

GATE Data Science and Artificial Intelligence

Sample Paper – 9

Duration: 128 Minutes

Maximum Marks: 72

Instructions

- This paper contains **46 questions** covering the core **Data Science and Artificial Intelligence** subject of the GATE paper conducted by the IITs and IISc (General Aptitude and Engineering Mathematics are provided as separate papers).
- **Q1–Q20 carry 1 mark each and Q21–Q46 carry 2 marks each**, for a total of **72 marks**.
- Each question carries **negative marking of one-third of its allotted marks**: a wrong 1-mark question costs **0.33** and a wrong 2-mark question costs **0.67**. Unattempted questions carry no penalty.
- Every question here is a single-correct **MCQ**. The real GATE paper also has Multiple Select (MSQ) and Numerical Answer Type (NAT) questions, and **those carry no negative marking**.
- Attempt this paper in one timed sitting of about **128 minutes**, the share this core subject earns at GATE's overall pace of 180 minutes for 65 questions. All logarithms in information-theoretic quantities (entropy, information gain) are to base 2 and expressed in bits unless stated otherwise.

Probability & Statistics

Q1. From a group of 4 statisticians and 6 programmers, a review committee of 4 members is to be formed that contains *exactly 2 statisticians*. The number of distinct committees possible is: [1 mark]

- (A) 210
- (B) 15
- (C) 120

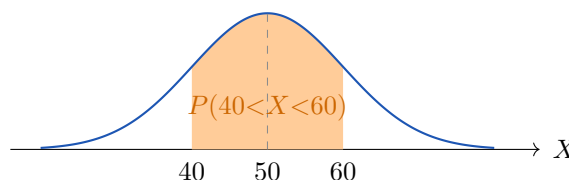


(D) 90

- Q2.** The contingency table below classifies 200 e-mails by whether they are spam and whether they contain a flagged keyword. Given that an e-mail contains the keyword, the probability that it is spam, $P(\text{Spam} \mid \text{Word})$, is: [1 mark]

	Word	No Word
Spam	40	20
Ham	10	130

- (A) 0.667
(B) 0.800
(C) 0.200
(D) 0.500
- Q3.** A screening test for a rare condition has sensitivity 0.99 (probability of a positive result given the condition) and a false-positive rate of 0.05. The condition affects 1% of the population. If a randomly chosen person tests positive, the probability that they actually have the condition is nearest to: [1 mark]
- (A) 0.990
(B) 0.950
(C) 0.500
(D) 0.167
- Q4.** A feature X is normally distributed with mean $\mu = 50$ and standard deviation $\sigma = 10$. Using the shaded central region of the standard normal curve shown, the probability $P(40 < X < 60)$ is nearest to: [1 mark]



- (A) 0.3413
- (B) 0.6826
- (C) 0.9544
- (D) 0.5000

Q5. Flaws occur along a roll of fabric according to a Poisson process at an average rate of 2 flaws per metre. The probability of finding *at least one* flaw in a given one-metre length is nearest to (take $e^{-2} = 0.135$):[1 mark]

- (A) 0.865
- (B) 0.135
- (C) 0.271
- (D) 0.500

Q6. A discrete random variable X takes the values 1, 2, 3 with probabilities 0.5, 0.3, 0.2 respectively. The expected value $E[X]$ is: [1 mark]

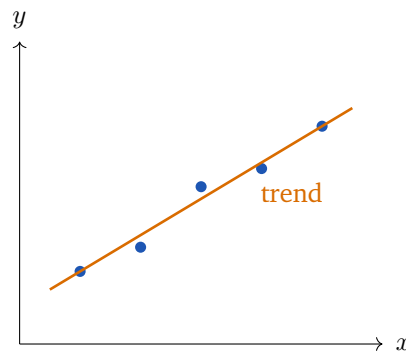
- (A) 1.5
- (B) 1.7
- (C) 2.0
- (D) 0.61

Q7. A sample of size $n = 64$ drawn from a population with known standard deviation $\sigma = 8$ has sample mean $\bar{x} = 52$. To test the null hypothesis $\mu_0 = 50$, the value of the standardized test statistic $z = \frac{\bar{x} - \mu_0}{\sigma/\sqrt{n}}$ is: [1 mark]

- (A) 0.25
- (B) 16
- (C) 2
- (D) 1



- Q8.** For a bivariate data set the sums of squares and cross-products about the means are $S_{xx} = 25$, $S_{yy} = 100$ and $S_{xy} = 30$, and the scatter is shown below. The Pearson correlation coefficient r between x and y is: [1 mark]



- (A) 1.200
(B) 0.300
(C) 0.600
(D) 0.360
- Q9.** A fair coin is tossed 3 times independently. The probability of obtaining *exactly 2 heads* is: [1 mark]
- (A) 0.125
(B) 0.375
(C) 0.500
(D) 0.250
- Q10.** The lifetime of a sensor follows an exponential distribution with a mean of 4 years. The probability that a sensor fails *within* the first 4 years, $P(X < 4)$, is nearest to (take $e^{-1} = 0.368$): [1 mark]
- (A) 0.368
(B) 0.250
(C) 0.135
(D) 0.632



Q11. Consider the matrix $A = \begin{bmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 1 & 0 & 1 \end{bmatrix}$. The rank of A is: [1 mark]

- (A) 1
- (B) 2
- (C) 3
- (D) 0

Q12. The determinant of the matrix $B = \begin{bmatrix} 1 & 2 & 0 \\ 0 & 3 & 1 \\ 2 & 0 & 4 \end{bmatrix}$ is: [1 mark]

- (A) 16
- (B) 10
- (C) 12
- (D) 0

Q13. The *smaller* eigenvalue of the matrix $C = \begin{bmatrix} 2 & 2 \\ 1 & 3 \end{bmatrix}$ is: [1 mark]

- (A) 4
- (B) 5
- (C) 2
- (D) 1

Q14. For a real matrix A , the matrix $A^T A$ has eigenvalues 16 and 4. The condition number of A in the spectral norm, defined as $\kappa = \sigma_{\max}/\sigma_{\min}$, is: [1 mark]

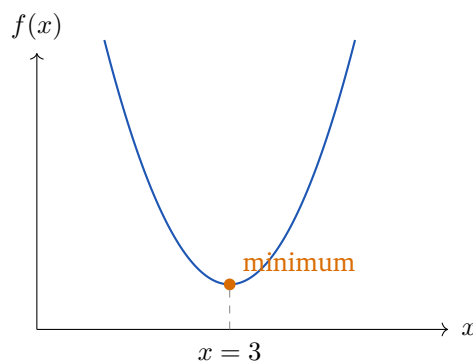
- (A) 4
- (B) 2
- (C) 8
- (D) 16



Q15. For the vector $v = (2, -3, 6)$, the Euclidean (ℓ_2) norm $\|v\|_2$ is: [1 mark]

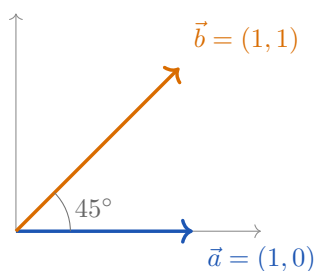
- (A) 7
- (B) 11
- (C) 49
- (D) 6

Q16. The function $f(x) = x^2 - 6x + 11$ is minimized over the real line, as suggested by the convex graph below. The value of x at which the minimum occurs is: [1 mark]



- (A) 6
- (B) 3
- (C) 11
- (D) 2

Q17. The cosine similarity between the vectors $\vec{a} = (1, 0)$ and $\vec{b} = (1, 1)$, shown below, defined as $\frac{\vec{a} \cdot \vec{b}}{\|\vec{a}\| \|\vec{b}\|}$, is nearest to: [1 mark]



- (A) 1.000



- (B) 0.500
- (C) 0.707
- (D) 0.000

Q18. Gradient descent is used to minimize $f(x) = x^2$ with a fixed learning rate $\eta = 0.25$. Starting from $x_0 = 8$, the value of x_2 after two updates $x_{t+1} = x_t - \eta f'(x_t)$ is: [1 mark]

- (A) 2
- (B) 4
- (C) 6
- (D) 0

Q19. Consider the function $f(x, y) = x^2 - y^2$ on $\mathbb{R}^2$. Based on its Hessian, the function is: [1 mark]

- (A) convex everywhere
- (B) concave everywhere
- (C) linear (affine)
- (D) neither convex nor concave

Programming, Data Structures & Algorithms

Q20. What is the output of the following Python snippet?

```
s = "data"
print(s[::-1])
```

[1 mark]

- (A) data
- (B) atad
- (C) dat
- (D) ata

Q21. What value is printed by the following Python code?

```
print(len([x for x in range(10) if x % 3 == 0]))
```

[2 marks]

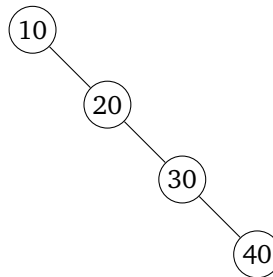


- (A) 3
- (B) 4
- (C) 5
- (D) 10

Q22. The following operations are performed on an initially empty *queue* (FIFO), in order: enqueue(1), enqueue(2), dequeue(), enqueue(3), enqueue(4), dequeue(), dequeue(). The value returned by the *last* dequeue() is: [2 marks]

- (A) 3
- (B) 1
- (C) 2
- (D) 4

Q23. The keys 10, 20, 30, 40 are inserted, in that order, into an initially empty binary search tree, producing the tree shown. The *height* of the resulting tree (measured as the number of edges on the longest root-to-leaf path) is: [2 marks]



- (A) 1
- (B) 2
- (C) 3
- (D) 4

Q24. A hash table of size 10 uses the hash function $h(k) = k \bmod 10$ with *linear probing* for collision resolution. The keys 12, 22, 32 are inserted in this order into an otherwise empty table. The final index occupied by key 32 is: [2 marks]



- (A) 4
- (B) 2
- (C) 3
- (D) 32

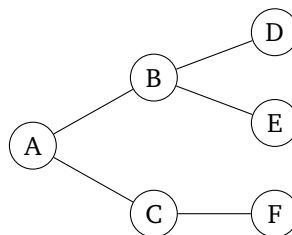
Q25. Binary search is performed on a sorted array of 63 distinct elements. The maximum number of key comparisons required in the worst case is: [2 marks]

- (A) 63
- (B) 32
- (C) 7
- (D) 6

Q26. The worst-case time complexity of Merge Sort on an array of n elements is: [2 marks]

- (A) $O(n^2)$
- (B) $O(n)$
- (C) $O(\log n)$
- (D) $O(n \log n)$

Q27. A depth-first search (DFS) is started from vertex A in the undirected graph shown, always visiting adjacent vertices in increasing alphabetical order and backtracking when no unvisited neighbour remains. The order in which the vertices are first visited is: [2 marks]



- (A) A, B, D, E, C, F
- (B) A, B, C, D, E, F



(C) A, C, F, B, D, E

(D) A, B, C, E, D, F

Q28. The running time of an algorithm satisfies the recurrence $T(n) = 2T(n/2) + n^2$, with $T(1) = 1$. By the Master theorem, $T(n)$ is: [2 marks]

(A) $\Theta(n^2)$

(B) $\Theta(n \log n)$

(C) $\Theta(n)$

(D) $\Theta(n^3)$

Database Management & Warehousing

Q29. Two union-compatible relations R and S have 12 and 8 tuples respectively, and exactly 3 tuples appear in both. The number of tuples in $R \cup S$ is: [2 marks]

(A) 20

(B) 3

(C) 17

(D) 5

Q30. A column grade of a table contains the values 5, 5, 10, 10, 10, NULL, 20 (seven rows). What does the query `SELECT COUNT(DISTINCT grade) FROM T;` return? [2 marks]

(A) 7

(B) 3

(C) 4

(D) 6

Q31. A relation $R(A, B, C)$ has the functional dependencies $AB \rightarrow C$ and $C \rightarrow B$. The candidate keys are AB and AC , so A , B and C are all prime attributes. The highest normal form that R satisfies is: [2 marks]



- (A) BCNF
- (B) 2NF
- (C) 3NF
- (D) 1NF

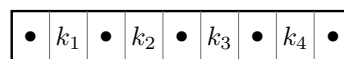
Q32. Relation $R(A, B)$ contains the tuples $\{(1, x), (2, y), (3, x)\}$ and relation $S(B, C)$ contains $\{(x, p), (x, q), (z, r)\}$. The number of tuples in the natural join $R \bowtie S$ (on the common attribute B) is: [2 marks]

- (A) 4
- (B) 9
- (C) 2
- (D) 6

Q33. In data warehousing, the schema in which a single central fact table is directly connected to a set of *denormalized* dimension tables, with the diagram resembling a central hub with radiating points, is called a: [2 marks]

- (A) snowflake schema
- (B) fact constellation schema
- (C) star schema
- (D) fully normalized schema

Q34. The figure shows an internal node of a B^+ tree of order $m = 5$ (an internal node holds up to m child pointers, drawn as $\bullet$). For a *non-root* internal node, the *minimum* number of child pointers it must contain, equal to $\lceil m/2 \rceil$, is: [2 marks]



order $m = 5$: up to 5 pointers ($\bullet$)

- (A) 5



- (B) 3
- (C) 2
- (D) 4

Q35. A relation has attributes $\{A, B, C, D, E\}$ with functional dependencies $A \rightarrow B$, $B \rightarrow C$ and $D \rightarrow E$. The attribute closure A^+ is: [2 marks]

- (A) $\{A, B\}$
- (B) $\{A, B, C, D, E\}$
- (C) $\{A, B, C\}$
- (D) $\{A\}$

Machine Learning & Artificial Intelligence

Q36. A logistic regression model has weight vector $w = (0.5, 0.5)$ and bias $b = 0$, and uses the sigmoid $\sigma(z) = 1/(1 + e^{-z})$. For the input $x = (1, 1)$, the predicted probability $\sigma(w^T x + b)$ is nearest to (take $e^{-1} = 0.368$): [2 marks]

- (A) 0.50
- (B) 0.27
- (C) 0.88
- (D) 0.73

Q37. A model is too simple to capture the underlying pattern in the data and therefore gives poor accuracy on *both* the training set and the test set. In terms of the bias–variance decomposition, this behaviour is best described as: [2 marks]

- (A) low bias, high variance
- (B) low bias, low variance
- (C) high bias, high variance
- (D) high bias, low variance



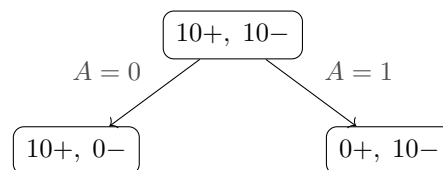
Q38. A k -nearest-neighbour classifier with $k = 3$ (Euclidean distance) is trained on the labelled points $A(1, 1)$: Class 0, $B(2, 1)$: Class 0, $C(5, 5)$: Class 1, $D(6, 5)$: Class 1, $E(6, 6)$: Class 1. The query point $Q(5, 6)$ is classified as:
[2 marks]

- (A) Class 1
- (B) Class 0
- (C) a tie (undecided)
- (D) cannot be determined

Q39. A leaf of a decision tree contains 8 training examples, all belonging to the *same* class. The entropy (in bits) of this node is: [2 marks]

- (A) 1.000
- (B) 0.500
- (C) 0.811
- (D) 0.000

Q40. A node in a decision tree contains 10 positive and 10 negative examples. A candidate attribute A splits it into two *pure* children as shown: $A = 0$ gives (10+, 0−) and $A = 1$ gives (0+, 10−). The information gain of this split (in bits) is: [2 marks]

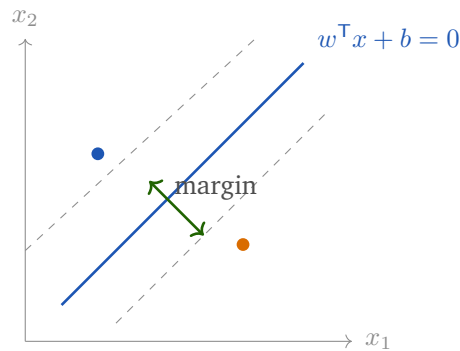


- (A) 0.500
- (B) 1.000
- (C) 0.000
- (D) 0.811

Q41. For the linear support vector machine shown, the separating hyperplane $w^T x + b = 0$ has two margin boundaries whose separation (the margin

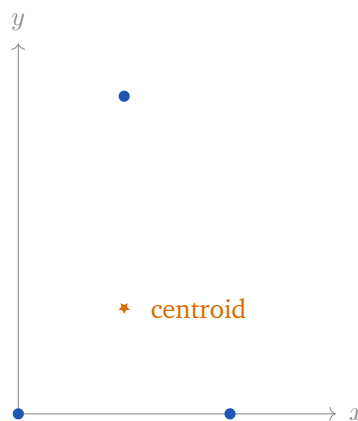


width) equals $2/\|w\|$. If the margin width is measured to be 0.5, the value of $\|w\|$ is: [2 marks]



- (A) 4
- (B) 0.5
- (C) 2
- (D) 1

Q42. In one iteration of the k -means algorithm, a cluster is assigned the three points $(0, 0)$, $(4, 0)$ and $(2, 6)$, shown below. The updated centroid of this cluster is: [2 marks]



- (A) $(3, 3)$
- (B) $(2, 3)$
- (C) $(2, 2)$
- (D) $(6, 6)$

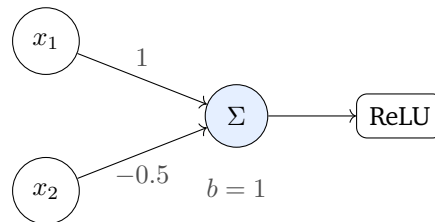
Q43. Principal component analysis is applied to a data set whose covariance matrix has eigenvalues 5, 3, 1, 1. The fraction of the total variance ex-



plained by the *first* principal component alone is:
[2 marks]

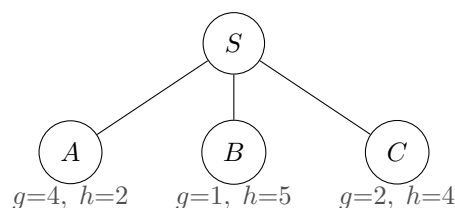
- (A) 80%
- (B) 50%
- (C) 30%
- (D) 90%

Q44. A single neuron with a ReLU activation computes $\text{ReLU}(w_1x_1 + w_2x_2 + b)$, as shown, where $\text{ReLU}(z) = \max(0, z)$. With inputs $x_1 = 2$, $x_2 = 1$, weights $w_1 = 1$, $w_2 = -0.5$ and bias $b = 1$, the output of the neuron is: [2 marks]



- (A) 0
- (B) -2.5
- (C) 2.5
- (D) 1.5

Q45. Greedy best-first search expands the frontier node with the *smallest heuristic value* h , ignoring the path cost g . The tree below shows three frontier nodes with their g and h values. The node greedy best-first search selects for expansion next is: [2 marks]

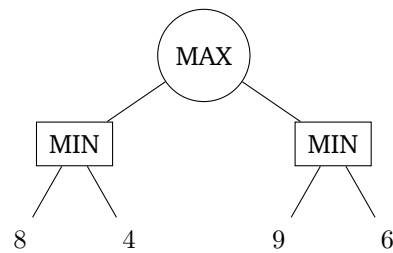


- (A) A
- (B) B



- (C) C
- (D) all three at once

Q46. In the two-ply game tree shown, the root is a MAX node and the second level are MIN nodes; the leaf values are given. The minimax value backed up to the root is: [2 marks]



- (A) 4
- (B) 8
- (C) 9
- (D) 6



Detailed Solutions

Q1.

Solution

Concept – choose within each group, then multiply: An “exactly 2 statisticians” committee fixes the split as 2 statisticians and 2 programmers.

Step 1 – choose 2 statisticians from 4: $\binom{4}{2} = \frac{4 \times 3}{2 \times 1}$

$$= \frac{12}{2}$$

$$= 6$$

Step 2 – choose 2 programmers from 6: $\binom{6}{2} = \frac{6 \times 5}{2 \times 1}$

$$= \frac{30}{2}$$

$$= 15$$

Step 3 – multiply the independent choices: $6 \times 15 = 90$

Final Answer: 90 committees $\Rightarrow$ D

Answer: (D) [Go Back to Q1](#)

Q2.

Solution

Concept – conditional probability from a table: $P(\text{Spam} \mid \text{Word}) = \frac{\text{Spam and Word}}{\text{all Word}}$.

Step 1 – count e-mails with the keyword: Word column total = $40 + 10 = 50$.

Step 2 – count spam among those: Spam and Word = 40.

Step 3 – form the ratio: $P(\text{Spam} \mid \text{Word}) = \frac{40}{50}$

Step 4 – simplify: $P(\text{Spam} \mid \text{Word}) = 0.80$

Final Answer: 0.80 $\Rightarrow$ B

Answer: (B) [Go Back to Q2](#)



Q3.

Solution

Concept – Bayes’ theorem: $P(\text{Cond} \mid +) = \frac{P(+ \mid \text{Cond}) P(\text{Cond})}{P(+)}$.

Step 1 – true-positive contribution: $P(+ \mid \text{Cond}) P(\text{Cond}) = 0.99 \times 0.01 = 0.0099$

Step 2 – false-positive contribution: $P(+ \mid \text{no Cond}) P(\text{no Cond}) = 0.05 \times 0.99 = 0.0495$

Step 3 – total probability of a positive test: $P(+) = 0.0099 + 0.0495 = 0.0594$

Step 4 – apply Bayes’ theorem: $P(\text{Cond} \mid +) = \frac{0.0099}{0.0594}$

Step 5 – divide: $P(\text{Cond} \mid +) = 0.167$

Final Answer: $0.167 \Rightarrow \boxed{\text{D}}$

Answer: (D) [Go Back to Q3](#)

Q4.

Solution

Concept – standardize both limits, then use symmetry: Convert X to $Z = \frac{X - \mu}{\sigma}$.

Step 1 – standardize the lower limit $X = 40$: $Z = \frac{40 - 50}{10} = \frac{-10}{10} = -1$

Step 2 – standardize the upper limit $X = 60$: $Z = \frac{60 - 50}{10} = \frac{10}{10} = 1$

Step 3 – write the required probability in Z : $P(40 < X < 60) = P(-1 < Z < 1)$

Step 4 – use the standard normal value $P(Z \leq 1) = 0.8413$: $P(-1 < Z < 1) = 0.8413 - (1 - 0.8413)$

$= 0.8413 - 0.1587$

$= 0.6826$

Final Answer: $P(40 < X < 60) = 0.6826 \Rightarrow \boxed{\text{B}}$

Answer: (B) [Go Back to Q4](#)



Q5.

Solution

Concept – “at least one” via the complement: $P(X \geq 1) = 1 - P(X = 0)$ for a Poisson variable.

Step 1 – list the parameter: $\lambda = 2$ per metre.

Step 2 – probability of zero flaws: $P(X = 0) = \frac{e^{-2} 2^0}{0!} = e^{-2}$

Step 3 – substitute $e^{-2} = 0.135$: $P(X = 0) = 0.135$

Step 4 – take the complement: $P(X \geq 1) = 1 - 0.135 = 0.865$

Final Answer: $P(X \geq 1) \approx 0.865 \Rightarrow \boxed{\text{A}}$

Answer: (A) [Go Back to Q5](#)

Q6.

Solution

Concept – expected value of a discrete variable: $E[X] = \sum_i x_i P(x_i)$.

Step 1 – multiply each value by its probability: $1 \times 0.5 = 0.5$

$$2 \times 0.3 = 0.6$$

$$3 \times 0.2 = 0.6$$

Step 2 – add the products: $E[X] = 0.5 + 0.6 + 0.6$

Step 3 – finish the arithmetic: $E[X] = 1.7$

Final Answer: $E[X] = 1.7 \Rightarrow \boxed{\text{B}}$

Answer: (B) [Go Back to Q6](#)

Q7.

Solution

Concept – one-sample z statistic: $z = \frac{\bar{x} - \mu_0}{\sigma/\sqrt{n}}$.

Step 1 – compute the standard error: $\frac{\sigma}{\sqrt{n}} = \frac{8}{\sqrt{64}}$

$$= \frac{8}{8}$$

$$= 1$$

Step 2 – compute the numerator: $\bar{x} - \mu_0 = 52 - 50 = 2$

Step 3 – divide: $z = \frac{2}{1} = 2$



Final Answer: $z = 2 \Rightarrow$ C

Answer: (C) [Go Back to Q7](#)

Q8.

Solution

Concept – Pearson correlation: $r = \frac{S_{xy}}{\sqrt{S_{xx} S_{yy}}}$.

Step 1 – compute the product $S_{xx} S_{yy}$: $S_{xx} S_{yy} = 25 \times 100 = 2500$

Step 2 – take the square root: $\sqrt{2500} = 50$

Step 3 – form the ratio: $r = \frac{30}{50}$

Step 4 – simplify: $r = 0.60$

Why option (A) is wrong: $1.20 = S_{xy}/S_{xx}$ is the regression *slope*, not the correlation; a correlation can never exceed 1.

Final Answer: $r = 0.60 \Rightarrow$ C

Answer: (C) [Go Back to Q8](#)

Q9.

Solution

Concept – binomial probability: $P(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}$.

Step 1 – list the parameters: $n = 3, k = 2, p = 0.5, 1 - p = 0.5$.

Step 2 – evaluate the binomial coefficient: $\binom{3}{2} = 3$

Step 3 – evaluate the powers: $p^2 = 0.5^2 = 0.25$ and $(1 - p)^1 = 0.5$

Step 4 – multiply the pieces: $P(X = 2) = 3 \times 0.25 \times 0.5$

$$= 3 \times 0.125$$

$$= 0.375$$

Final Answer: $P(X = 2) = 0.375 \Rightarrow$ B

Answer: (B) [Go Back to Q9](#)



Q10.

Solution

Concept – exponential CDF: For rate λ , $P(X < t) = 1 - e^{-\lambda t}$.

Step 1 – find the rate from the mean: Mean = $\frac{1}{\lambda} = 4$ years, so $\lambda = \frac{1}{4} = 0.25 \text{ yr}^{-1}$.

Step 2 – form the exponent: $\lambda t = 0.25 \times 4 = 1$

Step 3 – write the CDF value: $P(X < 4) = 1 - e^{-1}$

Step 4 – substitute $e^{-1} = 0.368$: $P(X < 4) = 1 - 0.368 = 0.632$

Final Answer: $P(X < 4) = 0.632 \Rightarrow \boxed{\text{D}}$

Answer: (D) [Go Back to Q10](#)

Q11.

Solution

Concept – rank via the determinant: For a 3×3 matrix, a non-zero determinant means all three rows are independent, giving full rank 3.

Step 1 – expand the determinant along the first row: $\det A = 1 \begin{vmatrix} 1 & 1 \\ 0 & 1 \end{vmatrix} - 1 \begin{vmatrix} 0 & 1 \\ 1 & 1 \end{vmatrix} + 0$

Step 2 – evaluate the 2×2 minors: $\begin{vmatrix} 1 & 1 \\ 0 & 1 \end{vmatrix} = (1)(1) - (1)(0) = 1$

$\begin{vmatrix} 0 & 1 \\ 1 & 1 \end{vmatrix} = (0)(1) - (1)(1) = -1$

Step 3 – combine: $\det A = 1(1) - 1(-1) + 0 = 1 + 1 = 2$

Step 4 – conclude the rank: $\det A = 2 \neq 0$, so all rows are independent and $\text{rank}(A) = 3$.

Final Answer: $\text{rank}(A) = 3 \Rightarrow \boxed{\text{C}}$

Answer: (C) [Go Back to Q11](#)

Q12.

Solution

Concept – cofactor expansion along the first row: $\det B = b_{11}M_{11} - b_{12}M_{12} + b_{13}M_{13}$.

Step 1 – first-row entries: $b_{11} = 1, b_{12} = 2, b_{13} = 0$.



Step 2 – minor for b_{11} : $\begin{vmatrix} 3 & 1 \\ 0 & 4 \end{vmatrix} = (3)(4) - (1)(0) = 12$

Step 3 – minor for b_{12} : $\begin{vmatrix} 0 & 1 \\ 2 & 4 \end{vmatrix} = (0)(4) - (1)(2) = -2$

Step 4 – the third term vanishes: Since $b_{13} = 0$, its contribution is 0.

Step 5 – combine with the sign pattern: $\det B = 1(12) - 2(-2) + 0$
 $= 12 + 4$
 $= 16$

Final Answer: $\det B = 16 \Rightarrow \boxed{\text{A}}$

Answer: (A) [Go Back to Q12](#)

Q13.

Solution

Concept – eigenvalues from trace and determinant: The characteristic equation is $\lambda^2 - \text{tr}(C)\lambda + \det(C) = 0$.

Step 1 – trace and determinant of C : $\text{tr}(C) = 2 + 3 = 5$ and $\det(C) = (2)(3) - (2)(1) = 6 - 2 = 4$.

Step 2 – write the characteristic equation: $\lambda^2 - 5\lambda + 4 = 0$

Step 3 – factor: $(\lambda - 4)(\lambda - 1) = 0$

Step 4 – read the roots: $\lambda = 4$ or $\lambda = 1$.

Step 5 – pick the smaller: The smaller eigenvalue is 1.

Final Answer: $\lambda_{\min} = 1 \Rightarrow \boxed{\text{D}}$

Answer: (D) [Go Back to Q13](#)

Q14.

Solution

Concept – condition number from singular values: $\sigma_i = \sqrt{\lambda_i(A^T A)}$ and $\kappa = \frac{\sigma_{\max}}{\sigma_{\min}}$.

Step 1 – largest singular value: $\sigma_{\max} = \sqrt{16} = 4$

Step 2 – smallest singular value: $\sigma_{\min} = \sqrt{4} = 2$

Step 3 – form the ratio: $\kappa = \frac{4}{2}$

Step 4 – simplify: $\kappa = 2$



Final Answer: $\kappa = 2 \Rightarrow$ B

Answer: (B) [Go Back to Q14](#)

Q15.

Solution

Concept – the ℓ_2 norm is the square root of the sum of squares: $\|v\|_2 = \sqrt{\sum_i v_i^2}$.

Step 1 – square each component: $2^2 = 4, (-3)^2 = 9, 6^2 = 36$.

Step 2 – add the squares: $4 + 9 + 36 = 49$

Step 3 – take the square root: $\|v\|_2 = \sqrt{49} = 7$

Final Answer: $\|v\|_2 = 7 \Rightarrow$ A

Answer: (A) [Go Back to Q15](#)

Q16.

Solution

Concept – minimizer of a parabola: Set the derivative to zero; for an upward parabola the stationary point is the global minimum.

Step 1 – differentiate: $f'(x) = 2x - 6$

Step 2 – set the derivative to zero: $2x - 6 = 0$

Step 3 – solve for x : $2x = 6$

$x = 3$

Step 4 – confirm it is a minimum: $f''(x) = 2 > 0$, so $x = 3$ is a minimum.

Final Answer: minimizer $x = 3 \Rightarrow$ B

Answer: (B) [Go Back to Q16](#)

Q17.

Solution

Concept – cosine similarity: $\cos \theta = \frac{\vec{a} \cdot \vec{b}}{\|\vec{a}\| \|\vec{b}\|}$.

Step 1 – compute the dot product: $\vec{a} \cdot \vec{b} = (1)(1) + (0)(1) = 1$

Step 2 – compute the magnitudes: $\|\vec{a}\| = \sqrt{1^2 + 0^2} = 1$



$$\|\vec{b}\| = \sqrt{1^2 + 1^2} = \sqrt{2}$$

Step 3 – form the ratio: $\cos \theta = \frac{1}{1 \times \sqrt{2}} = \frac{1}{\sqrt{2}}$

Step 4 – evaluate ($\sqrt{2} = 1.414$): $\cos \theta = \frac{1}{1.414} = 0.707$

Final Answer: cosine similarity = 0.707 $\Rightarrow$ **C**

Answer: (C) [Go Back to Q17](#)

Q18.

Solution

Concept – two gradient-descent steps: Apply $x_{t+1} = x_t - \eta f'(x_t)$ twice.

Step 1 – differentiate the objective: $f(x) = x^2 \Rightarrow f'(x) = 2x$

Step 2 – first update from $x_0 = 8$: $f'(8) = 2 \times 8 = 16$

$$x_1 = 8 - 0.25 \times 16 = 8 - 4 = 4$$

Step 3 – second update from $x_1 = 4$: $f'(4) = 2 \times 4 = 8$

$$x_2 = 4 - 0.25 \times 8 = 4 - 2 = 2$$

Final Answer: $x_2 = 2 \Rightarrow$ **A**

Answer: (A) [Go Back to Q18](#)

Q19.

Solution

Concept – convexity via the Hessian: A twice-differentiable function is convex only if its Hessian is positive semi-definite everywhere; an indefinite Hessian means the function is neither convex nor concave.

Step 1 – compute the second partials: $f_{xx} = 2, f_{yy} = -2, f_{xy} = f_{yx} = 0$.

Step 2 – write the Hessian: $H = \begin{bmatrix} 2 & 0 \\ 0 & -2 \end{bmatrix}$

Step 3 – read the eigenvalues: For this diagonal matrix the eigenvalues are 2 and -2.

Step 4 – classify: One eigenvalue is positive and one is negative, so H is indefinite (a saddle); f is neither convex nor concave.

Final Answer: neither convex nor concave $\Rightarrow$ **D**

Answer: (D) [Go Back to Q19](#)



Q20.

Solution

Concept – Python slice with a negative step: `s[::-1]` walks the string from the last character to the first, producing the reverse.

Step 1 – index the string: `s = d(0) a(1) t(2) a(3)`

Step 2 – apply the step -1 : Reading from index 3 down to 0 gives a, t, a, d.

Step 3 – join the characters: The reversed string is atad.

Final Answer: atad $\Rightarrow$

Answer: (B) [Go Back to Q20](#)

Q21.

Solution

Concept – filtered list comprehension: The comprehension keeps only the values of `x` for which the condition is true, then `len` counts them.

Step 1 – list the candidates: `range(10)` yields 0, 1, 2, 3, 4, 5, 6, 7, 8, 9.

Step 2 – apply the filter `x % 3 == 0`: The multiples of 3 in the range are 0, 3, 6, 9.

Step 3 – count the surviving elements: There are 4 such values.

Step 4 – evaluate `len`: `len([0, 3, 6, 9]) = 4`

Final Answer: 4 $\Rightarrow$

Answer: (B) [Go Back to Q21](#)

Q22.

Solution

Concept – a queue is first-in, first-out (FIFO): `dequeue` always removes the element that has waited longest (the front).

Step 1 – `enqueue(1), enqueue(2)`: Queue (front $\rightarrow$ rear): [1, 2].

Step 2 – `dequeue()`: Removes the front 1; queue becomes [2].

Step 3 – `enqueue(3), enqueue(4)`: Queue becomes [2, 3, 4].

Step 4 – `dequeue()`: Removes the front 2; queue becomes [3, 4].

Step 5 – `dequeue()` (the last one): Removes the front 3; this is the value returned.

Final Answer: last `dequeue()` returns 3 $\Rightarrow$

Answer: (A) [Go Back to Q22](#)



Q23.

Solution

Concept – inserting sorted keys makes a skewed BST: Each new key is larger than all previous ones, so it always goes to the right, forming a single chain.

Step 1 – insert 10: 10 becomes the root.

Step 2 – insert 20: $20 > 10$, so it becomes the right child of 10.

Step 3 – insert 30: $30 > 10$ then $30 > 20$, so it becomes the right child of 20.

Step 4 – insert 40: 40 goes right of 10, 20, 30, becoming the right child of 30.

Step 5 – count the edges on the longest path: The chain is $10 \rightarrow 20 \rightarrow 30 \rightarrow 40$, which has 3 edges.

Final Answer: height = 3 $\Rightarrow$

Answer: (C) [Go Back to Q23](#)

Q24.

Solution

Concept – linear probing scans the next free slot: On a collision at index i , try $i + 1, i + 2, \dots$ (modulo table size).

Step 1 – insert 12: $h(12) = 12 \bmod 10 = 2$; slot 2 is free, so 12 goes to index 2.

Step 2 – insert 22: $h(22) = 2$; index 2 is occupied, so probe index 3; it is free, so 22 goes to index 3.

Step 3 – insert 32: $h(32) = 2$; index 2 occupied $\rightarrow$ probe index 3 (occupied) $\rightarrow$ probe index 4 (free).

Step 4 – place 32: 32 is stored at index 4.

Final Answer: index 4 $\Rightarrow$

Answer: (A) [Go Back to Q24](#)

Q25.

Solution

Concept – binary search halves the range each step: The worst-case number of comparisons is $\lfloor \log_2 n \rfloor + 1$.

Step 1 – take the base-2 logarithm of n : $\log_2 63 \approx 5.98$

Step 2 – take the floor: $\lfloor 5.98 \rfloor = 5$

Step 3 – add one: $5 + 1 = 6$



Step 4 – cross-check by doubling: $2^5 = 32 \leq 63 < 64 = 2^6$, so 6 halvings suffice.

Final Answer: 6 comparisons $\Rightarrow$ D

Answer: (D) [Go Back to Q25](#)

Q26.

Solution

Concept – Merge Sort always splits in half: Merge Sort divides the array into two halves and merges them in linear time, regardless of the input order.

Step 1 – write the recurrence: $T(n) = 2T(n/2) + O(n)$

Step 2 – identify the Master-theorem parameters: $a = 2$, $b = 2$, $f(n) = n$, and $n^{\log_b a} = n^{\log_2 2} = n$.

Step 3 – apply Case 2 ($f(n) = \Theta(n^{\log_b a})$): $T(n) = \Theta(n \log n)$

Step 4 – note the input-independence: The split is always balanced, so best, average and worst cases are all $O(n \log n)$.

Final Answer: worst case = $O(n \log n) \Rightarrow$ D

Answer: (D) [Go Back to Q26](#)

Q27.

Solution

Concept – DFS goes deep before backtracking: Visit an unvisited neighbour (alphabetically first), recurse, and backtrack only when a vertex has no unvisited neighbour.

Step 1 – start at A: Visit A. Its neighbours are B and C; go to the smaller, B.

Step 2 – explore from B: Visit B. Its unvisited neighbours are D and E; go to D.

Step 3 – explore from D: Visit D. It has no unvisited neighbour, so backtrack to B.

Step 4 – continue from B: Visit E (the remaining neighbour of B); it has no unvisited neighbour, so backtrack to A.

Step 5 – continue from A, then C: Visit C, then its neighbour F.

Step 6 – assemble the order: A, B, D, E, C, F.

Final Answer: A, B, D, E, C, F $\Rightarrow$ A

Answer: (A) [Go Back to Q27](#)



Q28.

Solution

Concept – Master theorem, Case 3: For $T(n) = aT(n/b) + f(n)$, if $f(n)$ grows polynomially faster than $n^{\log_b a}$, then $T(n) = \Theta(f(n))$.

Step 1 – identify the parameters: $a = 2, b = 2, f(n) = n^2$.

Step 2 – compute the critical exponent: $n^{\log_b a} = n^{\log_2 2} = n^1 = n$

Step 3 – compare $f(n)$ with $n^{\log_b a}$: $f(n) = n^2$ grows faster than n ; indeed $n^2 = \Omega(n^{1+\varepsilon})$ with $\varepsilon = 1$ (Case 3).

Step 4 – apply Case 3: $T(n) = \Theta(f(n)) = \Theta(n^2)$

Final Answer: $T(n) = \Theta(n^2) \Rightarrow \boxed{A}$

Answer: (A) [Go Back to Q28](#)

Q29.

Solution

Concept – inclusion–exclusion for union: $|R \cup S| = |R| + |S| - |R \cap S|$.

Step 1 – read the counts: $|R| = 12, |S| = 8, |R \cap S| = 3$.

Step 2 – add the two relation sizes: $12 + 8 = 20$

Step 3 – subtract the double-counted common tuples: $20 - 3 = 17$

Final Answer: $|R \cup S| = 17$ tuples $\Rightarrow \boxed{C}$

Answer: (C) [Go Back to Q29](#)

Q30.

Solution

Concept – COUNT(DISTINCT col) counts distinct non-null values: It ignores NULLs and collapses duplicates.

Step 1 – list the values: 5, 5, 10, 10, 10, NULL, 20.

Step 2 – drop the NULL: The NULL is not counted, leaving 5, 5, 10, 10, 10, 20.

Step 3 – keep only distinct values: The distinct values are 5, 10, 20.

Step 4 – count them: There are 3 distinct values.

Final Answer: COUNT(DISTINCT grade) returns 3 $\Rightarrow \boxed{B}$

Answer: (B) [Go Back to Q30](#)



Q31.

Solution

Concept – 3NF vs BCNF: A relation is in BCNF if every FD's left side is a superkey; it is in 3NF if additionally each FD either has a superkey left side *or* its right side is a prime attribute.

Step 1 – list the given FDs and keys: $AB \rightarrow C$ and $C \rightarrow B$; candidate keys AB and AC ; prime attributes $\{A, B, C\}$.

Step 2 – test $AB \rightarrow C$ for BCNF: AB is a candidate key, hence a superkey; this FD is fine for BCNF.

Step 3 – test $C \rightarrow B$ for BCNF: C alone is *not* a superkey (it does not determine A), so $C \rightarrow B$ violates BCNF.

Step 4 – test $C \rightarrow B$ for 3NF: The right side B is a prime attribute (it belongs to the key AB), so 3NF is satisfied.

Step 5 – conclude the highest normal form: R is in 3NF but not BCNF, so the highest normal form is 3NF.

Final Answer: highest form = 3NF $\Rightarrow$

Answer: (C) [Go Back to Q31](#)

Q32.

Solution

Concept – natural join matches tuples on the common attribute: Combine a tuple of R with a tuple of S whenever their B values are equal.

Step 1 – match R 's tuple $(1, x)$: S has (x, p) and (x, q) with $B = x$, giving $(1, x, p)$ and $(1, x, q)$ – 2 tuples.

Step 2 – match R 's tuple $(2, y)$: S has no tuple with $B = y$, giving 0 tuples.

Step 3 – match R 's tuple $(3, x)$: S has (x, p) and (x, q) with $B = x$, giving $(3, x, p)$ and $(3, x, q)$ – 2 tuples.

Step 4 – total the matches: $2 + 0 + 2 = 4$

Final Answer: 4 tuples $\Rightarrow$

Answer: (A) [Go Back to Q32](#)



Q33.

Solution

Concept – warehouse schema shapes: A star schema keeps its dimension tables denormalized; a snowflake normalizes them into sub-dimensions.

Step 1 – match the described layout: One central fact table linked directly to denormalized dimensions, radiating outward like a star, is the **star schema**.

Step 2 – rule out the alternatives: A snowflake schema splits dimensions into normalized tables (extra branches); a fact constellation has multiple fact tables sharing dimensions; “fully normalized” is not the standard warehouse layout for fast OLAP.

Final Answer: star schema $\Rightarrow$

Answer: (C) [Go Back to Q33](#)

Q34.

Solution

Concept – minimum fill of a B^+ tree internal node: A non-root internal node of order m must hold at least $\lceil m/2 \rceil$ child pointers.

Step 1 – substitute the order: $m = 5$, so minimum pointers = $\lceil 5/2 \rceil$.

Step 2 – evaluate the ceiling: $\frac{5}{2} = 2.5$, and $\lceil 2.5 \rceil = 3$.

Step 3 – state the result: A non-root internal node must have at least 3 child pointers.

Final Answer: 3 pointers $\Rightarrow$

Answer: (B) [Go Back to Q34](#)

Q35.

Solution

Concept – attribute closure: A^+ is the set of all attributes determined by A , built by repeatedly applying the FDs.

Step 1 – start with the attribute itself: $A^+ = \{A\}$

Step 2 – apply $A \rightarrow B$: A determines B , so $A^+ = \{A, B\}$.

Step 3 – apply $B \rightarrow C$: B is in the closure, so it brings in C : $A^+ = \{A, B, C\}$.

Step 4 – check $D \rightarrow E$: D is not in the closure, so this FD cannot fire; D and E are never added.



Step 5 – conclude: No further FD applies, so $A^+ = \{A, B, C\}$.

Final Answer: $A^+ = \{A, B, C\} \Rightarrow \boxed{C}$

Answer: (C) [Go Back to Q35](#)

Q36.

Solution

Concept – logistic regression prediction: Compute the linear score $z = w^T x + b$, then apply the sigmoid.

Step 1 – compute the linear score: $z = (0.5)(1) + (0.5)(1) + 0$

$$= 0.5 + 0.5$$

$$= 1$$

Step 2 – write the sigmoid: $\sigma(1) = \frac{1}{1 + e^{-1}}$

Step 3 – substitute $e^{-1} = 0.368$: $\sigma(1) = \frac{1}{1 + 0.368} = \frac{1}{1.368}$

Step 4 – divide: $\sigma(1) = 0.731 \approx 0.73$

Final Answer: $\sigma(1) \approx 0.73 \Rightarrow \boxed{D}$

Answer: (D) [Go Back to Q36](#)

Q37.

Solution

Concept – underfitting in bias-variance terms: A model too simple for the data has high bias, and because it barely responds to the training sample, low variance.

Step 1 – interpret the poor training accuracy: Failing even on the training set means the model cannot represent the pattern: *high bias*.

Step 2 – interpret the similar test accuracy: Its predictions change little from one training sample to another: *low variance*.

Step 3 – combine: This is the signature of underfitting: high bias, low variance.

Final Answer: high bias, low variance $\Rightarrow \boxed{D}$

Answer: (D) [Go Back to Q37](#)



Q38.

Solution

Concept – kNN takes a majority vote over the k nearest points: Compute distances from Q to every training point, keep the 3 closest, and vote.

Step 1 – distance $Q(5, 6)$ to $C(5, 5)$: $\sqrt{(5-5)^2 + (6-5)^2} = \sqrt{0+1} = 1.00$

Step 2 – distance to $E(6, 6)$: $\sqrt{(5-6)^2 + (6-6)^2} = \sqrt{1+0} = 1.00$

Step 3 – distance to $D(6, 5)$: $\sqrt{(5-6)^2 + (6-5)^2} = \sqrt{1+1} = 1.41$

Step 4 – distances to A and B : $A(1, 1) : \sqrt{16+25} = 6.40$; $B(2, 1) : \sqrt{9+25} = 5.83$.

Step 5 – pick the 3 nearest and vote: Nearest three are C (1.00, Class 1), E (1.00, Class 1), D (1.41, Class 1). All three vote Class 1.

Step 6 – assign the majority label: Class 1 wins 3 to 0.

Final Answer: Q is Class 1 $\Rightarrow$ A

Answer: (A) [Go Back to Q38](#)

Q39.

Solution

Concept – entropy of a pure node: $H = -\sum_i p_i \log_2 p_i$; a node in which every example shares one class has no uncertainty.

Step 1 – write the class proportions: All 8 examples are the same class, so $p = 1$ for that class and 0 for the other.

Step 2 – substitute into the entropy formula: $H = -1 \cdot \log_2 1 - 0 \cdot \log_2 0$

Step 3 – evaluate each term: $\log_2 1 = 0$, so the first term is 0; the $0 \cdot \log_2 0$ term is taken as 0 by convention.

Step 4 – add: $H = 0 + 0 = 0$

Final Answer: $H = 0$ bits $\Rightarrow$ D

Answer: (D) [Go Back to Q39](#)



Q40.

Solution

Concept – information gain: $IG = H(\text{parent}) - \sum_{\text{children}} \frac{n_{\text{child}}}{n_{\text{parent}}} H(\text{child})$.

Step 1 – parent entropy: The parent has 10+ and 10– (balanced), so $H(\text{parent}) = 1.0$ bit.

Step 2 – entropy of each child: Each child is pure (all one class), so $H(\text{child}) = 0$.

Step 3 – weighted child entropy: $\frac{10}{20}(0) + \frac{10}{20}(0) = 0$

Step 4 – subtract from the parent entropy: $IG = 1.0 - 0 = 1.0$

Step 5 – interpret: A perfectly separating split removes all uncertainty, giving the maximum possible gain of 1 bit.

Final Answer: $IG = 1.0$ bit $\Rightarrow$ **B**

Answer: (B) [Go Back to Q40](#)

Q41.

Solution

Concept – SVM margin width and $\|w\|$: The margin width is $\frac{2}{\|w\|}$; rearranging gives $\|w\| = \frac{2}{\text{margin}}$.

Step 1 – write the relation: $\text{margin} = \frac{2}{\|w\|}$

Step 2 – solve for $\|w\|$: $\|w\| = \frac{2}{\text{margin}}$

Step 3 – substitute margin = 0.5: $\|w\| = \frac{2}{0.5}$

Step 4 – divide: $\|w\| = 4$

Final Answer: $\|w\| = 4 \Rightarrow$ **A**

Answer: (A) [Go Back to Q41](#)

Q42.

Solution

Concept – k-means centroid update: The new centroid is the coordinate-wise mean of the points assigned to the cluster.

Step 1 – average the x -coordinates: $\bar{x} = \frac{0 + 4 + 2}{3} = \frac{6}{3} = 2$



Step 2 – average the y -coordinates: $\bar{y} = \frac{0 + 0 + 6}{3} = \frac{6}{3} = 2$

Step 3 – assemble the centroid: $(\bar{x}, \bar{y}) = (2, 2)$

Final Answer: centroid = $(2, 2) \Rightarrow$ C

Answer: (C) [Go Back to Q42](#)

Q43.

Solution

Concept – variance explained in PCA: Each principal component explains variance proportional to its eigenvalue; the fraction is that eigenvalue over the total.

Step 1 – total variance (sum of all eigenvalues): $5 + 3 + 1 + 1 = 10$

Step 2 – variance in the first component: 5

Step 3 – form the fraction: $\frac{5}{10} = 0.50$

Step 4 – express as a percentage: $0.50 = 50\%$

Final Answer: 50% of the variance $\Rightarrow$ B

Answer: (B) [Go Back to Q43](#)

Q44.

Solution

Concept – neuron forward pass with ReLU: Compute the weighted sum plus bias, then apply $\text{ReLU}(z) = \max(0, z)$.

Step 1 – weighted sum of inputs: $w_1x_1 + w_2x_2 = (1)(2) + (-0.5)(1)$

$$= 2 - 0.5$$

$$= 1.5$$

Step 2 – add the bias: $z = 1.5 + 1 = 2.5$

Step 3 – apply ReLU: $\text{ReLU}(2.5) = \max(0, 2.5) = 2.5$

Final Answer: output = 2.5 $\Rightarrow$ C

Answer: (C) [Go Back to Q44](#)



Q45.

Solution

Concept – greedy best-first uses h only: It expands the frontier node with the smallest heuristic value h , ignoring the path cost g .

Step 1 – read the heuristic of node A : $h(A) = 2$

Step 2 – read the heuristic of node B : $h(B) = 5$

Step 3 – read the heuristic of node C : $h(C) = 4$

Step 4 – choose the smallest h : $\min(2, 5, 4) = 2$, which belongs to node A .

Note: node B has the smallest g , but greedy best-first ignores g , so it does not pick B .

Final Answer: greedy expands $A \Rightarrow \boxed{A}$

Answer: (A) [Go Back to Q45](#)

Q46.

Solution

Concept – minimax back-up: MIN nodes take the minimum of their children; the MAX root takes the maximum of the MIN results.

Step 1 – evaluate the left MIN node: $\min(8, 4) = 4$

Step 2 – evaluate the right MIN node: $\min(9, 6) = 6$

Step 3 – evaluate the MAX root: $\max(4, 6) = 6$

Final Answer: minimax value = 6 $\Rightarrow \boxed{D}$

Answer: (D) [Go Back to Q46](#)



Answer Key

Q	Ans	Q	Ans	Q	Ans	Q	Ans	Q	Ans
1	D	2	B	3	D	4	B	5	A
6	B	7	C	8	C	9	B	10	D
11	C	12	A	13	D	14	B	15	A
16	B	17	C	18	A	19	D	20	B
21	B	22	A	23	C	24	A	25	D
26	D	27	A	28	A	29	C	30	B
31	C	32	A	33	C	34	B	35	C
36	D	37	D	38	A	39	D	40	B
41	A	42	C	43	B	44	C	45	A
46	D								

