

GATE Data Science and Artificial Intelligence

Sample Paper – 1

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

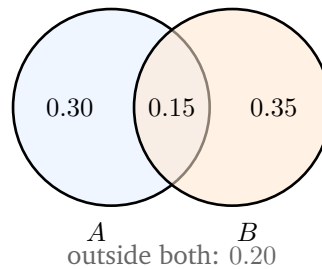
Q1. A project team of 3 members is to be selected from a pool of 5 data scientists and 4 data engineers. The number of distinct teams that contain *at least one* data engineer is: [1 mark]

- (A) 84
- (B) 10
- (C) 74

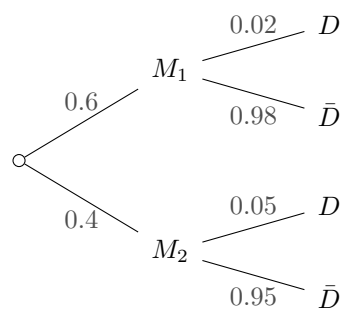


(D) 60

- Q2.** For two events A and B the probabilities are shown on the Venn diagram below (the region outside both circles has probability 0.20). The conditional probability $P(A | B)$ is: [1 mark]



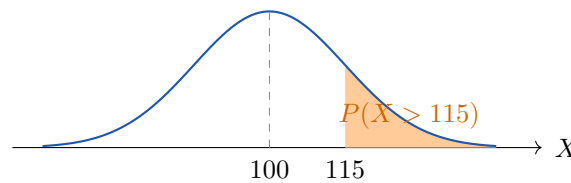
- (A) 0.50
 (B) 0.15
 (C) 0.60
 (D) 0.30
- Q3.** A workshop produces items on two machines. Machine M_1 makes 60% of the items with a defect rate of 2%; machine M_2 makes the remaining 40% with a defect rate of 5%. The probability tree is shown. An item drawn at random is found to be defective. The probability that it was produced by M_1 is: [1 mark]



- (A) 0.375
 (B) 0.625
 (C) 0.020
 (D) 0.600



- Q4.** A feature X is normally distributed with mean $\mu = 100$ and standard deviation $\sigma = 15$. Using the shaded right tail of the standard normal curve shown, the probability $P(X > 115)$ is nearest to: [1 mark]

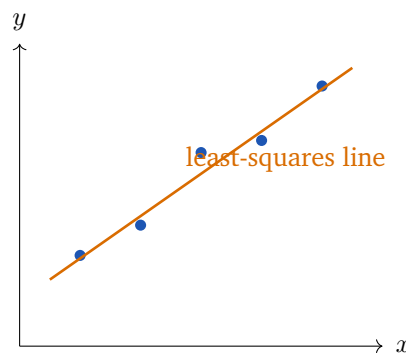


- (A) 0.8413
(B) 0.1587
(C) 0.3413
(D) 0.5000
- Q5.** Calls arrive at a support desk following a Poisson process at an average rate of 3 calls per hour. The probability of receiving *exactly* 2 calls in a given one-hour window is nearest to (take $e^{-3} = 0.0498$): [1 mark]
- (A) 0.149
(B) 0.224
(C) 0.050
(D) 0.100
- Q6.** A discrete random variable X takes the values 0, 1, 2 with probabilities 0.2, 0.5, 0.3 respectively. The variance of X is: [1 mark]
- (A) 1.10
(B) 1.70
(C) 0.49
(D) 0.70
- Q7.** A sample of size $n = 100$ drawn from a population with known standard deviation $\sigma = 10$ has sample mean $\bar{x} = 50$. For a 95% confidence interval for the population mean (using $z_{0.025} = 1.96$), the margin of error (half-width of the interval) is: [1 mark]



- (A) 1.96
- (B) 19.6
- (C) 0.196
- (D) 3.92

Q8. For a bivariate data set the sums of squares and cross-products about the means are $S_{xx} = 50$, $S_{yy} = 80$ and $S_{xy} = 40$, and the scatter with its least-squares line is shown. The slope of the least-squares regression line of y on x is: [1 mark]



- (A) 0.632
- (B) 1.250
- (C) 0.800
- (D) 0.500

Q9. A biased coin has probability 0.6 of landing heads. It is tossed 4 times independently. The probability of obtaining *exactly* 3 heads is: [1 mark]

- (A) 0.3456
- (B) 0.1296
- (C) 0.2304
- (D) 0.5000

Q10. The time between failures of a server follows an exponential distribution with a mean of 5 hours. The probability that the server runs for *more than* 10 hours without a failure is (take $e^{-2} = 0.1353$): [1 mark]



- (A) 0.8647
- (B) 0.6065
- (C) 0.1353
- (D) 0.2707

Linear Algebra, Calculus & Optimization

Q11. Consider the matrix $A = \begin{bmatrix} 1 & 2 & 3 \\ 2 & 4 & 6 \\ 1 & 1 & 1 \end{bmatrix}$. The rank of A is: [1 mark]

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

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

- (A) 10
- (B) 6
- (C) 8
- (D) 12

Q13. The larger eigenvalue of the matrix $C = \begin{bmatrix} 4 & 1 \\ 2 & 3 \end{bmatrix}$ is: [1 mark]

- (A) 2
- (B) 6
- (C) 7
- (D) 5

Q14. For a real matrix A , the matrix $A^T A$ has eigenvalues 9 and 4. The largest singular value of A (equivalently, the spectral norm $\|A\|_2$) is: [1 mark]



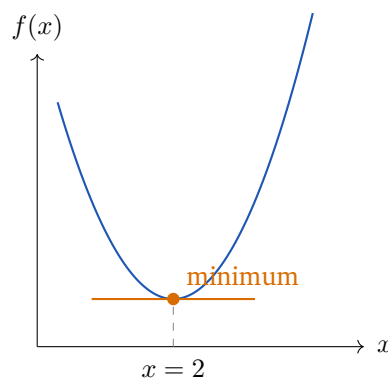
- (A) 3
- (B) 9
- (C) 6
- (D) 81

Q15. For the vector $v = (3, -4)$, the ℓ_1 norm $\|v\|_1$ is:

[1 mark]

- (A) 5
- (B) 4
- (C) 7
- (D) 25

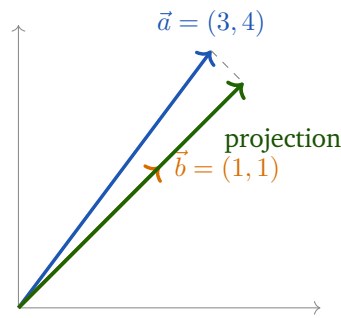
Q16. The function $f(x) = x^2 - 4x + 7$ is minimized over the real line, as suggested by its convex graph below. The minimum value of f is: [1 mark]



- (A) 2
- (B) 7
- (C) 3
- (D) 0

Q17. The scalar projection (component) of $\vec{a} = (3, 4)$ onto $\vec{b} = (1, 1)$, as shown, is nearest to: [1 mark]





- (A) 7.00
- (B) 4.95
- (C) 3.50
- (D) 2.47

Q18. Gradient descent is used to minimize $f(x) = x^2$ with a fixed learning rate $\eta = 0.1$. Starting from $x_0 = 4$, the value of x_1 after one update $x_1 = x_0 - \eta f'(x_0)$ is: [1 mark]

- (A) 3.2
- (B) 0.8
- (C) 3.6
- (D) 4.8

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

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

Programming, Data Structures & Algorithms

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

```
a = [1, 2, 3, 4]
print(a[1:3])
```

[1 mark]



- (A) [1, 2]
- (B) [2, 3, 4]
- (C) [2, 3]
- (D) [1, 2, 3]

Q21. Consider the following Python code:

```
x = [1, 2, 3]
y = x
y.append(4)
print(len(x))
```

The value printed is:

[2 marks]

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

Q22. The following operations are performed on an initially empty stack, in order: push(10), push(20), pop(), push(30), push(40), pop(), pop().

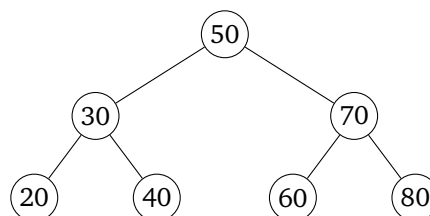
The value returned by the *last* pop() is:

[2 marks]

- (A) 40
- (B) 10
- (C) 20
- (D) 30

Q23. Starting from an empty binary search tree, the keys 50, 30, 70, 20, 40, 60, 80 are inserted in the given order, producing the tree shown. The number of *leaf* nodes in the resulting tree is:

[2 marks]



- (A) 4
- (B) 3
- (C) 2
- (D) 7

Q24. A hash table of size 10 uses the hash function $h(k) = k \bmod 10$ with *linear probing* for collision resolution. The keys 21, 31, 41 are inserted in this order into an otherwise empty table. The final index occupied by key 41 is: [2 marks]

- (A) 1
- (B) 2
- (C) 41
- (D) 3

Q25. Binary search is performed on a sorted array of 1000 distinct elements. The maximum number of key comparisons required in the worst case is: [2 marks]

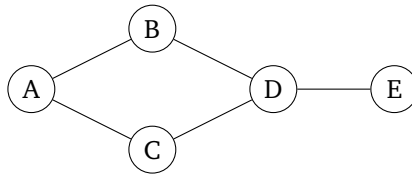
- (A) 1000
- (B) 500
- (C) 100
- (D) 10

Q26. The worst-case time complexity of the standard Quicksort algorithm (for example, when the pivot is always the smallest or largest element of the current sub-array) is: [2 marks]

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



- Q27.** A breadth-first search (BFS) is started from vertex A in the undirected graph shown, always visiting adjacent vertices in increasing alphabetical order. The order in which the vertices are first visited is: [2 marks]



- (A) A, B, D, C, E
(B) A, B, C, D, E
(C) A, C, B, E, D
(D) A, B, C, E, D
- Q28.** The running time of an algorithm satisfies the recurrence $T(n) = 2T(n/2) + n$, with $T(1) = 1$. By the Master theorem, $T(n)$ is: [2 marks]
- (A) $\Theta(n)$
(B) $\Theta(n^2)$
(C) $\Theta(\log n)$
(D) $\Theta(n \log n)$

Database Management & Warehousing

- Q29.** Relation R has 5 attributes and 8 tuples; relation S has 3 attributes and 4 tuples. The number of tuples in the Cartesian product $R \times S$ is: [2 marks]
- (A) 12
(B) 32
(C) 40
(D) 8
- Q30.** A table EMP has 10 rows. In the salary column, 2 of these rows contain NULL and the remaining 8 contain distinct non-null values. What does the query `SELECT COUNT(salary) FROM EMP;` return? [2 marks]



- (A) 10
- (B) 8
- (C) 2
- (D) 0

Q31. A relation $R(A, B, C, D)$ has the functional dependencies $A \rightarrow B$, $B \rightarrow C$, $C \rightarrow D$, with A as the only candidate key. The highest normal form that R satisfies is: [2 marks]

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

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

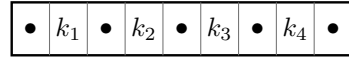
- (A) 3
- (B) 9
- (C) 2
- (D) 6

Q33. In OLAP, the operation that navigates from a summarized (coarser) level of a data cube to a more detailed (finer) level, such as from yearly totals to monthly totals, is called: [2 marks]

- (A) roll-up
- (B) slice
- (C) pivot
- (D) drill-down



- Q34.** The figure shows the structure of an internal node of a B^+ tree of order $m = 5$ (each internal node holds up to m child pointers, drawn as $\bullet$, separated by keys). The maximum number of keys an internal node of this tree can hold is: [2 marks]



up to 5 pointers ($\bullet$) and 4 keys

- (A) 5
(B) 2
(C) 4
(D) 6
- Q35.** A relation has attributes $\{A, B, C, D, E\}$ with functional dependencies $A \rightarrow BC$ and $C \rightarrow DE$. The attribute closure A^+ is: [2 marks]
- (A) $\{A, B, C\}$
(B) $\{A\}$
(C) $\{A, B, C, D, E\}$
(D) $\{A, B, C, D\}$

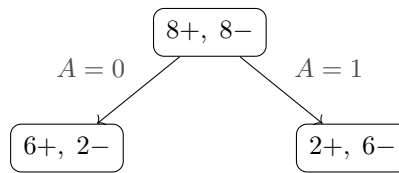
Machine Learning & Artificial Intelligence

- Q36.** A logistic regression model has weight vector $w = (1, 1)$ and bias $b = -1$, and uses the sigmoid $\sigma(z) = 1/(1 + e^{-z})$. For the input $x = (2, 1)$, the predicted probability $\sigma(w^T x + b)$ is nearest to (take $e^{-2} = 0.135$): [2 marks]
- (A) 0.50
(B) 0.88
(C) 0.12
(D) 0.73



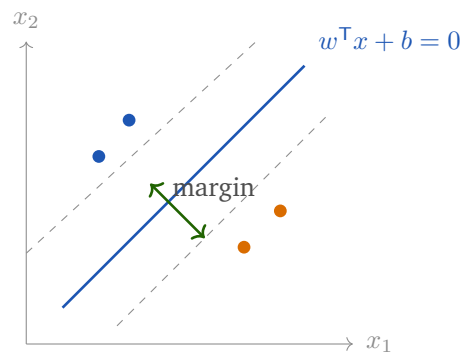
- Q37.** A flexible model attains very high accuracy on the training set but performs poorly on unseen test data. In terms of the bias–variance decomposition, this behaviour is best described as: [2 marks]
- (A) high bias, low variance
 - (B) high bias, high variance
 - (C) low bias, low variance
 - (D) low bias, high variance
- Q38.** A k -nearest-neighbour classifier with $k = 3$ (Euclidean distance) is trained on the labelled points $A(1, 2)$: Class 0, $B(2, 3)$: Class 0, $C(3, 3)$: Class 1, $D(6, 5)$: Class 1, $E(7, 7)$: Class 1. The query point $Q(3, 4)$ is classified as: [2 marks]
- (A) Class 1
 - (B) Class 0
 - (C) a tie (undecided)
 - (D) cannot be determined
- Q39.** A training set for a binary classification task contains 12 examples, of which 9 are positive and 3 are negative. The entropy of this set (in bits) is nearest to: [2 marks]
- (A) 0.811
 - (B) 1.000
 - (C) 0.500
 - (D) 0.750
- Q40.** A node in a decision tree contains 8 positive and 8 negative examples. A candidate attribute A splits it into two children as shown: $A = 0$ gives $(6+, 2-)$ and $A = 1$ gives $(2+, 6-)$. The information gain of this split is nearest to (each child has entropy 0.811 bits): [2 marks]





- (A) 0.189
- (B) 0.811
- (C) 1.000
- (D) 0.500

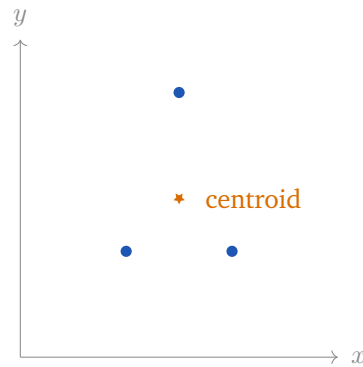
Q41. For a linear support vector machine the separating hyperplane $w^T x + b = 0$ is shown with its two margin boundaries. If $\|w\| = 0.5$, the width of the margin, equal to $2/\|w\|$, is: [2 marks]



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

Q42. In one iteration of the k -means algorithm, a cluster is assigned the three points $(2, 2)$, $(4, 2)$ and $(3, 5)$, shown below. The updated centroid of this cluster is: [2 marks]



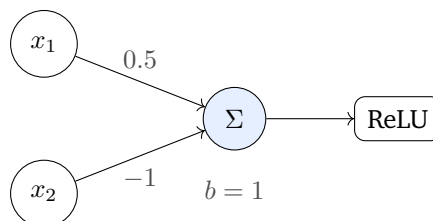


- (A) (3, 3)
- (B) (3, 4)
- (C) (2, 3)
- (D) (9, 9)

Q43. Principal component analysis is applied to a data set whose covariance matrix has eigenvalues 4, 3, 2, 1. The fraction of the total variance explained by the first two principal components is: [2 marks]

- (A) 40%
- (B) 30%
- (C) 90%
- (D) 70%

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 $x_1 = 1$, $x_2 = 2$, weights $w_1 = 0.5$, $w_2 = -1$ and bias $b = 1$, the output of the neuron is: [2 marks]



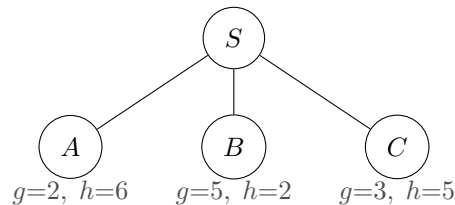
- (A) 0
- (B) -0.5



(C) 0.5

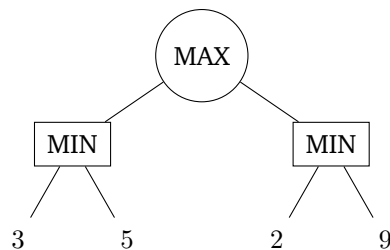
(D) 1

- Q45.** A* search maintains three frontier nodes A , B and C with g (path cost so far) and h (heuristic) values shown in the tree below, where $f = g + h$. The node A* 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) 2

(B) 3

(C) 5

(D) 9



Detailed Solutions

Q1.

Solution

Concept – count by complement: “At least one engineer” is easiest to count as (all teams) – (teams with no engineer).

Step 1 – total number of teams of 3 from 9 people: $\binom{9}{3} = \frac{9 \times 8 \times 7}{3 \times 2 \times 1}$

$$= \frac{504}{6}$$

$$= 84$$

Step 2 – teams with no engineer (all 3 from the 5 data scientists): $\binom{5}{3} = \frac{5 \times 4 \times 3}{3 \times 2 \times 1}$

$$= \frac{60}{6}$$

$$= 10$$

Step 3 – subtract to enforce “at least one engineer”: $84 - 10 = 74$

Final Answer: 74 teams $\Rightarrow$ C

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

Q2.

Solution

Concept – conditional probability: $P(A | B) = \frac{P(A \cap B)}{P(B)}$.

Step 1 – read the overlap from the Venn diagram: $P(A \cap B) = 0.15$ (the shared region).

Step 2 – read $P(B)$ as the whole of circle B : $P(B) = 0.35 + 0.15 = 0.50$

Step 3 – substitute: $P(A | B) = \frac{0.15}{0.50}$

Step 4 – divide: $P(A | B) = 0.30$

Step 5 – sanity check the whole sample space: $0.30 + 0.15 + 0.35 + 0.20 = 1.00$, so the diagram is consistent.

Final Answer: $P(A | B) = 0.30 \Rightarrow$ D

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



Q3.

Solution

Concept – Bayes' theorem (reverse the tree): $P(M_1 | D) = \frac{P(M_1) P(D | M_1)}{P(M_1) P(D | M_1) + P(M_2) P(D | M_2)}$.

Step 1 – probability from the M_1 branch: $P(M_1) P(D | M_1) = 0.60 \times 0.02 = 0.012$

Step 2 – probability from the M_2 branch: $P(M_2) P(D | M_2) = 0.40 \times 0.05 = 0.020$

Step 3 – total probability of a defective item: $P(D) = 0.012 + 0.020 = 0.032$

Step 4 – apply Bayes' theorem: $P(M_1 | D) = \frac{0.012}{0.032}$

Step 5 – divide: $P(M_1 | D) = 0.375$

Final Answer: $P(M_1 | D) = 0.375 \Rightarrow \boxed{A}$

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

Q4.

Solution

Concept – standardize, then use the normal table: Convert X to the standard normal $Z = \frac{X - \mu}{\sigma}$.

Step 1 – standardize the cut-off $X = 115$: $Z = \frac{115 - 100}{15}$

$$Z = \frac{15}{15} = 1$$

Step 2 – write the required probability in Z : $P(X > 115) = P(Z > 1)$

Step 3 – use the standard normal value: $P(Z \leq 1) = 0.8413$

Step 4 – take the right tail: $P(Z > 1) = 1 - 0.8413 = 0.1587$

Final Answer: $P(X > 115) = 0.1587 \Rightarrow \boxed{B}$

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

Q5.

Solution

Concept – Poisson probability: $P(X = k) = \frac{e^{-\lambda} \lambda^k}{k!}$, with mean rate λ .

Step 1 – list the parameters: $\lambda = 3$ per hour, $k = 2$.

Step 2 – substitute into the formula: $P(X = 2) = \frac{e^{-3} 3^2}{2!}$

Step 3 – evaluate the powers and factorial: $3^2 = 9$ and $2! = 2$



$$P(X = 2) = \frac{e^{-3} \times 9}{2}$$

Step 4 – substitute $e^{-3} = 0.0498$: $P(X = 2) = \frac{0.0498 \times 9}{2} = \frac{0.4482}{2}$

Step 5 – divide: $P(X = 2) = 0.224$

Final Answer: $P(X = 2) \approx 0.224 \Rightarrow \boxed{\text{B}}$

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

Q6.

Solution

Concept – variance from the definition: $\text{Var}(X) = E[X^2] - (E[X])^2$.

Step 1 – compute the mean $E[X]$: $E[X] = 0(0.2) + 1(0.5) + 2(0.3)$

$$= 0 + 0.5 + 0.6$$

$$= 1.1$$

Step 2 – compute $E[X^2]$: $E[X^2] = 0^2(0.2) + 1^2(0.5) + 2^2(0.3)$

$$= 0 + 0.5 + 1.2$$

$$= 1.7$$

Step 3 – square the mean: $(E[X])^2 = (1.1)^2 = 1.21$

Step 4 – subtract: $\text{Var}(X) = 1.7 - 1.21 = 0.49$

Final Answer: $\text{Var}(X) = 0.49 \Rightarrow \boxed{\text{C}}$

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

Q7.

Solution

Concept – margin of error of a confidence interval: For a known- σ mean, the margin of error is $E = z \frac{\sigma}{\sqrt{n}}$.

Step 1 – compute the standard error: $\frac{\sigma}{\sqrt{n}} = \frac{10}{\sqrt{100}}$

$$= \frac{10}{10}$$

$$= 1$$

Step 2 – multiply by the critical value: $E = 1.96 \times 1$

$$E = 1.96$$

Step 3 – interpret: The 95% interval is 50 ± 1.96 , i.e. (48.04, 51.96); its half-width



is 1.96.

Final Answer: margin of error = 1.96 $\Rightarrow$ **A**

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

Q8.

Solution

Concept – least-squares slope: The slope of the regression line of y on x is $b_1 = \frac{S_{xy}}{S_{xx}}$.

Step 1 – list the sums of squares: $S_{xx} = 50$, $S_{xy} = 40$ (the value S_{yy} is not needed for the slope).

Step 2 – substitute: $b_1 = \frac{40}{50}$

Step 3 – simplify: $b_1 = 0.80$

Step 4 – note on correlation: For interest, $r = \frac{S_{xy}}{\sqrt{S_{xx}S_{yy}}} = \frac{40}{\sqrt{50 \times 80}} = \frac{40}{63.25} = 0.632$; this is option (A) and is *not* the slope.

Final Answer: slope = 0.80 $\Rightarrow$ **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: $n = 4$, $k = 3$, $p = 0.6$, $1 - p = 0.4$.

Step 2 – evaluate the binomial coefficient: $\binom{4}{3} = 4$

Step 3 – evaluate the powers: $p^3 = 0.6^3 = 0.216$ and $(1 - p)^1 = 0.4$

Step 4 – multiply the pieces: $P(X = 3) = 4 \times 0.216 \times 0.4$

$$= 4 \times 0.0864$$

$$= 0.3456$$

Final Answer: $P(X = 3) = 0.3456 \Rightarrow$ **A**

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



Q10.

Solution

Concept – exponential survival probability: For an exponential variable with rate λ , $P(X > t) = e^{-\lambda t}$.

Step 1 – find the rate from the mean: Mean = $\frac{1}{\lambda} = 5$ hours, so $\lambda = \frac{1}{5} = 0.2 \text{ h}^{-1}$.

Step 2 – form the exponent: $\lambda t = 0.2 \times 10 = 2$

Step 3 – write the survival probability: $P(X > 10) = e^{-2}$

Step 4 – substitute $e^{-2} = 0.1353$: $P(X > 10) = 0.1353$

Final Answer: $P(X > 10) = 0.1353 \Rightarrow \boxed{\text{C}}$

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

Q11.

Solution

Concept – rank is the number of independent rows: Reduce the rows and count how many are linearly independent.

Step 1 – compare rows 1 and 2: Row 2 = $(2, 4, 6) = 2 \times (1, 2, 3) = 2R_1$, so row 2 is dependent on row 1.

Step 2 – eliminate row 2: $R_2 \leftarrow R_2 - 2R_1 = (0, 0, 0)$

Step 3 – test row 3 against row 1: Row 3 = $(1, 1, 1)$ is not a multiple of $(1, 2, 3)$, so it is independent.

Step 4 – count the non-zero independent rows: Independent rows are R_1 and R_3 , giving rank = 2.

Final Answer: $\text{rank}(A) = 2 \Rightarrow \boxed{\text{B}}$

Answer: (B) [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: $b_{11} = 2, b_{12} = 1, b_{13} = 0$.

Step 2 – minor for b_{11} : $\begin{vmatrix} 3 & 1 \\ 1 & 2 \end{vmatrix} = (3)(2) - (1)(1) = 6 - 1 = 5$



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

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

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

$$= 10 - 2$$

$$= 8$$

Final Answer: $\det B = 8 \Rightarrow \boxed{\text{C}}$

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

Q13.

Solution

Concept – eigenvalues from the characteristic equation: Solve $\det(C - \lambda I) = 0$; use trace and determinant as a shortcut.

Step 1 – trace and determinant of C : $\text{tr}(C) = 4 + 3 = 7$ and $\det(C) = (4)(3) - (1)(2) = 12 - 2 = 10$.

Step 2 – write the characteristic equation: $\lambda^2 - \text{tr}(C)\lambda + \det(C) = 0$

$$\lambda^2 - 7\lambda + 10 = 0$$

Step 3 – factor: $(\lambda - 5)(\lambda - 2) = 0$

Step 4 – read the roots: $\lambda = 5$ or $\lambda = 2$.

Step 5 – pick the larger: The larger eigenvalue is 5.

Final Answer: $\lambda_{\max} = 5 \Rightarrow \boxed{\text{D}}$

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

Q14.

Solution

Concept – singular values are square roots of the eigenvalues of $A^T A$: $\sigma_i = \sqrt{\lambda_i(A^T A)}$, and $\|A\|_2 = \sigma_{\max}$.

Step 1 – take the square roots of the given eigenvalues: $\sigma_1 = \sqrt{9} = 3$ and $\sigma_2 = \sqrt{4} = 2$.

Step 2 – identify the largest singular value: $\sigma_{\max} = 3$.

Step 3 – relate to the spectral norm: $\|A\|_2 = \sigma_{\max} = 3$.

Final Answer: $\sigma_{\max} = 3 \Rightarrow \boxed{\text{A}}$



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

Q15.

Solution

Concept – the ℓ_1 norm is the sum of absolute values: $\|v\|_1 = \sum_i |v_i|$.

Step 1 – take absolute values of the components: $|3| = 3$ and $|-4| = 4$.

Step 2 – add them: $\|v\|_1 = 3 + 4 = 7$

Step 3 – contrast with the other norms: $\|v\|_2 = \sqrt{3^2 + 4^2} = 5$ and $\|v\|_\infty = \max(3, 4) = 4$; the ℓ_1 value is 7.

Final Answer: $\|v\|_1 = 7 \Rightarrow \boxed{\text{C}}$

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

Q16.

Solution

Concept – minimize by setting the derivative to zero: For a convex parabola, the stationary point is the global minimum.

Step 1 – differentiate: $f'(x) = 2x - 4$

Step 2 – set the derivative to zero: $2x - 4 = 0$

$$x = 2$$

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

Step 4 – evaluate f at $x = 2$: $f(2) = 2^2 - 4(2) + 7$

$$= 4 - 8 + 7$$

$$= 3$$

Final Answer: minimum value = 3 $\Rightarrow \boxed{\text{C}}$

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

Q17.

Solution

Concept – scalar projection: The component of $\vec{a}$ along $\vec{b}$ is $\text{comp}_{\vec{b}} \vec{a} = \frac{\vec{a} \cdot \vec{b}}{\|\vec{b}\|}$.

Step 1 – compute the dot product: $\vec{a} \cdot \vec{b} = (3)(1) + (4)(1) = 3 + 4 = 7$

Step 2 – compute the magnitude of $\vec{b}$: $\|\vec{b}\| = \sqrt{1^2 + 1^2} = \sqrt{2}$



Step 3 – form the scalar projection: $\text{comp}_{\vec{b}} \vec{a} = \frac{7}{\sqrt{2}}$

Step 4 – evaluate ($\sqrt{2} = 1.414$): $\frac{7}{1.414} = 4.95$

Final Answer: scalar projection = 4.95 $\Rightarrow$ **B**

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

Q18.

Solution

Concept – one gradient-descent step: $x_1 = x_0 - \eta f'(x_0)$.

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

Step 2 – evaluate the gradient at $x_0 = 4$: $f'(4) = 2 \times 4 = 8$

Step 3 – apply the update with $\eta = 0.1$: $x_1 = 4 - 0.1 \times 8$

Step 4 – finish the arithmetic: $x_1 = 4 - 0.8 = 3.2$

Final Answer: $x_1 = 3.2 \Rightarrow$ **A**

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

Q19.

Solution

Concept – convexity via the Hessian: A twice-differentiable function is convex iff its Hessian is positive semi-definite everywhere.

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

Step 2 – write the Hessian: $H = \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix}$

Step 3 – test with the leading principal minors: First minor: $2 > 0$. Second minor: $\det H = (2)(2) - (-1)(-1) = 4 - 1 = 3 > 0$.

Step 4 – conclude: Both leading minors are positive, so H is positive definite and f is convex everywhere.

Final Answer: f is convex $\Rightarrow$ **C**

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



Q20.

Solution

Concept – Python list slicing: $a[i:j]$ returns elements from index i up to *but not including* index j (0-based indexing).

Step 1 – index the list: $a = \left[\underbrace{1}_0, \underbrace{2}_1, \underbrace{3}_2, \underbrace{4}_3 \right]$

Step 2 – apply the slice $a[1:3]$: Take indices 1 and 2 (index 3 is excluded).

Step 3 – read off the elements: Index 1 $\rightarrow$ 2, index 2 $\rightarrow$ 3, giving $[2, 3]$.

Final Answer: $[2, 3] \Rightarrow \boxed{C}$

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

Q21.

Solution

Concept – lists are mutable objects held by reference: $y = x$ makes y an *alias* of the same list object, not a copy.

Step 1 – after $x = [1, 2, 3]$: x refers to a list of length 3.

Step 2 – after $y = x$: y and x point to the *same* underlying list.

Step 3 – after $y.append(4)$: The shared list becomes $[1, 2, 3, 4]$; the change is visible through both names.

Step 4 – evaluate $len(x)$: x now has 4 elements, so $len(x)$ is 4.

Final Answer: $4 \Rightarrow \boxed{D}$

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

Q22.

Solution

Concept – a stack is last-in, first-out (LIFO): Track the top of the stack after each operation.

Step 1 – $push(10), push(20)$: Stack (bottom $\rightarrow$ top): $[10, 20]$.

Step 2 – $pop()$: Removes the top 20; stack becomes $[10]$.

Step 3 – $push(30), push(40)$: Stack becomes $[10, 30, 40]$.

Step 4 – $pop()$: Removes the top 40; stack becomes $[10, 30]$.

Step 5 – $pop()$ (the last one): Removes the top 30; this is the value returned.

Final Answer: last $pop()$ returns 30 $\Rightarrow \boxed{D}$



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

Q23.

Solution

Concept – a leaf has no children: Build the BST and count nodes with no descendants.

Step 1 – place the root: 50 is inserted first and becomes the root.

Step 2 – place the next level: $30 < 50$ goes left; $70 > 50$ goes right.

Step 3 – place the remaining keys: $20 < 30$ (left of 30); 40 is between 30 and 50 (right of 30); 60 is between 50 and 70 (left of 70); $80 > 70$ (right of 70).

Step 4 – identify the leaves: The nodes with no children are 20, 40, 60, 80.

Step 5 – count them: There are 4 leaf nodes.

Final Answer: 4 leaves $\Rightarrow$

Answer: (A) [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 21: $h(21) = 21 \bmod 10 = 1$; slot 1 is free, so 21 goes to index 1.

Step 2 – insert 31: $h(31) = 1$; index 1 is occupied, so probe index 2; it is free, so 31 goes to index 2.

Step 3 – insert 41: $h(41) = 1$; index 1 occupied $\rightarrow$ probe index 2 (occupied) $\rightarrow$ probe index 3 (free).

Step 4 – place 41: 41 is stored at index 3.

Final Answer: index 3 $\Rightarrow$

Answer: (D) [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 : $\log_2 1000 \approx 9.97$

Step 2 – take the floor: $\lfloor 9.97 \rfloor = 9$

Step 3 – add one: $9 + 1 = 10$

Step 4 – cross-check by doubling: $2^9 = 512 < 1000 \leq 1024 = 2^{10}$, so 10 halvings suffice.

Final Answer: 10 comparisons $\Rightarrow$ **D**

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

Q26.

Solution

Concept – Quicksort's worst case comes from unbalanced partitions: When the pivot is always the extreme element, one partition is empty and the other has $n - 1$ items.

Step 1 – write the worst-case recurrence: $T(n) = T(n - 1) + O(n)$

Step 2 – unroll the recurrence: $T(n) = O(n) + O(n - 1) + \dots + O(1)$

Step 3 – sum the series: $n + (n - 1) + \dots + 1 = \frac{n(n + 1)}{2}$

Step 4 – express in big-O: $\frac{n(n + 1)}{2} = O(n^2)$

Why the other options are wrong: $O(n \log n)$ is the *average* case and the best case (balanced pivots), not the worst case.

Final Answer: worst case = $O(n^2) \Rightarrow$ **A**

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

Q27.

Solution

Concept – BFS explores level by level using a queue: Visit the source, then all its neighbours, then their unvisited neighbours, and so on.

Step 1 – start at A : Visit A . Its neighbours are B and C .

Step 2 – visit neighbours of A in alphabetical order: Enqueue and visit B , then



C.

Step 3 – process *B*, then *C*: *B*'s unvisited neighbour is *D*; *C*'s neighbour *D* is already queued. Visit *D*.

Step 4 – process *D*: *D*'s unvisited neighbour is *E*. Visit *E*.

Step 5 – assemble the order: *A*, *B*, *C*, *D*, *E*.

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

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

Q28.

Solution

Concept – Master theorem: For $T(n) = aT(n/b) + f(n)$, compare $f(n)$ with $n^{\log_b a}$.

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

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

Step 3 – compare $f(n)$ with $n^{\log_b a}$: $f(n) = n = \Theta(n^{\log_b a})$, which is Master-theorem Case 2.

Step 4 – apply Case 2: $T(n) = \Theta(n^{\log_b a} \log n) = \Theta(n \log n)$

Final Answer: $T(n) = \Theta(n \log n) \Rightarrow$ **D**

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

Q29.

Solution

Concept – Cartesian product pairs every tuple of *R* with every tuple of *S*: $|R \times S| = |R| \times |S|$, and its degree is the sum of the two degrees.

Step 1 – read the tuple counts: $|R| = 8$, $|S| = 4$.

Step 2 – multiply: $|R \times S| = 8 \times 4 = 32$

Step 3 – note the degree (not asked, for completeness): Number of attributes $= 5 + 3 = 8$.

Final Answer: 32 tuples $\Rightarrow$ **B**

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



Q30.

Solution

Concept – COUNT(column) ignores NULLs: COUNT(expr) counts only rows where expr is non-null, unlike COUNT(*) which counts every row.

Step 1 – total rows in the table: 10 rows.

Step 2 – rows with a NULL salary: 2 rows are NULL and are skipped by COUNT(salary).

Step 3 – subtract: $10 - 2 = 8$ non-null salaries are counted.

Final Answer: COUNT(salary) returns 8 $\Rightarrow$ **B**

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

Q31.

Solution

Concept – normal forms up to 3NF: 2NF forbids partial dependencies; 3NF additionally forbids transitive dependencies of non-prime attributes on the key.

Step 1 – identify prime and non-prime attributes: The only candidate key is A , so A is prime and B, C, D are non-prime.

Step 2 – check 2NF (partial dependency): The key $\{A\}$ has a single attribute, so no proper subset exists; there can be no partial dependency. Hence R is in 2NF.

Step 3 – check 3NF (transitive dependency): $A \rightarrow B$ and $B \rightarrow C$ give the transitive dependency $A \rightarrow C$ through the non-prime attribute B . This violates 3NF.

Step 4 – conclude the highest normal form: R satisfies 2NF but not 3NF, so the highest normal form is 2NF.

Final Answer: highest form = 2NF $\Rightarrow$ **D**

Answer: (D) [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 (y, r) with $B = y$, giving $(2, y, r) - 1$ tuple.

Step 3 – match R 's tuple $(3, z)$: S has no tuple with $B = z$, giving 0 tuples.

Step 4 – total the matches: $2 + 1 + 0 = 3$

Final Answer: 3 tuples $\Rightarrow$

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

Q33.

Solution

Concept – OLAP navigation operations: Roll-up aggregates to a coarser level; drill-down does the opposite, moving to finer detail.

Step 1 – describe the direction of movement: Going from yearly totals to monthly totals *adds* detail (less aggregation).

Step 2 – match the operation: Moving to a more detailed level is **drill-down**.

Why the other options are wrong: Roll-up is the reverse (more aggregation); slice fixes one dimension to a single value; pivot rotates the cube's axes.

Final Answer: drill-down $\Rightarrow$

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

Q34.

Solution

Concept – order and keys in a B^+ tree: An internal node of order m holds at most m child pointers, and the number of keys is always one fewer than the number of pointers.

Step 1 – maximum pointers: Order $m = 5 \Rightarrow$ up to 5 child pointers.

Step 2 – keys are one fewer than pointers: Maximum keys $= m - 1 = 5 - 1$

Step 3 – evaluate: $= 4$

Final Answer: 4 keys $\Rightarrow$

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



Q35.

Solution

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

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

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

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

Step 4 – check for more: All five attributes are present; no further FD adds anything.

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

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: $z = (1)(2) + (1)(1) + (-1)$

$$= 2 + 1 - 1$$

$$= 2$$

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

Step 3 – substitute $e^{-2} = 0.135$: $\sigma(2) = \frac{1}{1 + 0.135} = \frac{1}{1.135}$

Step 4 – divide: $\sigma(2) = 0.881 \approx 0.88$

Final Answer: $\sigma(2) \approx 0.88 \Rightarrow \boxed{B}$

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

Q37.

Solution

Concept – overfitting in bias–variance terms: A model that fits training data closely but generalizes poorly has low bias and high variance.

Step 1 – interpret the high training accuracy: Fitting the training set very well means the model is flexible and has *low bias*.



Step 2 – interpret the poor test accuracy: Sensitivity to the particular training sample (poor generalization) means *high variance*.

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

Final Answer: low bias, high variance $\Rightarrow$ 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 $Q(3, 4)$ to $A(1, 2)$: $\sqrt{(3-1)^2 + (4-2)^2} = \sqrt{4+4} = \sqrt{8} = 2.83$

Step 2 – distance to $B(2, 3)$: $\sqrt{(3-2)^2 + (4-3)^2} = \sqrt{1+1} = \sqrt{2} = 1.41$

Step 3 – distance to $C(3, 3)$: $\sqrt{(3-3)^2 + (4-3)^2} = \sqrt{0+1} = 1.00$

Step 4 – distances to D and E : $D(6, 5) : \sqrt{9+1} = 3.16$; $E(7, 7) : \sqrt{16+9} = 5.00$.

Step 5 – pick the 3 nearest and vote: Nearest three are C (1.00, Class 1), B (1.41, Class 0), A (2.83, Class 0). Votes: Class 0 twice, Class 1 once.

Step 6 – assign the majority label: Class 0 wins 2 to 1.

Final Answer: Q is Class 0 $\Rightarrow$ B

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

Q39.

Solution

Concept – binary entropy: $H = -p \log_2 p - (1-p) \log_2 (1-p)$, with p the fraction of positives.

Step 1 – compute the class proportions: $p = \frac{9}{12} = 0.75$ positive, $1-p = \frac{3}{12} = 0.25$ negative.

Step 2 – take the base-2 logs: $\log_2 0.75 = -0.415$ and $\log_2 0.25 = -2$.

Step 3 – form each term: $-0.75 \times (-0.415) = 0.311$

$-0.25 \times (-2) = 0.500$

Step 4 – add the terms: $H = 0.311 + 0.500 = 0.811$ bits

Final Answer: $H \approx 0.811$ bits $\Rightarrow$ A



Answer: (A) [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 8+ and 8– (balanced), so $H(\text{parent}) = 1.0$ bit.

Step 2 – child sizes and entropies: Each child has 8 examples out of 16; each child entropy is 0.811 bits (given, from a 6:2 split).

Step 3 – weighted child entropy: $\frac{8}{16}(0.811) + \frac{8}{16}(0.811) = 0.5(0.811) + 0.5(0.811) = 0.811$

Step 4 – subtract from the parent entropy: $IG = 1.0 - 0.811 = 0.189$

Final Answer: $IG \approx 0.189$ bits $\Rightarrow$ **A**

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

Q41.

Solution

Concept – SVM margin width: For the hyperplane $w^T x + b = 0$, the full margin (distance between the two boundaries) is $\frac{2}{\|w\|}$.

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

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

Step 3 – divide: $\text{margin} = 4$

Final Answer: margin width = 4 $\Rightarrow$ **B**

Answer: (B) [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: $\bar{x} = \frac{2 + 4 + 3}{3} = \frac{9}{3} = 3$



Step 2 – average the y -coordinates: $\bar{y} = \frac{2 + 2 + 5}{3} = \frac{9}{3} = 3$

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

Final Answer: centroid = $(3, 3) \Rightarrow$ A

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

Q43.

Solution

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

Step 1 – total variance (sum of all eigenvalues): $4 + 3 + 2 + 1 = 10$

Step 2 – variance in the first two components: $4 + 3 = 7$

Step 3 – form the fraction: $\frac{7}{10} = 0.70$

Step 4 – express as a percentage: $0.70 = 70\%$

Final Answer: 70% of the variance $\Rightarrow$ D

Answer: (D) [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: $w_1x_1 + w_2x_2 = (0.5)(1) + (-1)(2)$

$= 0.5 - 2$

$= -1.5$

Step 2 – add the bias: $z = -1.5 + 1 = -0.5$

Step 3 – apply ReLU: $\text{ReLU}(-0.5) = \max(0, -0.5) = 0$

Final Answer: output = 0 $\Rightarrow$ A

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



Q45.

Solution

Concept – A* expands the node of least $f = g + h$: Compute f for each frontier node and pick the minimum.

Step 1 – f for node A : $f(A) = g + h = 2 + 6 = 8$

Step 2 – f for node B : $f(B) = 5 + 2 = 7$

Step 3 – f for node C : $f(C) = 3 + 5 = 8$

Step 4 – choose the smallest f : $\min(8, 7, 8) = 7$, which belongs to node B .

Final Answer: A* expands $B \Rightarrow \boxed{B}$

Answer: (B) [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: $\min(3, 5) = 3$

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

Step 3 – evaluate the MAX root: $\max(3, 2) = 3$

Final Answer: minimax value = 3 $\Rightarrow \boxed{B}$

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



Answer Key

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

