

Time : 0

Correct Marks : 2 Max. Selectable Options : 0

Question Label : Multiple Select Question

Which of the following JavaScript method(s) is/are used to remove data from the session storage?

Options :

6406532734047. ✔ clear()

6406532734048. ✔ removeItem()

6406532734049. ✖ deleteItem()

6406532734050. ✖ unSet()

MLT

Section Id :	64065356696
Section Number :	11
Section type :	Online
Mandatory or Optional :	Mandatory
Number of Questions :	17
Number of Questions to be attempted :	17
Section Marks :	50
Display Number Panel :	Yes
Section Negative Marks :	0
Group All Questions :	No
Enable Mark as Answered Mark for Review and Clear Response :	Yes
Maximum Instruction Time :	0
Sub-Section Number :	1
Sub-Section Id :	640653118943
Question Shuffling Allowed :	No

Is Section Default? : null

Question Number : 318 Question Id : 640653816191 Question Type : MCQ Is Question Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 0

Question Label : Multiple Choice Question

THIS IS QUESTION PAPER FOR THE SUBJECT "DIPLOMA LEVEL : MACHINE LEARNING TECHNIQUES (COMPUTER BASED EXAM)"

ARE YOU SURE YOU HAVE TO WRITE EXAM FOR THIS SUBJECT?
CROSS CHECK YOUR HALL TICKET TO CONFIRM THE SUBJECTS TO BE WRITTEN.

(IF IT IS NOT THE CORRECT SUBJECT, PLS CHECK THE SECTION AT THE [TOP](#) FOR THE SUBJECTS REGISTERED BY YOU)

Options :

6406532734051. ✓ YES

6406532734052. ✗ NO

Sub-Section Number : 2

Sub-Section Id : 640653118944

Question Shuffling Allowed : Yes

Is Section Default? : null

Question Number : 319 Question Id : 640653816194 Question Type : MCQ Is Question Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 3

Question Label : Multiple Choice Question

Imagine a dataset characterized by two features, Feature 1 and Feature 2, demonstrating a perfect negative correlation of -1. When applying k-means clustering with $k = 3$ to this dataset, what is the most likely arrangement of cluster centers that minimizes the within-cluster sum of squares

(WCSS)?

Options :

6406532734058. ✖ An equilateral triangle centered around the mean of the data.

6406532734059. ✔ Cluster centers positioned along a straight line.

6406532734060. ✖ A triangle with two acute angles, positioned strategically within the data distribution.

6406532734061. ✖ A right-angled triangle with one center at the origin.

Question Number : 320 Question Id : 640653816197 Question Type : MCQ Is Question Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 3

Question Label : Multiple Choice Question

Consider the following two models fitted on a one-dimensional dataset:

Model 1: $\hat{y} = w_0 + w_1x$

Model 2: $\hat{y} = w_1x^2 + w_2x + w_3$

If both models are trained on the same one-dimensional dataset and evaluated on the same test dataset, which model is more likely to have higher bias and lower variance?

Options :

6406532734067. ✔ Model 1

6406532734068. ✖ Model 2

6406532734069. ✖ Both models are equally sensitive to outliers

6406532734070. ✖ Insufficient data

Sub-Section Number :	3
Sub-Section Id :	640653118945
Question Shuffling Allowed :	Yes
Is Section Default? :	null

Question Number : 321 Question Id : 640653816201 Question Type : MCQ Is Question Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 4

Question Label : Multiple Choice Question

Consider a logistic regression model for a binary classification problem with two features x_1 and x_2 . The feature vector is $\begin{bmatrix} x_1 \\ x_2 \end{bmatrix}$ and labels lie in $\{0, 1\}$. The threshold for inference is 0.5. The dummy feature and the weight corresponding to it can be ignored for this problem. Let x_1 be the horizontal axis and x_2 be the vertical axis. You are given two feature vectors:

$$\mathbf{x}_1 = \begin{bmatrix} 1 \\ \sqrt{3} \end{bmatrix}, \mathbf{x}_2 = \begin{bmatrix} -1 \\ \sqrt{3} \end{bmatrix}$$

The weight vector makes an angle of θ with the positive x_1 axis (horizontal). Each θ corresponds to a different classifier. For what range of values of θ are both $\mathbf{x}_1$ and $\mathbf{x}_2$ predicted to belong to class-1?

Hints:

- To draw the weight vector $\mathbf{w} = \begin{bmatrix} w_1 \\ w_2 \end{bmatrix}$, plot the point (w_1, w_2) and draw an arrow starting at the origin to this point.
- $\tan(60^\circ) = \sqrt{3}$

Options :

6406532734074. ✓ $30^\circ < \theta < 150^\circ$

6406532734075. ✗ $0^\circ < \theta < 60^\circ$

6406532734076. ✗ $60^\circ < \theta < 180^\circ$

6406532734077. ✗ $0^\circ < \theta < 180^\circ$

6406532734078. ✗ $0 < \theta < 360^\circ$

Sub-Section Number : 4
Sub-Section Id : 640653118946
Question Shuffling Allowed : Yes
Is Section Default? : null

Question Number : 322 Question Id : 640653816193 Question Type : MSQ Is Question Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 3 Max. Selectable Options : 0

Question Label : Multiple Select Question

Which of the following statements accurately describes the characteristics of kernel functions?
Assume the dataset to be mean-centered.

Options :

6406532734054. ✖ Kernel PCA can reconstruct original PCA if the kernel function is $k(x_i, x_j) = (x_i^T x_j + 1)^2$.

6406532734055. ✖ The dimensionality of the transformed dataset $\phi(X)$, computed using the kernel function, is always smaller than the original feature space.

6406532734056. ✔ The dimensionality of the transformed dataset $\phi(X)$, computed using the kernel function, can exceed the original feature space.

6406532734057. ✔ The dimension of the transformed dataset $\phi(X)$, whose inner products the kernel function computes, can be infinite.

Question Number : 323 Question Id : 640653816196 Question Type : MSQ Is Question Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 3 Max. Selectable Options : 0

Question Label : Multiple Select Question

Let X be a data matrix of the shape (d, n) and y be the associated label vector of shape $(n, 1)$. Assume that a linear regression model with loss as the sum of squared error is trained on the data $\{X, y\}$. In which of the following cases, the loss on the training data will necessarily be zero? Assume that the solution of the model is obtained by the normal equation that is $w^* = (XX^T)^{-1}Xy$.

Options :

6406532734063. ✓ If y lies in the space spanned of row vectors of X .

6406532734064. ✗ If y lies in the space spanned of row vectors of X^T .

6406532734065. ✓ If all the data points satisfy the equality $x_1 + x_2 + \dots + x_d = 0$, where x_i is the i th feature and $y = 0$ for all the data points.

6406532734066. ✗ If all the data points satisfy the equality $x_1^3 + x_2^3 + \dots + x_d^3 = 0$, where x_i is the i th feature and $y = 0$ for all the data points.

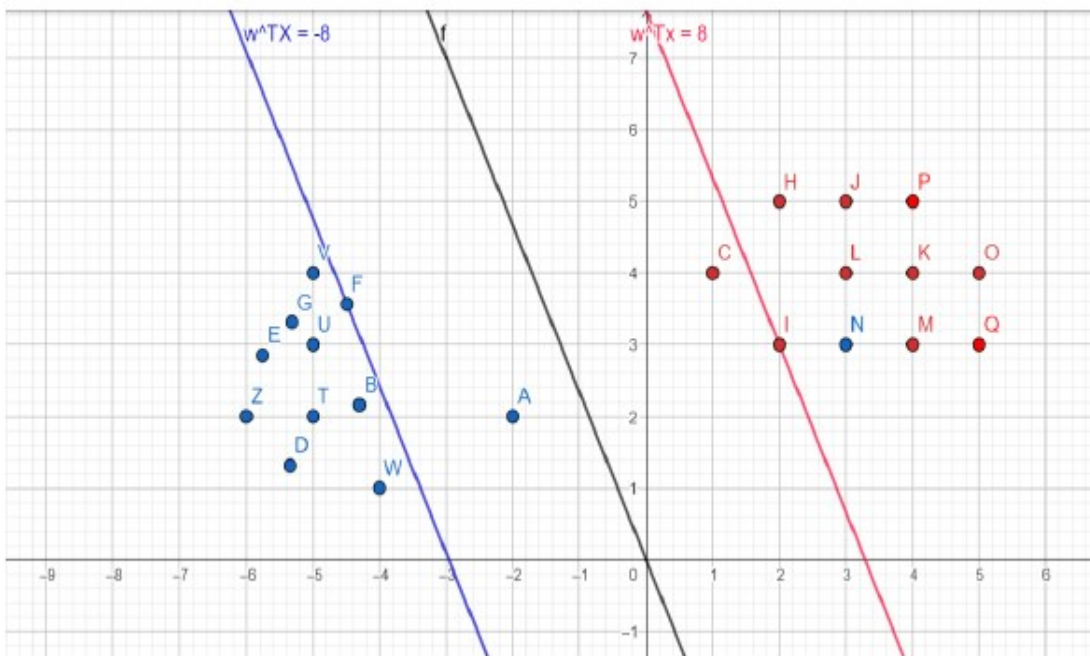
Question Number : 324 Question Id : 640653816203 Question Type : MSQ Is Question

Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 3 Max. Selectable Options : 0

Question Label : Multiple Select Question

Consider the following dataset on which the soft margin SVM is applied.



Which of the following statements is/are true about this dataset?

Options :

6406532734085. ✖ Points {F, A, C, I} are the only support vectors.

6406532734086. ✔ Points {A, N} are a subset of support vectors.

6406532734087. ✔ Points {F, A, N} are a subset of support vectors.

6406532734088. ✔ Points except {F, A, C, I, N} do not play any role in determining optimal weight vector.

6406532734089. ✖ Points except {F, A, C, I} do not play any role in determining optimal weight vector.

Question Number : 325 Question Id : 640653816205 Question Type : MSQ Is Question

Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 3 Max. Selectable Options : 0

Question Label : Multiple Select Question

Given a two-dimensional data set where points from class 1 are:

$\{(-2, 3), (-1, 1), (-1, 2), (-1, 4)\}$

And points from class 0 are:

$\{(1, 3), (1, 4), (2, 4), (2, 2)\}$

Which of the following statements are true?

Options :

6406532734091. ✓ The given data points from classes 1 and 0 can be linearly separated using a Hard-margin SVM.

6406532734092. ✓ A perceptron model and a hard margin SVM can give different decision boundary for this dataset.

6406532734093. ✗ A Soft-margin SVM would be a more robust choice than a Hard-margin SVM for this dataset as the dataset is not linearly separable.

6406532734094. ✗ The width of the separation between the two supporting hyperplanes is 4.
(Hint: Calculate width using formulae $\frac{2}{\|w\|}$)

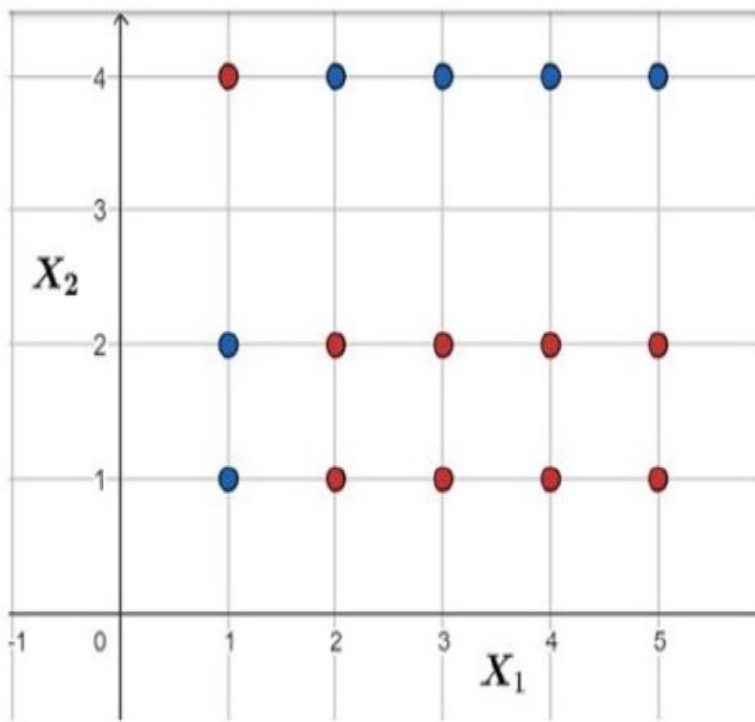
Question Number : 326 Question Id : 640653816206 Question Type : MSQ Is Question

Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 3 Max. Selectable Options : 0

Question Label : Multiple Select Question

Consider the following two-dimensional dataset with two classes: +1 for blue points and -1 for red points. An AdaBoost algorithm was run on this dataset using decision stumps as weak learners.



When training the new weak learner $h_t(x)$ (decision stump at t^{th} iteration), we choose the split that minimizes the weighted miss-classification error with respect to current weights D_t i.e. choose h_t that minimizes $\sum_{i=1}^n D_t(i) \mathbb{1}(h_t(x_i) \neq y_i)$.
Based on the above data, answer the below given question.

To train the second decision stump, which pair of points will be assigned equal weights to create the training data-set?

Options :

6406532734095. ✖ $[2, 2]^T, [1, 2]^T$

6406532734096. ✖ $[2, 2]^T, [1, 4]^T$

6406532734097. ✔ $[1, 1]^T, [1, 4]^T$

6406532734098. ✔ $[3, 1]^T, [3, 4]^T$

Sub-Section Number :

5

Sub-Section Id :

640653118947

Question Shuffling Allowed :

Yes

Is Section Default? :

null

Question Number : 327 Question Id : 640653816202 Question Type : MSQ Is Question

Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 4 Max. Selectable Options : 0

Question Label : Multiple Select Question

Consider a soft-margin Support Vector Machine (SVM) for a binary classification problem with a dataset in a two-dimensional space. The optimization problem for the soft-margin SVM is formulated as:

$$\text{Minimize } \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^N \xi_i$$

subject to the constraints:

$$y_i(\mathbf{w} \cdot \mathbf{x}_i + b) \geq 1 - \xi_i \text{ and } \xi_i \geq 0 \text{ for all } i$$

Where C is a positive constant.

Let w^* , ξ^* be the optimal solutions, and α^* , β^* be the optimal dual solutions of the soft margin SVM problem.

Which of the following statements about the soft-margin SVM is correct?

Options :

6406532734079. ✖ If i^{th} data point lies on one of the supporting hyperplanes, then $\alpha_i^* = 0$.

6406532734080. ✔ If i^{th} data point lies on the correct supporting hyperplane, it does **not** pay any bribes.

6406532734081. ✖ A smaller value of C allows for a larger margin, potentially leading to less misclassifications on the training data.

6406532734082. ✔ For a dataset with n data-points, there are $2n$ constraints for soft-margin SVM.

6406532734083. ✔ As C approaches ∞ the soft margin SVM is equal to the hard margin SVM.

6406532734084. ✖ C can be negative, as long as the bribe(ξ) each data point pays is non-negative.

Sub-Section Number : 6
Sub-Section Id : 640653118948
Question Shuffling Allowed : Yes
Is Section Default? : null

Question Number : 328 Question Id : 640653816192 Question Type : SA Calculator : None
Response Time : N.A Think Time : N.A Minimum Instruction Time : 0
Correct Marks : 3

Question Label : Short Answer Question

Consider a dataset X in $\mathbb{R}^3$. The dataset X consists of 4 samples with 3 features each. The covariance matrix C of this dataset has three non-zero eigenvalues which follow the given linear equations:

$$2\lambda_1 + 3\lambda_2 - \lambda_3 = 5$$

$$\lambda_1 - 2\lambda_2 + 4\lambda_3 = 8$$

$$3\lambda_1 + \lambda_2 - 2\lambda_3 = 3$$

Determine the variance of the given dataset.

Response Type : Numeric

Evaluation Required For SA : Yes

Show Word Count : Yes

Answers Type : Equal

Text Areas : PlainText

Possible Answers :

5

Question Number : 329 Question Id : 640653816195 Question Type : SA Calculator : None
Response Time : N.A Think Time : N.A Minimum Instruction Time : 0
Correct Marks : 3

Question Label : Short Answer Question

Consider a dataset of n observations $\{x_1, x_2, \dots, x_n\}$, where each x_i follows a Bernoulli distribution with parameter p , i.e., $x_i \sim \text{Bernoulli}(p)$ for $i = 1, 2, \dots, n$. However, you have reason to believe that the parameter p might differ for two distinct groups within the dataset. You suspect that there are two groups in the dataset, each with its own parameter (p_1 and p_2). Now, develop an algorithm to estimate the parameters p_1 and p_2 using maximum likelihood estimation. Then, apply your algorithm to a dataset with the following observations and corresponding group labels: $\{0, 1, 1, 0, 1\}$ and $\{1, 0, 1, 0, 1\}$ for group 1 and group 2 respectively.

Calculate the maximum likelihood estimates of p_1 and rounded to two decimal places.

Response Type : Numeric

Evaluation Required For SA : Yes

Show Word Count : Yes

Answers Type : Equal

Text Areas : PlainText

Possible Answers :

0.6

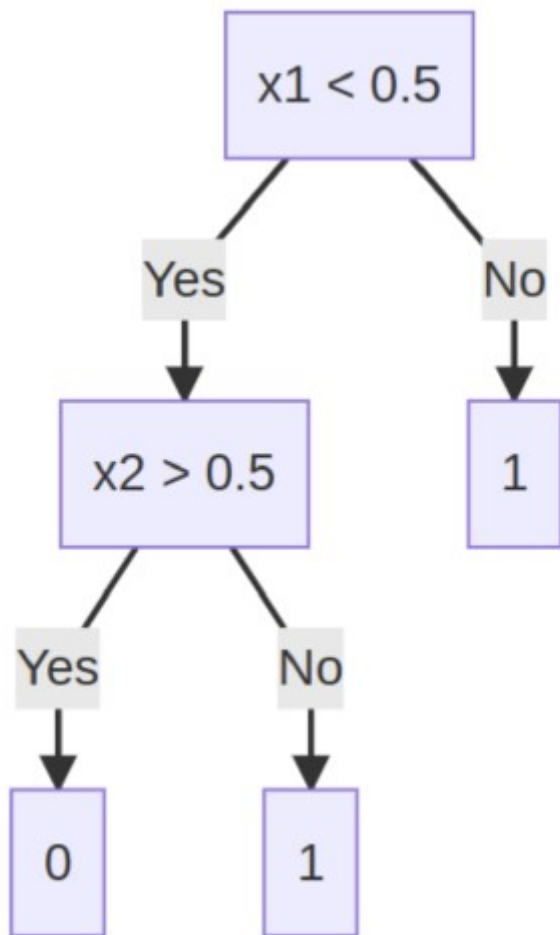
Question Number : 330 **Question Id :** 640653816198 **Question Type :** SA **Calculator :** None

Response Time : N.A **Think Time :** N.A **Minimum Instruction Time :** 0

Correct Marks : 3

Question Label : Short Answer Question

Consider the following decision tree for a classification problem in which all the data-points are constrained to lie in the unit square in the first quadrant. That is $0 \leq x_1, x_2 \leq 1$. If a point is picked uniformly at random from the unit square, what is the probability that the decision tree predicts this point as belonging to class 1?



Response Type : Numeric

Evaluation Required For SA : Yes

Show Word Count : Yes

Answers Type : Equal

Text Areas : PlainText

Possible Answers :

0.75

Question Number : 331 **Question Id :** 640653816199 **Question Type :** SA **Calculator :** None

Response Time : N.A **Think Time :** N.A **Minimum Instruction Time :** 0

Correct Marks : 3

Question Label : Short Answer Question

Consider the following dataset with 6 samples along with the corresponding labels. Each sample has three binary features f_1, f_2 and f_3 .

sample	f_1	f_2	f_3	y
x_1	1	1	0	1
x_2	0	1	0	1
x_3	1	1	1	0
x_4	0	1	1	0
x_5	1	0	1	0
x_6	1	1	1	1

Assume that the features are conditionally independent given the label y . Suppose the test sample is $x_{test} = [0, 1, 0]^T$.

What is the estimated probability that the test point belongs to class 0 (that is, $p(y = 0|x_{test})$)?

Response Type : Numeric

Evaluation Required For SA : Yes

Show Word Count : Yes

Answers Type : Equal

Text Areas : PlainText

Possible Answers :

0

Question Number : 332 **Question Id :** 640653816200 **Question Type :** SA **Calculator :** None

Response Time : N.A **Think Time :** N.A **Minimum Instruction Time :** 0

Correct Marks : 3

Question Label : Short Answer Question

Consider a linearly separable binary classification data set with 1000 data points and 100 features. Assume that there exists a w such that $\|w\| = 1, y_i(w^T x_i) \geq 0.5 \forall i$. Also assume that $\|x\|_2 \leq 2 \forall i$ What is the maximum number of mistakes that the Perceptron algorithm can make in this data set?

Response Type : Numeric

Evaluation Required For SA : Yes

Show Word Count : Yes

Answers Type : Equal

Text Areas : PlainText

Possible Answers :

16

Question Number : 333 **Question Id :** 640653816204 **Question Type :** SA **Calculator :** None

Response Time : N.A **Think Time :** N.A **Minimum Instruction Time :** 0

Correct Marks : 3

Question Label : Short Answer Question

Consider a single iteration of the AdaBoost algorithm that was run on three sample points, starting with uniform weights on the sample points. The labels are either +1 or -1 In the table below, some values have been omitted.

Data point	True label	Predicted label	Initial weight	Updated weight
x_1	?	1	$\frac{1}{3}$	$\frac{1}{2}$
x_2	-1	-1	$\frac{1}{3}$	?
x_3	-1	?	$\frac{1}{3}$	$\frac{1}{4}$

Based on the above data, what will be the updated weight for point x_2 ?

Response Type : Numeric

Evaluation Required For SA : Yes

Show Word Count : Yes

Answers Type : Equal

Text Areas : PlainText

Possible Answers :

0.25

Question Number : 334 **Question Id :** 640653816207 **Question Type :** SA **Calculator :** None

Response Time : N.A **Think Time :** N.A **Minimum Instruction Time :** 0

Correct Marks : 3

Question Label : Short Answer Question

Consider a simple neural network with one hidden layer. The network has the following architecture:

Input layer with 3 neurons. Hidden layer with 2 neurons, using the sigmoid activation function.

Output layer with 1 neuron, using the linear activation function.

The weights and biases for the network are as follows:

Hidden Layer:

Neuron 1: Weights: [0.5, -0.2, 0.8]

Bias: 0.1

Neuron 2: Weights: [0.4, 0, 0.2]

Bias: -0.4

Output Layer:

Neuron 1: Weights: [0.2, 0.4]

Bias: -0.3

Assume that the input values are [0.6, 0.3, 0.8].

Calculate output of Neuron 1 in hidden layer

Response Type : Numeric

Evaluation Required For SA : Yes

Show Word Count : Yes

Answers Type : Range

Text Areas : PlainText

Possible Answers :

0.70 to 0.80

BDM

Section Id :	64065356697
Section Number :	12
Section type :	Online
Mandatory or Optional :	Mandatory
Number of Questions :	20
Number of Questions to be attempted :	20
Section Marks :	30
Display Number Panel :	Yes
Section Negative Marks :	0
Group All Questions :	No
Enable Mark as Answered Mark for Review and Clear Response :	Yes
Maximum Instruction Time :	0
Sub-Section Number :	1
Sub-Section Id :	640653118949
Question Shuffling Allowed :	No
Is Section Default? :	null

Question Number : 335 Question Id : 640653816208 Question Type : MCQ Is Question Mandatory : No Calculator : None Response Time : N.A Think Time : N.A Minimum Instruction Time : 0

Correct Marks : 0

Question Label : Multiple Choice Question

THIS IS QUESTION PAPER FOR THE SUBJECT **"DIPLOMA LEVEL : BUSINESS DATA MANAGEMENT (COMPUTER BASED EXAM)"**

ARE YOU SURE YOU HAVE TO WRITE EXAM FOR THIS SUBJECT?

CROSS CHECK YOUR HALL TICKET TO CONFIRM THE SUBJECTS TO BE WRITTEN.