entropy.

Ans: Entropy is a measure of uncertainty or impurity in a dataset, commonly used in decision trees to determine the best attribute for splitting the data.

f) List any two algorithms which use bagging and boosting.

Ans: Bagging: Random Forest, Bagged decision trees
Boosting: AdaBoost, Gradient Boosting.

g) Perform single-point crossover at position 3 between binary strings: A = 110101, B = 001110

Ans: Single-point crossover at position 3

A = 110|101

B = 001|110

Offspring 1: 110110

Offspring 2: 001101

h) LDA is best suited for which type of problem: classification or Regression.

Ans: Linear Discriminant Analysis (LDA) is best suited for classification problems.

i) Write state update rule in Q-Learning?

Ans: $Q(s,a) \leftarrow Q(s,a) + \alpha [r + \gamma \max_{a'} Q(s',a') - Q(s,a)]$

Where α = learning rate

γ = discount factor

r = reward

s' = next state.

2

1) Differentiate between Sampling with and without Replacement.

Ans: With Replacement: An item can be selected more than once in the sample.

Without Replacement: An item cannot be selected again once it is chosen.

PART B

2) Describe the Linear Regression algorithm. Derive the formula for the weights of a linear regression model.

Ans: Linear Regression Algorithm and Derivation of Weight Formula.

(i) Introduction to Linear Regression.

Linear Regression is a supervised learning algorithm used to model the relationship between a dependent variable (output) and one or more independent variables (inputs) by fitting a linear equation to observed data.

It is mainly used for Prediction and Regression Problems.

(ii) Linear Regression Model;

For a dataset with n samples and m features;

$$y = w_0 + w_1x_1 + w_2x_2 + \dots + w_mx_m$$

Where:

y = Predicted output
 $x_1, x_2, \dots, x_m$ = i/p features.

w_0 = bias
 $w_1, w_2, \dots, w_m$ = weights (Parameters)

(iii) Objective of Linear Regression:

The goal is to find weights w such that the difference between Predicted and actual values is minimized. This difference is measured using the Mean Squared Error (MSE) cost function:

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

(iv) Error Function

Using matrix function

$$J(w) = \frac{1}{2} (y - Xw)^T (y - Xw)$$

(v) Derivation of weight Formula (Normal equation)

Take derivative

Step 1 : $\frac{\partial J}{\partial w} = -X^T (y - Xw)$

Step 2 : Set derivative to zero

$$X^T (y - Xw) = 0$$

Step 3 : Expand

$$X^T y = X^T X w$$

Step 4 : Solve for w

$$w = (X^T X)^{-1} X^T y$$

(vi) Final weight Formula : $w = (X^T X)^{-1} X^T y$

(vii) Algorithmic steps for Linear Regression:

1. Collect training dataset (x, y)
2. Assume linear model $(y = Xw)$
3. Define cost function (MSE)
4. Minimize cost function
5. Compute weight using normal equation
6. Use the model for Prediction

3) Find the Version Space for the given data using Candidate Elimination algorithm.

Origin	manufacturer	Color	Decade	Type	Example Type.
Japan	Honda	Blue	1980	Economy	Positive
Japan	Toyota	Green	1970	Sports	Negative
Japan	Toyota	Blue	1990	Economy	Positive
USA	Chrysler	Red	1980	Economy	Negative
Japan	Honda	White	1980	Economy	Positive.

Ans: Using the Candidate Elimination Algorithm, we find the Version Space by computing the specific boundary (s) and general boundary (g).

Given attributes.

$\langle \text{Origin, Manufacturer, Color, Decade, Type} \rangle$

Training examples (from the table)

1. $\langle \text{Japan, Honda, Blue, 1980, Economy} \rangle \rightarrow \text{Positive}$
2. $\langle \text{Japan, Toyota, Green, 1970, Sports} \rangle \rightarrow \text{Negative}$
3. $\langle \text{Japan, Toyota, Blue, 1990, Economy} \rangle \rightarrow \text{Positive}$
4. $\langle \text{USA, Chrysler, Red, 1980, Economy} \rangle \rightarrow \text{Negative}$
5. $\langle \text{Japan, Honda, White, 1980, Economy} \rangle \rightarrow \text{Positive}$

Step 1. Initialize.

$$S_0 = (\emptyset, \emptyset, \emptyset, \emptyset, \emptyset)$$

$$G_0 = (?, ?, ?, ?, ?)$$

Step 2: Process Positive examples.

After Processing all Positive examples, the specific boundary becomes

$$S = \langle \text{Japan}, ?, ?, 1980, \text{Economy} \rangle$$

Step 3: Process Negative examples.

After removing hypotheses in G that cover negative examples and keeping only those consistent with S , the general boundary is

$$G = \{ \langle \text{Japan}, ?, ?, ?, \text{Economy} \rangle, \langle ?, ?, ?, 1980, \text{Economy} \rangle \}$$

The version space is the set of all hypotheses between S and G where:

- Specific boundary (S): $\langle \text{Japan}, ?, ?, 1980, \text{Economy} \rangle$
- General boundary (G):

$$\{ \langle \text{Japan}, ?, ?, ?, \text{Economy} \rangle, \langle ?, ?, ?, 1980, \text{Economy} \rangle \}$$

This represents all hypotheses consistent with the given training data.

4) Derive back Propagation rule for both hidden and o/p layers,

Ans: Consider a simple feedforward neural network;

→ Input: x

→ Hidden layer $h = f(z^h) = f(w^h x + b^h)$

→ Output layer $y_{\text{Pred}} = g(z^o) = g(w^o h + b^o)$

→ Loss function: $L(y_{\text{Pred}}, y)$

Here: w^h, b^h = weights and biases of hidden layer.

w^o, b^o = weights and biases of output layer.

f = activation function in hidden layer

g = activation function in output layer.

We aim to compute the gradients;

$$\frac{\partial L}{\partial w^o}, \frac{\partial L}{\partial b^o}, \frac{\partial L}{\partial w^h}, \frac{\partial L}{\partial b^h}$$

(i) Gradients for o/p layer.

For o/p layer, using chain rule:

$$\frac{\partial L}{\partial w^o} = \frac{\partial L}{\partial z^o} \cdot \frac{\partial z^o}{\partial w^o}$$

$$\text{but } z^o = w^o h + b^o \Rightarrow \frac{\partial z^o}{\partial w^o} = h^T$$

Define the error term at o/p layer as: $\delta^o = \frac{\partial L}{\partial z^o} = \frac{\partial L}{\partial y_{\text{Pred}}} \odot g'(z^o)$

$$\text{Then } \frac{\partial L}{\partial w^o} = \delta^o (h)^T \Rightarrow \frac{\partial L}{\partial b^o} = \delta^o$$

(ii) Gradients for Hidden Layer

For the hidden layer, the error propagates backward from the output:

$$\frac{\partial L}{\partial w^h} = \frac{\partial L}{\partial z^h} \cdot \frac{\partial z^h}{\partial w^h}$$

$$z^h = w^h x + b^h \Rightarrow \frac{\partial z^h}{\partial w^h} = x^T$$

Error term in hidden layer

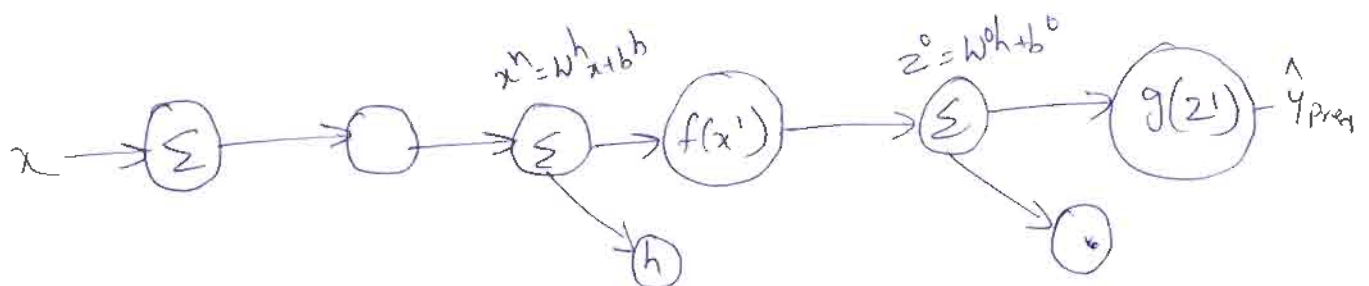
$$\delta^h = \frac{\partial L}{\partial z^h} = (w^0)^T \delta^0 \odot f'(z^h)$$

So the gradient updates are:

$$\frac{\partial L}{\partial w^h} = \delta^h (x)^T$$

$$\frac{\partial L}{\partial b^h} = \delta^h$$

forward pass



$$\delta^h = (w^0)^T \delta^0 \odot f'(z^h)$$

$$\frac{\partial L}{\partial w^h} = \delta^h x^T$$

$$\frac{\partial L}{\partial b^h} = \delta^h$$

$$\delta^0 = \frac{\partial L}{\partial z^0} \odot g'(z^0)$$

$$\frac{\partial L}{\partial w^0} = \delta^0 h^T$$

$$\frac{\partial L}{\partial b^0} = \delta^0$$

backward pass

5) Explain the role of basis functions in interpolation. Describe how splines are used as basis functions for smooth interpolation of data points.

Ans: Interpolation is the process of estimating values of a function $f(x)$ at points where it is not explicitly known, using a set of known data points (x_i, y_i) .

A basis function is a building block used to construct the interpolating function. Mathematically, the interpolated function is expressed as:

$$f(x) = \sum_{i=1}^n c_i \phi_i(x)$$

where $\phi_i(x)$ = basis functions

c_i = Coefficients determined by the known data points
 n = No. of data points.

Role of basis functions:

- (i) They define the shape of the interpolating function b/w points.
- (ii) They allow interpolation to be written as a linear combination, making computations systematic.
- (iii) They determine the smoothness and continuity of the interpolated function.
- (iv) Different choices of basis functions lead to different interpolation methods (e.g., Polynomials, Piecewise linear, Splines, Radial basis functions).

example: In Polynomial interpolation, the basis functions are $\phi_i(x) = x^{i-1}$.

6) Explain the Concepts of bagging and boosting in ensemble learning. Compare their mechanisms and use cases.

Ans: Ensemble learning combines multiple models to improve overall performance compared to a single model. The main idea is:

Final Prediction - Aggregation of Predictions from multiple models,

Two Popular ensemble techniques are Bagging and Boosting.

→ Bagging (Bootstrap aggregating):

Bagging is designed to reduce variance and avoid overfitting by training models on different subsets of data.

Example algorithm for Bagging is Random forest algorithm.

→ Boosting

Boosting is designed to reduce bias by sequentially training models, each focusing on the mistakes of the previous ones.

Example algorithms are Adaboost, Gradient Boosting Machine, XG Boost, LightGBM, CatBoost.

So finally choose Bagging, when models overfit (high variance) boosting when models underfit (high bias).

1) Construct decision tree for the following data.

Id	outlook	Temperature	Humidity	wind	Play Tennis.
1.	Sunny	Hot	High	weak	0
2.	Sunny	Hot	High	strong	0
3.	Overcast	Hot	High	weak	1
4.	Rainy	Mild	High	weak	1
5.	Rainy	Cool	Normal	weak	1
6.	Rainy	Cool	Normal	strong	0
7.	Overcast	Cool	Normal	strong	1
8.	Sunny	Mild	High	weak	0

Ans: step 1. Identify the Target and Attributes.

→ Target : Play Tennis (Yes/No)

→ Attributes : Outlook, Temperature, Humidity, Wind

step 2: Calculate the Entropy of the Entire Dataset.

$$\text{Entropy}(S) = -P_+ \log_2 P_+ - P_- \log_2 P_-$$

P_+ = Proportion of Yes (play tennis = 1) = $5/8$

P_- = Proportion of No (Play tennis = 0) = $3/8$

$$\text{Entropy}(S) = -\frac{5}{8} \log_2 \frac{5}{8} - \frac{3}{8} \log_2 \frac{3}{8} \approx 0.954$$

Step 3: Compute Entropy for each Attribute.

a) Outlook

Outlook	Play=Yes	Play=No	Entropy
Sunny	2	1	$E = -2/3 \log_2 2/3 - 1/3 \log_2 1/3$ ≈ 0.918
Overcast	2	0	$E = 0$
Rainy	1	2	$E = 0.918$

$$\text{Info}(S, \text{Outlook}) = \frac{3}{8} (0.918) + \frac{2}{8} (0) + \frac{3}{8} (0.918) \approx 0.689$$

$$\text{Gain}(S, \text{Outlook}) = 0.954 - 0.689 = 0.265$$

b) Temperature.

Temperature	Play=Yes	Play=No	Entropy
Hot	2	2	$E = 1$
Mild	2	0	$E = 0$
Cool	1	1	$E = 1$

$$\text{Info}(S, \text{Temp}) = \frac{4}{8} (1) + \frac{2}{8} (0) + \frac{2}{8} (1) = 1$$

$$\text{Gain}(S, \text{Temp}) = 0.954 - 1 = -0.046 \approx 0$$

c) Humidity

Humidity	Play(Yes)	Play(No)	Entropy
High	3	2	$E = 0.971$
Normal	2	1	$E = 0.918$

$$\text{Info}(S, \text{Humidity}) = \frac{5}{8} (0.971) + \frac{3}{8} (0.918) \approx 0.954$$

$$\text{Gain}(S, \text{Humidity}) = 0.954 - 0.954 \approx 0$$

Wind

Wind	Play(Yes)	Play(No)	Entropy
Weak	4	1	$E=0.722$
Strong	1	2	$E=0.918$

$$Info(S, Wind) = \frac{5}{8} (0.722) + \frac{3}{8} (0.918) \approx 0.788$$

$$Gain(S, Wind) = 0.954 - 0.788 = 0.166$$

∴ Outlook has highest gain = 0.265 → Root node.

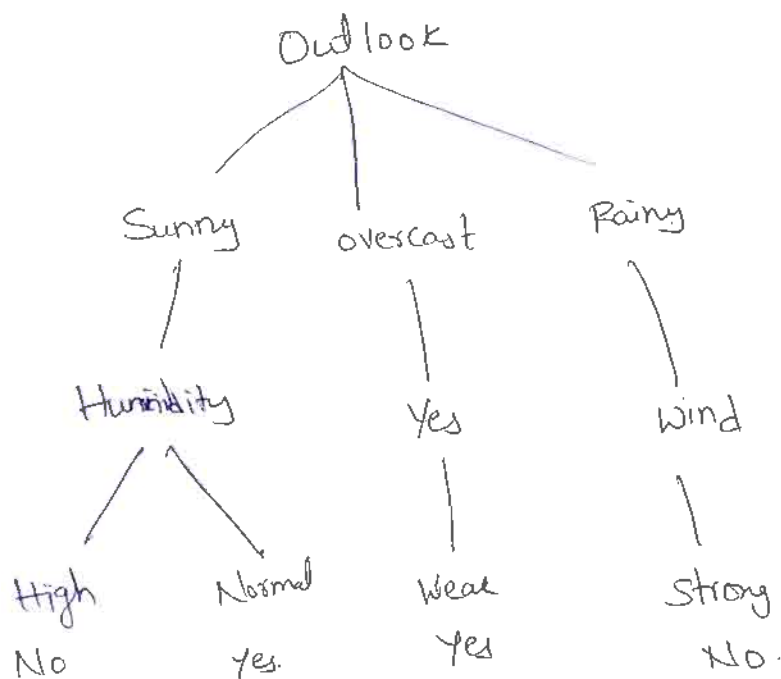


Fig: Decision Tree.

8) Describe the LDA (linear discriminant analysis) algorithm.
How does LDA Preserve local neighbourhood information while reducing dimensions.

Ans: LDA is a supervised dimensionality reduction technique used to project high dimensional data onto a lower-dimensional space while maximizing class separability.

Algorithm steps:

1. Compute the mean vectors:

For each class i , calculate the mean vector μ_i .

2. Compute Scatter matrices:

→ within-class scatter matrix (S_W): measures the spread of samples within each class:

$$S_W = \sum_{i=1}^C \sum_{x \in C_i} (x - \mu_i)(x - \mu_i)^T$$

→ Between-class scatter matrix (S_B): measures the spread b/w

class means

$$S_B = \sum_{i=1}^C n_i (\mu_i - \mu)(\mu_i - \mu)^T$$

where μ = overall mean, n_i = no. of samples in class i .

3. Compute the Projection Matrix:

Find the matrix W that maximizes the ratio of b/w-class to within-class scatter:

$$W = \arg \max |W^T S_B W| / |W^T S_W W|$$

generalized eigenvalue = $S_B W = \lambda S_W W$

4. Project data:

Transform the original data x into lower-dimensional space: $Y = W^T x$.

9. Explain the concept of PCA (Principal Component Analysis) to reduce the dataset from 2D to 1D.

Ans: → PCA is a dimensionality reduction technique used in data analysis. It is unsupervised learning algorithm.

→ It transform the original features into a new set of uncorrelated variable called Principal components,

→ The goal is to capture as much variance as possible in fewer dimensions.

→ Suppose you have a dataset with two features: x_1 and x_2 .

→ Sometimes, these two features are correlated, so a lot of the information is redundant.

→ PCA finds a new axis (Principal Component) that best represents the data in a single dimension.

Steps for Project 2D → 1D PCA

Step 1: Center the data.

→ Subtract the mean of each feature to center the data at the origin

$$X_{\text{centered}} = X - \text{mean}(X)$$

Step 2: Compute the Covariance matrix

→ Covariance measures how features vary together. For 2D,

$$\text{Cov}(X) = \begin{bmatrix} \text{Var}(x_1) & \text{Cov}(x_1, x_2) \\ \text{Cov}(x_1, x_2) & \text{Var}(x_2) \end{bmatrix}$$

Step 3: Compute eigenvectors and eigenvalues.

Step 4: Project the data onto first Principal component.

Steps: Result,

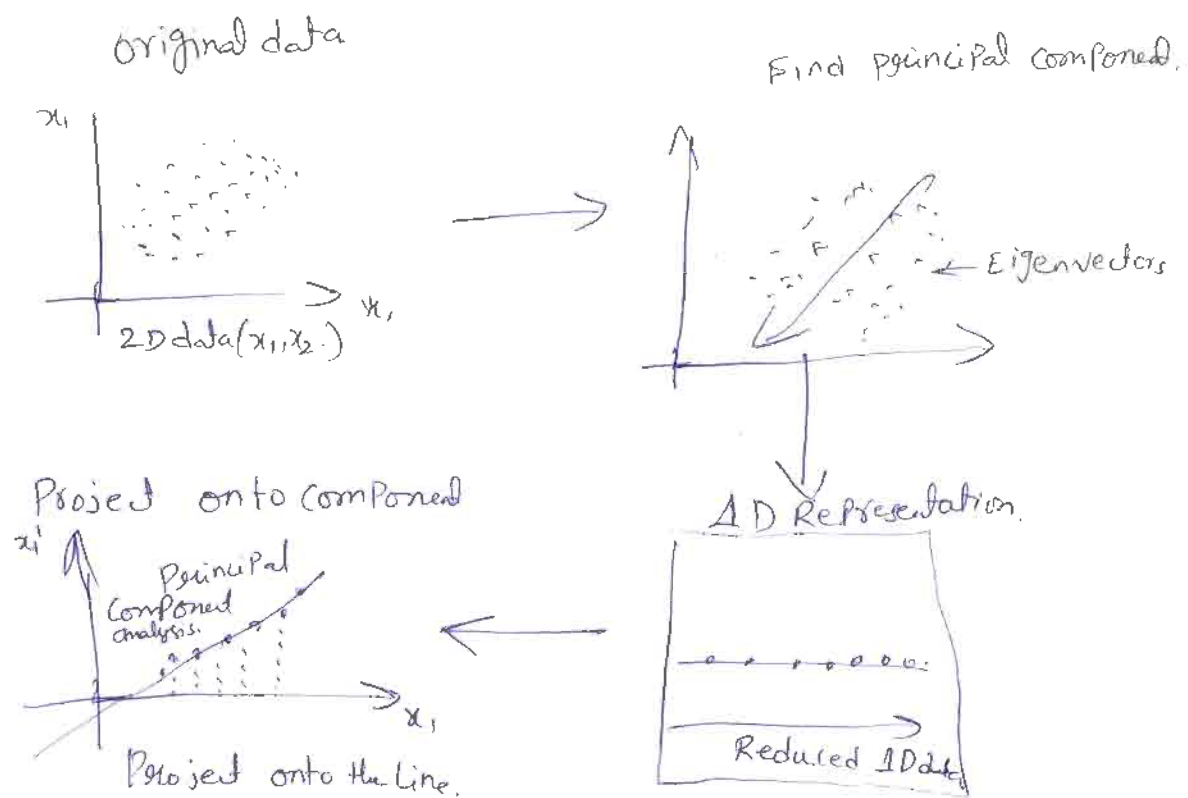


Fig: PCA: Reducing 2D Data to 1D

10) Compare and Contrast Reinforcement Learning with Supervised and unsupervised Learning,

Ans: Definition.

Learning Type	Definition
Supervised Learning (SL)	Learns a mapping from i/p to o/p using labeled data.
Unsupervised Learning (UL)	Finds patterns or structures in unlabeled data without explicit outputs.
Reinforcement Learning (RL)	Learns by interacting with an environment, taking actions to maximize cumulative reward.

Data Requirement :

Learning type

SL

UL

RL

Data Requirement

Requires labeled dataset
(features + target labels)

Works with unlabeled dataset
(only features)

Needs an environment and feedback
(reward / punishment) for actions,
not explicit labels.

Feed back type :

Learning type

SL

UL

RL

Feedback.

Immediate and exact

No explicit feedback

Delayed feedback.

Example Applications

Learning type

SL

UL

RL

Example

Image classification, regression,
Spam detection

Customer segmentation, anomaly
detection, PCA,

Game AI (Chess, Go), autonomous
driving.

11) Describe the importance of the Proposal distribution in MCMC (Markov chain monte carlo) algorithms,

Ans: In MCMC algorithms, the Proposal distribution plays a central role in determining how efficiently and accurately the target probability distribution is sampled.

→ The Proposal distribution, usually denoted as $q(x'/x)$, is the rule used to generate a candidate (Proposed) sample x' from the current state x .

→ MCMC algorithms such as Metropolis-Hastings use this Proposed sample and then decide whether to accept or reject it.

Affects Acceptance Probability.

→ In Metropolis-Hastings, the acceptance Probability depends on both the target distribution and the Proposal distribution.

$$\alpha = \min\left(1, \frac{\pi(x') q(x/x')}{\pi(x) q(x'/x)}\right)$$

A well chosen Proposal increases the chances of accepting useful moves, improving efficiency.