

GATE 2024 CS and IT Set 2 Question Paper with Solution

Time Allowed :3 Hours	Maximum Marks :100	Total Questions :65
------------------------------	---------------------------	----------------------------

General Instructions

Read the following instructions very carefully and strictly follow them:

1. The GATE Exam will be structured with a total of 100 marks.
2. The exam mode is Online CBT (Computer Based Test)
3. The total duration of Exam is 3 Hours.
4. It will include 65 questions , divided in 3 sections.
5. Section 1 : General Aptitude.
6. Section 2 : Engineering Mathematics.
7. Section 3 : Subject Based Questions.
8. The marking scheme is as such : 1 and 2 marks Questions. Each correct answer will carry marks as specified in the question paper. Incorrect answers may carry negative marks, as indicated in the question paper.
9. Question Types: The exam will include a mix of Multiple Choice Questions (MCQs), Multiple Select Questions (MSQs), and Numerical Answer Type (NAT). questions.

GENERAL APTITUDE

Question 1-5 carry one mark each

1. If $\rightarrow$ denotes increasing order of intensity, then the meaning of the words [walk $\rightarrow$ jog $\rightarrow$ sprint] is analogous to [bothered $\rightarrow$ _____ $\rightarrow$ daunted].

- (A) phased
- (B) phrased
- (C) fazed
- (D) fused

Correct Answer: (C) fazed

Solution:

The analogy in the first pair represents a progression in intensity:

walk $\rightarrow$ jog $\rightarrow$ sprint.

This pattern suggests a shift from a low to a high level of physical activity.

Similarly, the second pair should represent an emotional progression:

bothered $\rightarrow$ ___ $\rightarrow$ daunted.

Step 1: Understand the emotional progression - "Bothered" signifies a mild emotional reaction, while "daunted" indicates a significant or severe emotional impact. - The missing term must logically represent an intermediate stage between "bothered" and "daunted."

Step 2: Evaluate the given options (1) **Phased**: Incorrect, as it relates to phases or stages, not emotional intensity. (2) **Phrased**: Irrelevant, as it pertains to wording or expressions. (3)

Fazed: Correct, as it describes being disconcerted or disturbed to a moderate degree, fitting the progression. (4) **Fused**: Incorrect, as it has no relation to emotions.

Conclusion: The term "fazed" logically completes the progression:

bothered $\rightarrow$ fazed $\rightarrow$ daunted.

Final Answer:

fazed

Quick Tip

When solving analogy questions, identify the type of relationship (e.g., progression, contrast, cause-effect) in the first pair to guide your choice for the second. Focus on meaning and context to eliminate irrelevant options.

2. Two wizards try to create a spell using all the four elements, water, air, fire, and earth. For this, they decide to mix all these elements in all possible orders. They also decide to work independently. After trying all possible combinations of elements, they conclude that the spell does not work.

How many attempts does each wizard make before coming to this conclusion, independently?

- (A) 24
- (B) 48
- (C) 16
- (D) 12

Correct Answer: (A) 24

Solution:

To determine the number of possible orders for n elements, we use the concept of permutations. The total number of permutations is calculated using the factorial of n , denoted as $n!$.

Step 1: Determine permutations for 4 elements The formula for permutations of n elements is:

$$n! = n \times (n - 1) \times (n - 2) \dots \times 1$$

For $n = 4$, this becomes:

$$4! = 4 \times 3 \times 2 \times 1 = 24.$$

Step 2: Analyze the scenario Both wizards independently attempt all possible orders of the 4 elements. Since the total number of permutations for 4 elements is 24, each wizard also makes 24 attempts.

Conclusion: The total number of possible orders for 4 elements is:

24

Quick Tip

For any problem involving arrangements, remember that the total number of permutations of n elements is calculated as $n!$. This formula is essential for determining order-based possibilities.

3. In an engineering college of 10,000 students, 1,500 like neither their core branches nor other branches. The number of students who like their core branches is $\frac{1}{4}$ of the number of students who like other branches. The number of students who like both their core and other branches is 500.

The number of students who like their core branches is:

- (A) 1,800
- (B) 3,500
- (C) 1,600
- (D) 1,500

Correct Answer: (A) 1,800

Solution:

To solve this problem, we use set theory concepts to analyze the given relationships. Let Ω represent the universal set (total students), C represent students who like their core branch, and O represent students who like other branches.

We are given:

$$|\Omega| = 10,000, \quad |C^c \cap O^c| = 1,500, \quad |O| = 4|C|, \quad |C \cap O| = 500.$$

The goal is to determine $|C|$.

Step 1: Use complement relationships to find $|C \cup O|$ The complement $|C^c \cap O^c|$ represents students who like neither their core branch nor other branches. Using the complement formula for the union:

$$|C \cup O| = |\Omega| - |C^c \cap O^c|.$$

Substitute the given values:

$$|C \cup O| = 10,000 - 1,500 = 8,500.$$

Step 2: Relate $|C|$ and $|O|$ using the union formula The union formula states:

$$|C \cup O| = |C| + |O| - |C \cap O|.$$

Substitute $|C \cup O| = 8,500$, $|O| = 4|C|$, and $|C \cap O| = 500$:

$$8,500 = |C| + 4|C| - 500.$$

Simplify the equation:

$$8,500 = 5|C| - 500 \implies 5|C| = 8,500 + 500 = 9,000.$$

Solve for $|C|$:

$$|C| = \frac{9,000}{5} = 1,800.$$

Conclusion: The number of students who like their core branch is:

$$\boxed{1,800}.$$

Quick Tip

In problems involving set theory, carefully apply union and intersection formulas.
Use given ratios and complements systematically to simplify calculations.

4. For positive non-zero real variables x and y , if

$$\ln\left(\frac{x+y}{2}\right) = \frac{1}{2} [\ln(x) + \ln(y)],$$

then the value of $\frac{x}{y} + \frac{y}{x}$ is:

- (A) 1
- (B) $\frac{1}{2}$
- (C) 2
- (D) 4

Correct Answer: (C) 2

Solution:

We are given the equation:

$$\ln\left(\frac{x+y}{2}\right) = \frac{1}{2} [\ln(x) + \ln(y)].$$

Step 1: Rewrite the equation using logarithmic properties The right-hand side can be rewritten using the logarithmic property $\ln(a) + \ln(b) = \ln(ab)$:

$$\frac{1}{2} [\ln(x) + \ln(y)] = \ln(\sqrt{xy}).$$

This simplifies the equation to:

$$\ln\left(\frac{x+y}{2}\right) = \ln(\sqrt{xy}).$$

Step 2: Remove the logarithms by exponentiating Exponentiating both sides removes the logarithms:

$$\frac{x+y}{2} = \sqrt{xy}.$$

Step 3: Manipulate the equation to find $\frac{x}{y} + \frac{y}{x}$ Square both sides to eliminate the square root:

$$\left(\frac{x+y}{2}\right)^2 = xy.$$

Expanding and simplifying:

$$\frac{(x+y)^2}{4} = xy \implies x^2 + y^2 + 2xy = 4xy.$$

Rearrange to isolate $x^2 + y^2$:

$$x^2 + y^2 = 2xy.$$

Now calculate $\frac{x}{y} + \frac{y}{x}$:

$$\frac{x}{y} + \frac{y}{x} = \frac{x^2 + y^2}{xy} = \frac{2xy}{xy} = 2.$$

Conclusion: The value of $\frac{x}{y} + \frac{y}{x}$ is:

2

Quick Tip

To solve logarithmic equations, simplify using properties like $\ln(a) + \ln(b) = \ln(ab)$ and $\ln(a^b) = b \ln(a)$. Exponentiate to remove logarithms and manipulate algebraically to find the desired result.

5. In the sequence 6, 9, 14, x , 30, 41, a possible value of x is:

- (A) 25
- (B) 21
- (C) 18
- (D) 20

Correct Answer: (B) 21

Solution:

Step 1: Analyze the sequence and identify a pattern The given sequence is:

$$6, 9, 14, x, 30, 41.$$

To find the missing value x , calculate the differences between consecutive terms:

$$9 - 6 = 3, \quad 14 - 9 = 5.$$

The differences form an increasing sequence: 3, 5, ... The next difference is expected to be 7.

Step 2: Calculate the missing value Using the pattern, the difference between 14 and x should be 7:

$$x - 14 = 7 \implies x = 14 + 7 = 21.$$

Step 3: Verify the pattern for subsequent terms Check if the differences continue to follow the sequence:

$$30 - 21 = 9, \quad 41 - 30 = 11.$$

The differences 7, 9, 11 confirm the pattern.

Conclusion: The missing value in the sequence is:

21

Quick Tip
To solve sequence problems, systematically calculate the differences or ratios between terms. Look for consistent patterns, such as arithmetic or geometric progressions, to determine the missing values.

Question 6-10 carry two marks each

6. Sequence the following sentences in a coherent passage.

P: This fortuitous geological event generated a colossal amount of energy and heat that resulted in the rocks rising to an average height of 4 km across the contact zone.

Q: Thus, the geophysicists tend to think of the Himalayas as an active geological event rather than as a static geological feature.

R: The natural process of the cooling of this massive edifice absorbed large quantities of atmospheric carbon dioxide, altering the earth's atmosphere and making it better suited for life.

S: Many millennia ago, a breakaway chunk of bedrock from the Antarctic Plate collided with the massive Eurasian Plate.

(A) QPSR

(B) QSPR

(C) SPRQ

(D) SRPQ

Correct Answer: (C) SPRQ

Solution:

To arrange the sentences into a coherent passage, we analyze their logical and chronological flow.

Step 1: Determine the starting point The passage begins with a significant historical event. Sentence **S** introduces the collision between the Antarctic Plate and the Eurasian Plate, making it the logical starting point.

Step 2: Identify the immediate consequence Sentence **P** follows logically, describing the geological consequences of the collision mentioned in **S**.

Step 3: Address subsequent outcomes Sentence **R** explains the natural cooling process and its environmental effects, which follow the events in **P**.

Step 4: Provide a concluding perspective Sentence **Q** concludes the passage by offering a modern interpretation of the geological activity, rounding off the narrative.

Step 5: Confirm the sequence The logical sequence of sentences is:

$$S \rightarrow P \rightarrow R \rightarrow Q.$$

Conclusion: The correct arrangement of sentences is:

SPRQ

Quick Tip

When arranging sentences, focus on identifying the introduction, cause-effect relationships, and the conclusion to form a coherent flow.

7. A person sold two different items at the same price. He made 10% profit in one item, and 10% loss in the other item. In selling these two items, the person made a total of:

- (A) 1% profit
- (B) 2% profit
- (C) 1% loss
- (D) 2% loss

Correct Answer: (C) 1% loss

Solution:

Solution:

To analyze the situation, let the selling price (SP) of each item be 100 units. We'll determine the profit or loss by comparing the overall cost price (CP) and SP using the percentage values.

Step 1: Analyze the profit and loss effect - For the first item sold at a 10% profit:

$$\text{Profit} = \frac{\text{Profit Percentage}}{100} \times \text{Cost Price}.$$

Rewriting for cost price (CP_1):

$$CP_1 = \frac{SP}{1 + \frac{\text{Profit Percentage}}{100}} = \frac{100}{1.1} \approx 90.91 \text{ units}.$$

- For the second item sold at a 10% loss:

$$\text{Loss} = \frac{\text{Loss Percentage}}{100} \times \text{Cost Price}.$$

Rewriting for cost price (CP_2):

$$CP_2 = \frac{SP}{1 - \frac{\text{Loss Percentage}}{100}} = \frac{100}{0.9} \approx 111.11 \text{ units.}$$

Step 2: Calculate total cost price and selling price Add the cost prices of the two items:

$$\text{Total CP} = CP_1 + CP_2 = 90.91 + 111.11 = 202.02 \text{ units.}$$

Add the selling prices:

$$\text{Total SP} = 100 + 100 = 200 \text{ units.}$$

Step 3: Determine the net result Calculate the loss:

$$\text{Loss} = \text{Total CP} - \text{Total SP} = 202.02 - 200 = 2.02 \text{ units.}$$

Calculate the loss percentage:

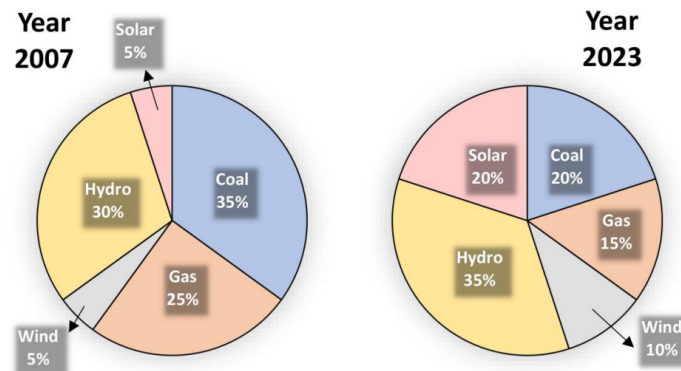
$$\text{Loss Percentage} = \frac{\text{Loss}}{\text{Total CP}} \times 100 = \frac{2.02}{202.02} \times 100 \approx 1\%.$$

Conclusion: The overall result is:

1% loss

Quick Tip
To calculate overall profit or loss, consider the individual effects of profit and loss percentages on the cost price. Add the total cost price and compare it with the total selling price.

8. The pie charts depict the shares of various power generation technologies in the total electricity generation of a country for the years 2007 and 2023.



The renewable sources of electricity generation consist of Hydro, Solar, and Wind. Assuming that the total electricity generated remains the same from 2007 to 2023, what is the percentage increase in the share of the renewable sources of electricity generation over this period?

- (1) 25%
- (2) 50%
- (3) 77.5%
- (4) 62.5%

Correct Answer: (4) 62.5%

Solution:

To determine the percentage increase in the share of renewable energy sources from 2007 to 2023, we break the problem into logical steps.

Step 1: Tabulate the renewable energy shares for 2007 and 2023 In 2007, the renewable energy shares are:

Hydro: 30%, Solar: 5%, Wind: 5%.

Total renewable energy share in 2007:

$$30\% + 5\% + 5\% = 40\%.$$

In 2023, the renewable energy shares are:

Hydro: 35%, Solar: 20%, Wind: 10%.

Total renewable energy share in 2023:

$$35\% + 20\% + 10\% = 65\%.$$

Step 2: Compute the absolute increase in renewable energy share The absolute increase is:

$$\text{Increase} = 65\% - 40\% = 25\%.$$

Step 3: Calculate the percentage increase relative to 2007 The percentage increase relative to the 2007 share is:

$$\text{Percentage Increase} = \frac{\text{Increase}}{\text{2007 Share}} \times 100 = \frac{25\%}{40\%} \times 100 = 62.5\%.$$

Conclusion: The percentage increase in the share of renewable energy sources from 2007 to 2023 is:

(4) 62.5%

Quick Tip

To compute percentage increases, first determine the absolute difference and then divide it by the initial value. Multiply by 100 to express it as a percentage.

9. A cube is to be cut into 8 pieces of equal size and shape. Here, each cut should be straight, and it should not stop till it reaches the other end of the cube. The minimum number of such cuts required is:

- (A) 3
- (B) 4
- (C) 7
- (D) 8

Correct Answer: (A) 3

Solution:

To determine how to divide a cube into 8 equal parts (smaller cubes), we analyze the process geometrically.

Step 1: Understand the goal We aim to create $8 = 2^3$ smaller cubes. This requires dividing the cube along three dimensions: length, breadth, and height. Each division doubles the number of smaller parts.

Step 2: Divide systematically along dimensions

- **First Cut:** A single cut along the length divides the cube into 2 equal parts.
- **Second Cut:** A cut along the breadth divides each part from the first cut into 2, resulting in $2 \times 2 = 4$ parts.
- **Third Cut:** A cut along the height divides each of the 4 parts into 2, producing $4 \times 2 = 8$ smaller cubes.

Step 3: Verify the total number of cuts Each cut along a different dimension doubles the number of parts, and 3 cuts are sufficient to achieve $8 = 2^3$ smaller cubes.

Conclusion: The minimum number of cuts required to divide the cube into 8 equal parts is:

3

Quick Tip

For dividing a cube into 2^n smaller cubes, n cuts are required, with one cut along each dimension (length, breadth, and height).

10. In the 4×4 array shown below, each cell of the first three rows has either a cross (X) or a number.

1	X	4	3
X	5	5	4
3	X	6	X
?	?	?	?

The number in a cell represents the count of the immediate neighboring cells (left, right, top, bottom, diagonals) NOT having a cross (X). Given that the last row has no crosses (X), the sum of the four numbers to be filled in the last row is:

(A) 11

- (B) 10
- (C) 12
- (D) 9

Correct Answer: (A) 11

Solution: The task is to find the numbers in the last row, where each number represents the count of adjacent cells that do not contain a cross (X). The last row itself contains no crosses (X), and its adjacent cells are in the third row.

For the first cell in the last row (?), the adjacent cells are: 3, X. Only 3 is not a cross, so this cell receives a count of 1.

For the second cell in the last row (?), the adjacent cells are: 3, X, 6. Both 3 and 6 are not crosses, so this cell gets a count of 2.

For the third cell in the last row (?), the adjacent cells are: X, 6, X. Only 6 is not a cross, so this cell receives a count of 1.

For the fourth cell in the last row (?), the adjacent cells are: 6, X. Only 6 is not a cross, so this cell gets a count of 1.

The numbers in the last row are: 2, 4, 3, 2. Their sum is:

$$2 + 4 + 3 + 2 = 11.$$

Final Answer:

11

Quick Tip
When solving array-based problems, be sure to carefully inspect the neighboring cells for each position and only count those that satisfy the given condition.

Question 11-35 carry one mark each

11. Consider a computer with a 4MHz processor. Its DMA controller can transfer 8 bytes in 1 cycle from a device to main memory through cycle stealing at regular inter-

vals. Which one of the following is the data transfer rate (in *bits per second*) of the DMA controller if 1% of the processor cycles are used for DMA?

- (A) 2, 56, 000
- (B) 3, 200
- (C) 25, 60, 000
- (D) 32, 000

Correct Answer: (C) 25, 60, 000

Solution:

To calculate the total data transfer rate using the DMA process, we approach the problem by analyzing the processor's performance and the DMA allocation step-by-step.

Step 1: Determine the total processor cycles per second The processor speed is given as:

$$4 \text{ MHz} = 4 \times 10^6 \text{ cycles per second.}$$

Step 2: Calculate the cycles allocated to DMA DMA uses 1% of the processor cycles. The cycles allocated to DMA per second are:

$$\text{DMA cycles} = \frac{1}{100} \times 4 \times 10^6 = 4 \times 10^4 \text{ cycles per second.}$$

Step 3: Determine the data transferred per DMA cycle Each DMA cycle transfers 8 bytes. Converting this to bits:

$$8 \text{ bytes} = 8 \times 8 = 64 \text{ bits.}$$

Step 4: Calculate the total data transfer rate The total data transfer rate is given by:

$$\text{Rate} = \text{DMA cycles} \times \text{Data per cycle.}$$

Substitute the values:

$$\text{Rate} = 4 \times 10^4 \times 64 = 25, 60, 000 \text{ bits per second.}$$

Conclusion: The data transfer rate using DMA is:

25, 60, 000 bits per second

Quick Tip

To compute the DMA transfer rate, identify the percentage of processor cycles allocated, calculate the cycles per second, and multiply by the data transferred per cycle.

12. Let p and q be the following propositions:

p : Fail grade can be given.

q : Student scores more than 50% marks.

Consider the statement: “Fail grade cannot be given when student scores more than 50% marks.”

Which one of the following is the CORRECT representation of the above statement in propositional logic?

(A) $q \rightarrow \neg p$

(B) $q \rightarrow p$

(C) $p \rightarrow q$

(D) $\neg p \rightarrow q$

Correct Answer: (A) $q \rightarrow \neg p$

Solution:

The problem involves interpreting the given statement and translating it into propositional logic. We approach this step by step.

Step 1: Understand the given statement The statement implies:

If a student scores more than 50% marks (q), then the fail grade (p) cannot be given ($\neg p$).

This relationship can be represented logically as:

$$q \rightarrow \neg p.$$

Step 2: Validate each option against the interpretation

- (A) $q \rightarrow \neg p$: This matches the interpretation of the statement and is correct.

- **(B)** $q \rightarrow p$: This implies that scoring $> 50\%$ leads to a fail grade, which contradicts the statement.
- **(C)** $p \rightarrow q$: This implies that failing (p) leads to scoring $> 50\%$, which is not equivalent to the given statement.
- **(D)** $\neg p \rightarrow q$: This implies that not failing ($\neg p$) leads to scoring $> 50\%$, which does not align with the original statement.

Step 3: Conclusion The correct representation of the statement is:

$$q \rightarrow \neg p.$$

Quick Tip

When translating statements into propositional logic, focus on the conditional relationships and correctly use logical symbols to represent the intended meaning.

13. Consider the following C program. Assume parameters to a function are evaluated from right to left.

```
#include <stdio.h>

int g(int p) { printf("%d", p); return p; }
int h(int q) { printf("%d", q); return q; }

void f(int x, int y) {
    g(x);
    h(y);
}

int main() {
    f(g(10), h(20));
}
```

Which one of the following options is the CORRECT output of the above C program?

- (A) 20101020
- (B) 10202010
- (C) 20102010
- (D) 10201020

Correct Answer: (A) 20101020

Solution:

To determine the output of the given C program, we analyze the sequence of function calls and the order of evaluation.

Step 1: Understand the function call sequence The `main()` function calls:

$$f(g(10), h(20))$$

In C, function arguments are evaluated **from right to left**. This means `h(20)` is evaluated before `g(10)`.

Step 2: Evaluate the parameters of `f()`

- **First: Evaluate `h(20)`** `h(20)` prints 20 and returns 20.
- **Second: Evaluate `g(10)`** `g(10)` prints 10 and returns 10.

The first part of the output is:

20 10.

Step 3: Execute the function `f()` The function `f()` is called with arguments 10 and 20. Within `f()`, the following occurs:

- `g(x)` is called with $x = 10$. This prints 10.
- `h(y)` is called with $y = 20$. This prints 20.

The second part of the output is:

10 20.

Step 4: Combine the output Concatenating both parts, the complete output is:

20101020

Conclusion: The total output of the program is:

20101020

Quick Tip

When analyzing the output of C programs, pay close attention to the order of parameter evaluation and the sequence of function calls.

14. The format of a single-precision floating-point number as per the IEEE 754 standard is:

Sign (1 bit)	Exponent (8 bits)	Mantissa (23 bits)
---------------------	--------------------------	---------------------------

Choose the largest floating-point number among the following options:

- (A) **Sign:** 0, **Exponent:** 0111 1111, **Mantissa:** 1111 1111 1111 1111 1111 111
(B) **Sign:** 0, **Exponent:** 1111 1110, **Mantissa:** 1111 1111 1111 1111 1111 111
(C) **Sign:** 0, **Exponent:** 1111 1111, **Mantissa:** 1111 1111 1111 1111 1111 111
(D) **Sign:** 0, **Exponent:** 0111 1111, **Mantissa:** 0000 0000 0000 0000 0000 000

Correct Answer: (B)

Solution:

To identify the largest floating-point number according to the IEEE 754 standard, we analyze the structure of floating-point numbers and evaluate the options.

Step 1: IEEE 754 Representation The value of a floating-point number is given by:

$$(-1)^{\text{Sign}} \times (1.\text{Mantissa}) \times 2^{\text{Exponent}-127}.$$

- The **Exponent** determines the range. Higher values of the exponent produce larger numbers.
- The **Mantissa** determines precision. A mantissa close to 1 (all bits set) yields the largest value for a given exponent.

Step 2: Evaluate each option

- **Option (A):** Exponent = 127, Mantissa = All 1's. Value:

$$(1.111...1) \times 2^{127-127} = 1.999...9 \times 2^0.$$

- **Option (B):** Exponent = 254, Mantissa = All 1's. Value:

$$(1.111...1) \times 2^{254-127} = 1.999...9 \times 2^{127}.$$

- **Option (C):** Exponent = 255, Mantissa = All 1's. This represents *NaN* (Not a Number), as 255 is a special case reserved for *NaN* and *Infinity*.

- **Option (D):** Exponent = 127, Mantissa = All 0's. Value:

$$(1.0) \times 2^{127-127} = 1 \times 2^0 = 1.$$

Step 3: Determine the largest value

- Option (B) has the largest exponent (254) among valid numbers, making it the largest.
- Option (C) is invalid (*NaN*), and Option (D) is the smallest due to a zero mantissa.

Conclusion: The largest floating-point number is:

(B)

Quick Tip

In IEEE 754 floating-point representation, the largest valid number has the maximum exponent (< 255) and a mantissa filled with 1's. Special cases like *NaN* and *Infinity* must be excluded.

15. Let $T(n)$ be the recurrence relation defined as follows:

$$T(0) = 1, \quad T(1) = 2, \quad T(n) = 5T(n-1) - 6T(n-2) \quad \text{for } n \geq 2$$

Which one of the following statements is TRUE?

- (A) $T(n) = \Theta(2^n)$
- (B) $T(n) = \Theta(n2^n)$
- (C) $T(n) = \Theta(3^n)$
- (D) $T(n) = \Theta(n3^n)$

Correct Answer: (A) $T(n) = \Theta(2^n)$

Solution:

To solve the recurrence relation:

$$T(n) = 5T(n-1) - 6T(n-2),$$

we analyze its behavior step-by-step.

Step 1: Derive the characteristic equation Assume a solution of the form $T(n) = r^n$. Substituting into the recurrence relation gives:

$$r^n = 5r^{n-1} - 6r^{n-2}.$$

Dividing through by r^{n-2} , we get the characteristic equation:

$$r^2 - 5r + 6 = 0.$$

Factoring yields:

$$(r-3)(r-2) = 0.$$

Thus, the roots are:

$$r_1 = 3, \quad r_2 = 2.$$

Step 2: Write the general solution The general solution for the recurrence relation is:

$$T(n) = C_1(3^n) + C_2(2^n),$$

where C_1 and C_2 are constants determined by the initial conditions.

Step 3: Apply initial conditions Given $T(0) = 1$ and $T(1) = 2$, substitute these into the general solution:

$$T(0) = C_1(3^0) + C_2(2^0) = C_1 + C_2 = 1, \tag{1}$$

$$T(1) = C_1(3^1) + C_2(2^1) = 3C_1 + 2C_2 = 2. \tag{2}$$

Solve equations (1) and (2) simultaneously:

$$C_1 + C_2 = 1,$$

$$3C_1 + 2C_2 = 2.$$

From the first equation:

$$C_2 = 1 - C_1.$$

Substitute into the second equation:

$$3C_1 + 2(1 - C_1) = 2 \implies 3C_1 + 2 - 2C_1 = 2 \implies C_1 = 0.$$

From $C_2 = 1 - C_1$, we find:

$$C_2 = 1.$$

Step 4: Simplify the solution Substituting $C_1 = 0$ and $C_2 = 1$ into the general solution gives:

$$T(n) = 2^n.$$

Step 5: Analyze asymptotic behavior As n grows, the solution is dominated by 2^n . Thus, the asymptotic complexity is:

$$T(n) = \Theta(2^n).$$

Conclusion: The asymptotic complexity of the recurrence relation is:

$$\boxed{\Theta(2^n)}.$$

Quick Tip

When solving recurrence relations, find the characteristic equation, solve for the roots, and use initial conditions to determine constants in the general solution.

16. Let $f(x)$ be a continuous function from $\mathbb{R}$ to $\mathbb{R}$ such that:

$$f(x) = 1 - f(2 - x)$$

Which one of the following options is the CORRECT value of $\int_0^2 f(x) dx$?

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

Correct Answer: (B) 1

Solution:

The given functional equation is:

$$f(x) + f(2 - x) = 1$$

Step 1: Split the integral The integral can be split into two parts:

$$\int_0^2 f(x) dx = \int_0^1 f(x) dx + \int_1^2 f(x) dx.$$

Step 2: Apply substitution for the second term In the second term, let $u = 2 - x$. Then:

$$du = -dx, \quad x = 1 \implies u = 1, \quad x = 2 \implies u = 0.$$

The integral becomes:

$$\int_1^2 f(x) dx = \int_1^0 f(2 - u)(-du) = \int_0^1 f(2 - u) du.$$

Step 3: Use the functional equation From $f(x) + f(2 - x) = 1$, we know:

$$f(2 - u) = 1 - f(u).$$

Substitute this into the integral:

$$\int_1^2 f(x) dx = \int_0^1 (1 - f(u)) du.$$

Step 4: Simplify the expression Substitute back into the original split integral:

$$\int_0^2 f(x) dx = \int_0^1 f(x) dx + \int_0^1 (1 - f(u)) du.$$

Simplify:

$$\int_0^2 f(x) dx = \int_0^1 f(x) dx + \int_0^1 1 du - \int_0^1 f(u) du.$$

Note that $\int_0^1 f(x) dx = \int_0^1 f(u) du$, so these terms cancel:

$$\int_0^2 f(x) dx = \int_0^1 1 du.$$

Step 5: Evaluate the remaining integral

$$\int_0^1 1 du = 1.$$

Conclusion: The value of the integral is:

$$\boxed{1}.$$

Quick Tip

Leverage symmetry and functional equations to simplify integrals. Substitution can help reduce terms and reveal cancellations.

17. Let A be the adjacency matrix of a simple undirected graph G . Suppose A is its own inverse. Which one of the following statements is always TRUE?

- (A) G is a cycle
- (B) G is a perfect matching
- (C) G is a complete graph
- (D) There is no such graph G

Correct Answer: (B) G is a perfect matching

Solution:

If A is the adjacency matrix of an undirected graph G , and it satisfies $A \cdot A = I$ (where I is the identity matrix), the following conclusions can be drawn:

Step 1: Implications of $A \cdot A = I$

- The diagonal elements of $A \cdot A = I$ indicate that no node in G connects back to itself.
- The off-diagonal zeros in $A \cdot A$ imply that there are no paths of length 2 between any two nodes in G . This means that each node is connected to exactly one other node.

Step 2: Interpretation of the graph G These conditions describe a perfect matching in G , where each node is paired with exactly one other node, and there are no additional edges or connections.

Step 3: Analyze other options

- **(A) G is a cycle:** A cycle graph would not satisfy $A \cdot A = I$ because there would be paths of length 2 between adjacent nodes in the cycle.
- **(C) G is a complete graph:** A complete graph would also not satisfy $A \cdot A = I$ because every node is connected to multiple other nodes.

- **(D) There is no such graph G :** This is incorrect because graphs with perfect matchings do exist and satisfy the given condition.

Conclusion: The graph G represents a perfect matching:

Perfect Matching

Quick Tip

When $A \cdot A = I$ for the adjacency matrix of a graph, each node is connected to exactly one other node, indicating a perfect matching.

18. When six unbiased dice are rolled simultaneously, the probability of getting all distinct numbers (i.e., 1, 2, 3, 4, 5, and 6) is

- (A) $\frac{1}{324}$
- (B) $\frac{5}{324}$
- (C) $\frac{7}{324}$
- (D) $\frac{11}{324}$

Correct Answer: (B) $\frac{5}{324}$

Solution:

To determine the probability of rolling six distinct numbers using six dice, we analyze the problem step-by-step using the concept of permutations and probabilities.

Step 1: Total possible outcomes Each die can show any one of 6 faces, and since there are 6 dice:

$$\text{Total Outcomes} = 6^6 = 46656.$$

Step 2: Favorable outcomes To roll six distinct numbers, all 6 faces of the dice must appear exactly once. The number of ways to arrange 6 distinct faces is given by the factorial:

$$\text{Favorable Outcomes} = 6! = 6 \cdot 5 \cdot 4 \cdot 3 \cdot 2 \cdot 1 = 720.$$

Step 3: Probability calculation The probability of rolling six distinct numbers is the ratio of favorable outcomes to total outcomes:

$$P = \frac{\text{Favorable Outcomes}}{\text{Total Outcomes}} = \frac{720}{46656}.$$

Simplify the fraction:

$$P = \frac{720}{46656} = \frac{5}{324}.$$

Conclusion: The probability of rolling six distinct numbers is:

$$\frac{5}{324}$$

Quick Tip

For probability problems involving dice, carefully consider whether the arrangement of outcomes matters and use factorials or combinations accordingly.

19. Once the DBMS informs the user that a transaction has been successfully completed, its effect should persist even if the system crashes before all its changes are reflected on disk. This property is called

- (A) durability
- (B) atomicity
- (C) consistency
- (D) isolation

Correct Answer: (A) durability

Solution:

The durability property ensures that once the database management system (DBMS) confirms a transaction as successfully completed, the changes made by that transaction are permanent. These changes persist even in the event of a system failure, making durability a crucial aspect of the ACID properties.

Step 1: Identify and evaluate the ACID properties

- **Atomicity:** Ensures a transaction is fully completed or fully rolled back. However, it does not guarantee that the changes persist after a crash.
- **Consistency:** Ensures that a database transitions from one valid state to another, preserving data integrity. This does not address persistence after a failure.

- **Isolation:** Ensures that transactions execute independently, without interference. This is unrelated to the permanent storage of changes.
- **Durability:** Ensures that once a transaction is committed, its changes are permanent and resilient to system crashes.

Step 2: Validate the correct property Only durability directly addresses the requirement that changes persist even after a system failure. Other ACID properties ensure correctness and reliability of transactions but do not focus on ensuring persistence.

Conclusion: The property that guarantees the permanence of committed changes is:

durability

Quick Tip

To understand transaction management, remember that durability ensures changes from committed transactions persist, even in the event of system failures.

20. In the context of owner and weak entity sets in the ER (Entity-Relationship) data model, which one of the following statements is TRUE?

- (A) The weak entity set **MUST** have total participation in the identifying relationship
- (B) The owner entity set **MUST** have total participation in the identifying relationship
- (C) Both weak and owner entity sets **MUST** have total participation in the identifying relationship
- (D) Neither weak entity set nor owner entity set **MUST** have total participation in the identifying relationship

Correct Answer: (A) The weak entity set **MUST** have total participation in the identifying relationship

Solution:

In the ER (Entity-Relationship) model, a weak entity set depends on an owner entity set for its existence. This dependency enforces the condition that every instance of a weak entity must be associated with at least one instance of the owner entity set through an identifying

relationship. This requirement translates to total participation of the weak entity set in the identifying relationship.

Step 1: Examine the given options

- **(A): The weak entity set must have total participation.** This is correct because every instance of a weak entity must participate in the identifying relationship with an owner entity to ensure its existence.
- **(B): The owner entity set must have total participation.** This is incorrect. The owner entity set does not require total participation in the identifying relationship; it may have partial participation.
- **(C): Both entity sets must have total participation.** This is incorrect. Only the weak entity set requires total participation, not the owner entity set.
- **(D): Neither entity set requires total participation.** This is incorrect. The weak entity set must have total participation to ensure the relationship is valid.

Conclusion: The correct statement is:

The weak entity set **MUST** have total participation.

Quick Tip

In the ER model, always ensure that a weak entity set participates fully in the identifying relationship with its owner entity set to maintain data integrity.

21. Consider the following two sets:

Set X		Set Y	
P.	Lexical Analyzer	1.	Abstract Syntax Tree
Q.	Syntax Analyzer	2.	Token
R.	Intermediate Code Generator	3.	Parse Tree
S.	Code Optimizer	4.	Constant Folding

Which one of the following options is the CORRECT match from Set X to Set Y?

- (A) $P - 4; Q - 1; R - 3; S - 2$
- (B) $P - 2; Q - 3; R - 1; S - 4$
- (C) $P - 2; Q - 1; R - 3; S - 4$
- (D) $P - 4; Q - 3; R - 2; S - 1$

Correct Answer: (B)

Solution:

To determine the correct mapping between Set X and Set Y, let us examine the functionality of each compiler component.

P. Lexical Analyzer: The lexical analyzer scans the input program and converts it into a sequence of tokens. Tokens represent keywords, identifiers, operators, etc.

$$P - 2 \quad (\text{Tokens}).$$

Q. Syntax Analyzer: The syntax analyzer uses the tokens provided by the lexical analyzer to construct a parse tree, representing the syntactical structure of the input program.

$$Q - 3 \quad (\text{Parse Tree}).$$

R. Intermediate Code Generator: The intermediate code generator produces a simplified and machine-independent representation of the input program, often represented as an abstract syntax tree (AST).

$$R - 1 \quad (\text{Abstract Syntax Tree (AST)}).$$

S. Code Optimizer: The code optimizer improves the intermediate code by performing operations like constant folding, loop unrolling, and dead code elimination to make it more efficient.

$$S - 4 \quad (\text{Constant Folding}).$$

Conclusion: The correct matching of Set X with Set Y is:

$$P - 2, Q - 3, R - 1, S - 4.$$

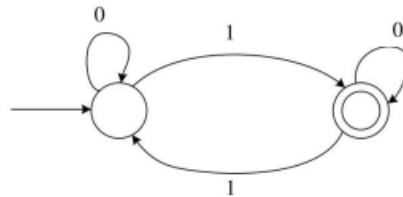
Final Answer:

(B)

Quick Tip

To match compiler phases with their functions, focus on the key outputs of each phase, such as tokens, parse trees, intermediate representations, and optimizations.

22. Which one of the following regular expressions is equivalent to the language accepted by the DFA given below?



- (A) $0^*1(0 + 10^*1)^*$
- (B) $0^*(10^*11)^*0^*$
- (C) $0^*1(010^*1)^*0^*$
- (D) $0(1 + 0^*10^*1)^*0^*$

Correct Answer: (A)

Solution:

The DFA in question consists of three states:

- q_0 : The start state.
- q_1 : A state reached after processing a 1.
- q_2 : The accepting state.

Step 1: Analyze the transitions

- From q_0 : On reading 1, the DFA moves to q_1 . Reading 0 keeps it in q_0 .
- From q_1 : On reading 0, it transitions to q_2 . On reading 1, it returns to q_0 .
- From q_2 : Reading 0 loops back to q_2 . Reading 1 returns to q_1 .

Step 2: Construct the regular expression The DFA accepts strings that satisfy the following conditions:

- The string may begin with any number of 0's: 0^* .
- A 1 must be encountered to leave q_0 : 1.
- After the initial 1, the DFA alternates between q_0 and q_1 with transitions like 0 or 10^*1 .
- Finally, the string may end with any number of 0's to remain in q_2 : 0^* .

Combining these, the regular expression is:

$$0^*1(0 + 10^*1)^*$$

Conclusion: The equivalent regular expression for the given DFA is:

$$0^*1(0 + 10^*1)^*$$

Quick Tip

To derive a regular expression from a DFA, trace the paths from the start state to the accepting state(s), accounting for all loops and transitions.

23. Node X has a TCP connection open to node Y. The packets from X to Y go through an intermediate IP router R. Ethernet switch S is the first switch on the network path between X and R. Consider a packet sent from X to Y over this connection.

Which of the following statements is/are TRUE about the destination IP and MAC addresses on this packet at the time it leaves X?

- (A) The destination IP address is the IP address of R.
- (B) The destination IP address is the IP address of Y.
- (C) The destination MAC address is the MAC address of S.
- (D) The destination MAC address is the MAC address of Y.

Correct Answer: (B)

Solution: In a TCP/IP packet, the destination **IP address** remains unchanged and always refers to the final destination node. In this scenario, the final destination is Y, meaning the destination IP address corresponds to Y. Thus, **Option (B)** is correct.

On the other hand, the **MAC address** changes at each hop and identifies the next device in

the communication path. When the packet departs from X , the destination MAC address is updated to the MAC address of R (the next-hop router) rather than S or Y .

Hence, the correct choice is **Option (B)**.

Final Answer:

(B)

Quick Tip

In a TCP/IP network, the IP address identifies the ultimate destination, while the MAC address is updated at each hop to specify the next device along the route.

24. Which of the following tasks is/are the responsibility/responsibilities of the memory management unit (MMU) in a system with paging-based memory management?

- (A) Allocate a new page table for a newly created process
- (B) Translate a virtual address to a physical address using the page table
- (C) Raise a trap when a virtual address is not found in the page table
- (D) Raise a trap when a process tries to write to a page marked with read-only permission in the page table

Correct Answer: (B), (C), (D)

Solution: **(A) Allocate a new page table for a newly created process:** Allocating a page table for a new process is generally the responsibility of the operating system (OS), not the Memory Management Unit (MMU). The OS handles the creation of the page table when a new process is initialized. Therefore, **Option (A)** is incorrect.

(B) Translate a virtual address to a physical address using the page table: This is a key function of the MMU. The MMU is responsible for translating virtual addresses into physical addresses by referencing the page table, which is why **Option (B)** is correct.

(C) Raise a trap when a virtual address is not found in the page table: When a virtual address is not found in the page table, the MMU triggers a page fault. This raises a trap, prompting the operating system to take appropriate actions, such as loading the page into memory. Hence, **Option (C)** is correct.

(D) Raise a trap when a process tries to write to a page marked with read-only permission in the page table: The MMU is responsible for enforcing memory protection. If a process attempts to write to a read-only page, the MMU raises a trap to prevent the violation, signaling an error. Therefore, **Option (D)** is correct.

Final Answer:

(B), (C), (D)

Quick Tip

The MMU is responsible for address translation, memory protection, and raising traps for invalid memory accesses, while the operating system manages tasks like page table allocation.

25. Consider a process P running on a CPU. Which one or more of the following events will always trigger a context switch by the OS that results in process P moving to a non-running state (e.g., ready, blocked)?

- (A) P makes a blocking system call to read a block of data from the disk.
- (B) P tries to access a page that is in the swap space, triggering a page fault.
- (C) An interrupt is raised by the disk to deliver data requested by some other process.
- (D) A timer interrupt is raised by the hardware.

Correct Answer: (A), (B)

Solution: When process P makes a blocking system call, such as waiting for data from a disk, it cannot continue until the I/O operation is complete. As a result, the operating system places P in the *blocked* state, which triggers a context switch. Hence, **Option (A)** is correct. If P attempts to access a page that is not currently in memory but is present in the swap space, a page fault occurs. In this case, the operating system retrieves the page from disk and places P in the *blocked* state while waiting for the I/O operation to finish. This triggers a context switch as well. Therefore, **Option (B)** is correct.

An interrupt caused by a disk operation to deliver data to another process is not directly related to P 's execution. The interrupt is handled by the OS, but it does not necessarily lead to

a context switch for P . Therefore, **Option (C)** is incorrect.

A timer interrupt, triggered by hardware, does not always force a context switch. The OS scheduler evaluates whether P should continue executing or be preempted. As a result, **Option (D)** is also incorrect.

Final Answer:

(A), (B)

Quick Tip

Context switches are triggered by blocking system calls and page faults, where the process must wait for I/O operations to complete.

26. Which of the following file organizations is/are I/O efficient for the scan operation in DBMS?

- (A) Sorted
- (B) Heap
- (C) Unclustered tree index
- (D) Unclustered hash index

Correct Answer: (A), (B)

Solution: In a scan operation, sequential data access reduces disk I/O, as fewer disk seeks are required. Sorted file organizations offer this type of sequential access, making them efficient for scan operations. Therefore, **Option (A)** is correct.

Heap files, though unordered, also allow sequential access for a full scan, where every block is read one after another. This sequential nature of access makes heap files I/O efficient for scan operations, meaning **Option (B)** is correct.

Unclustered tree indexes do not provide sequential access to data blocks. Since data blocks can be scattered across the disk, scan operations tend to involve random I/O, which is inefficient. Thus, **Option (C)** is incorrect.

Unclustered hash indexes are optimized for equality searches, not for scan operations. They

involve random I/O for accessing data blocks, making them inefficient for scans. Hence, **Option (D)** is incorrect.

Final Answer:

(A), (B)

Quick Tip

Both heap and sorted file organizations support efficient scans due to their sequential access, which minimizes random disk I/O operations.

27. Which of the following statements about the Two Phase Locking (2PL) protocol is/are TRUE?

- (A) 2PL permits only serializable schedules
- (B) With 2PL, a transaction always locks the data item being read or written just before every operation and always releases the lock just after the operation
- (C) With 2PL, once a lock is released on any data item inside a transaction, no more locks on any data item can be obtained inside that transaction
- (D) A deadlock is possible with 2PL

Correct Answer: (A), (C), (D)

Solution: The Two-Phase Locking (2PL) protocol ensures that transactions follow two distinct phases: 1. The growing phase (where locks are acquired), and 2. The shrinking phase (where locks are released). This protocol guarantees that once a lock is released, no new locks can be acquired.

Option (A): 2PL guarantees conflict-serializability, meaning only serializable schedules are allowed. This ensures that the transactions' final outcomes are consistent with a serial execution order. Thus, **Option (A)** is correct.

Option (B): This statement is incorrect because, in 2PL, locks are held until the shrinking phase. They are not acquired and released immediately for every operation; this would violate the protocol's two-phase nature. Hence, **Option (B)** is incorrect.

Option (C): Once a lock is released during the shrinking phase in 2PL, no new locks can be

acquired. This is a fundamental property of the protocol. Thus, **Option (C)** is correct.

Option (D): While 2PL ensures serializability, it does not prevent deadlocks. Deadlocks may occur if there are circular waiting conditions among transactions. Thus, **Option (D)** is correct.

Final Answer:

(A), (C), (D)

Quick Tip

2PL guarantees serializability, but deadlocks are still possible. Use techniques like timeout or wait-die to resolve deadlocks in 2PL.

28. Which of the following statements about IPv4 fragmentation is/are TRUE?

- (A) The fragmentation of an IP datagram is performed only at the source of the datagram.
- (B) The fragmentation of an IP datagram is performed at any IP router which finds that the size of the datagram to be transmitted exceeds the MTU.
- (C) The reassembly of fragments is performed only at the destination of the datagram.
- (D) The reassembly of fragments is performed at all intermediate routers along the path from the source to the destination.

Correct Answer: (B), (C)

Solution: IPv4 fragmentation happens when a datagram's size exceeds the Maximum Transmission Unit (MTU) of a network link.

Option (A): Incorrect. Fragmentation can take place at any router along the path, not exclusively at the source.

Option (B): Correct. Any IP router that encounters a datagram larger than the MTU will perform fragmentation to ensure it fits within the network's constraints.

Option (C): Correct. Reassembly of fragmented datagrams is performed only at the destination, where the original datagram is reconstructed.

Option (D): Incorrect. Intermediate routers do not reassemble fragments; they forward them as they are, without making changes to their structure.

Final Answer:

(B), (C)

Quick Tip

Fragmentation can occur at any router where the MTU is exceeded, but reassembly is always done only at the destination to reconstruct the original datagram.

29. Which of the following statements is/are FALSE?

- (A) An attribute grammar is a syntax-directed definition (SDD) in which the functions in the semantic rules have no side effects

- (B) The attributes in an L-attributed definition cannot always be evaluated in a depth-first order
- (C) Synthesized attributes can be evaluated by a bottom-up parser as the input is parsed
- (D) All L-attributed definitions based on LR(1) grammar can be evaluated using a bottom-up parsing strategy

Correct Answer: (B), (D)

Solution: Option (A): True. In attribute grammars, semantic rules must ensure that functions have no side effects, meaning they should not interfere with other computations.

Option (B): False. L-attributed definitions allow evaluation in a depth-first order because inherited attributes are passed along the parse tree edges, which can be computed in a single depth-first pass.

Option (C): True. Synthesized attributes are computed based on the values of child nodes, making them suitable for evaluation in a bottom-up parsing strategy.

Option (D): False. While L-attributed definitions can be evaluated with top-down parsing, they are not inherently designed for bottom-up parsing strategies.

Final Answer:

(B), (D)

Quick Tip

L-attributed grammars are optimized for top-down parsing, while synthesized attributes are naturally suited for bottom-up parsing.

30. For a Boolean variable x , which of the following statements is/are FALSE?

- (A) $x \cdot 1 = x$
- (B) $x + 1 = x$
- (C) $x \cdot x = 0$
- (D) $x + \bar{x} = 1$

Correct Answer: (B), (C)

Solution: Option (A): True. In Boolean algebra, the AND operation with 1 always results in the variable itself: $x \cdot 1 = x$.

Option (B): False. The OR operation with 1 always results in 1, regardless of the value of x : $x + 1 = 1$.

Option (C): False. The AND operation with the variable itself gives the variable: $x \cdot x = x$, not 0.

Option (D): True. The OR operation with the complement of a variable always results in 1: $x + \bar{x} = 1$.

Final Answer:

(B), (C)

Quick Tip

In Boolean algebra, key rules include $x + 1 = 1$, $x \cdot x = x$, and $x + \bar{x} = 1$. Review these fundamental identities for clarity.

31. An instruction format has the following structure:

Instruction Number: *Opcode destination reg, source reg-1, source reg-2*

Consider the following sequence of instructions to be executed in a pipelined processor:

I₁: DIV R3, R1, R2

I₂: SUB R5, R3, R4

I₃: ADD R3, R5, R6

I₄: MUL R7, R3, R8

Which of the following statements is/are TRUE?

(A) There is a RAW dependency on R3 between I₁ and I₂

(B) There is a WAR dependency on R3 between I₁ and I₃

(C) There is a RAW dependency on R3 between I₂ and I₃

(D) There is a WAW dependency on R3 between I₃ and I₄

Correct Answer: (A)

Solution: Option (A): True. Instruction I_1 writes to R3, and instruction I_2 reads from R3, forming a Read After Write (RAW) dependency between I_1 and I_2 .

Option (B): False. Instruction I_3 writes to R3, and I_1 also writes to R3, creating a Write After Write (WAW) dependency, not a Write After Read (WAR) dependency.

Option (C): False. Although instruction I_3 writes to R3, I_2 does not read from R3. Therefore, there is no RAW dependency between I_2 and I_3 .

Option (D): False. Both I_3 and I_4 write to R3, creating a Write After Write (WAW) dependency, but this dependency does not exist in the given instruction sequence.

Final Answer:

(A)

Quick Tip

To detect dependencies, examine the source and destination registers of consecutive instructions carefully.

32. Which of the following fields of an IP header is/are always modified by any router before it forwards the IP packet?

- (A) Source IP Address
- (B) Protocol
- (C) Time to Live (TTL)
- (D) Header Checksum

Correct Answer: (C), (D)

Solution: Option (A): The source IP address remains unchanged by routers during the packet's journey. It does not get modified. Therefore, **Option (A)** is incorrect.

Option (B): The protocol field, which specifies the higher-level protocol (e.g., TCP or UDP), is not altered by routers. Hence, **Option (B)** is incorrect.

Option (C): The Time to Live (TTL) field is decreased by 1 by each router as a mechanism to prevent infinite loops. Since this field is always modified by routers, **Option (C)** is correct.

Option (D): When the TTL field is modified, the Header Checksum must be recomputed by the router to maintain data integrity. Therefore, **Option (D)** is correct.

Final Answer:

(C), (D)

Quick Tip

The TTL field helps prevent routing loops, and any changes to the header require the checksum to be recalculated.

33. Consider the following C function definition.

```
int fX(char *a) {  
    char *b = a;  
    while(*b)  
        b++;  
    return b - a;  
}
```

Which of the following statements is/are TRUE?

- (A) The function call `fX("abcd")` will always return a value
- (B) Assuming a character array `c` is declared as `char c[] = "abcd"` in `main()`, the function call `fX(c)` will always return a value
- (C) The code of the function will not compile
- (D) Assuming a character pointer `c` is declared as `char *c = "abcd"` in `main()`, the function call `fX(c)` will always return a value

Correct Answer: (A), (B), (D)

Solution: Option (A): The string literal "abcd" is correctly formed and the function calculates the string's length by iterating until the null-terminating character. Hence, **Option (A)** is correct.

Option (B): The character array `c` is a valid input to the function, and the function will successfully calculate the length of the string. Therefore, **Option (B)** is correct.

Option (C): This statement is false. The code is valid and will compile without any errors. Hence, **Option (C)** is incorrect.

Option (D): The character pointer `c`, which points to the string literal "abcd", is a valid input, and the function correctly computes the string length. Hence, **Option (D)** is correct.

Final Answer:

(A), (B), (D)

Quick Tip

Both string literals and character arrays are valid inputs for functions that calculate string lengths, as long as they are properly null-terminated.

34. Let P be the partial order defined on the set $\{1, 2, 3, 4\}$ as follows:

$$P = \{(x, x) \mid x \in \{1, 2, 3, 4\}\} \cup \{(1, 2), (3, 2), (3, 4)\}$$

The number of total orders on $\{1, 2, 3, 4\}$ that contain P is _____.

Correct Answer: 5

Solution: A total order extends a partial order by ensuring that every pair of elements is comparable. Given the partial order P , we have the following relations:

$$1 \leq 2, \quad 3 \leq 2, \quad 3 \leq 4.$$

From these relations, the valid total orders that extend P are as follows:

1. $1 \leq 3 \leq 4 \leq 2$
2. $1 \leq 3 \leq 2 \leq 4$
3. $1 \leq 4 \leq 3 \leq 2$
4. $3 \leq 1 \leq 4 \leq 2$
5. $3 \leq 1 \leq 2 \leq 4$

Thus, there are 5 distinct total orders that extend P .

Final Answer:

5

Quick Tip

To find the number of total orders from a given partial order, examine the transitive closure and list all valid permutations that satisfy the given relations.

35. Let A be an array containing integer values. The distance of A is defined as the minimum number of elements in A that must be replaced with another integer so that the resulting array is sorted in non-decreasing order. The distance of the array $[2, 5, 3, 1, 4, 2, 6]$ is _____.

Correct Answer: 3

Solution: To find the distance, we must determine the length of the Longest Increasing Subsequence (LIS) of the array. The LIS represents the longest subset of elements that are already in sorted order.

The given array is $[2, 5, 3, 1, 4, 2, 6]$. The LIS for this array is:

$[2, 3, 4, 6]$ (length = 4).

The distance is then calculated as the difference between the length of the array and the length of the LIS:

$$\text{Distance} = \text{Length of array} - \text{Length of LIS} = 7 - 4 = 3.$$

Final Answer:

3

Quick Tip

To calculate the distance of an array, subtract the length of its Longest Increasing Subsequence (LIS) from the total number of elements.

Question 36-65 carry two mark each

36. What is the output of the following C program?

```
#include <stdio.h>

int main() {
    double a[2] = {20.0, 25.0}, *p, *q;
    p = a;
    q = p + 1;
    printf("%d,%d", (int)(q - p), (int)(*q - *p));
    return 0; }
```

- (A) 4, 8
- (B) 1, 5
- (C) 8, 5
- (D) 1, 8

Correct Answer: (B)

Solution: The program initializes an array $a[2] = \{20.0, 25.0\}$, along with two pointers p and q . The steps are as follows:

1. $p = a$: Pointer p now points to the first element of the array, which is $a[0] = 20.0$.
2. $q = p + 1$: Pointer q now points to the second element of the array, i.e., $a[1] = 25.0$.
3. The expression $q - p$ calculates the difference in pointer positions. Since q points to the second element and p to the first, we get $q - p = 1$.
4. The expression $*q - *p$ calculates the difference between the values pointed to by q and p . Thus:

$$*q - *p = a[1] - a[0] = 25.0 - 20.0 = 5.0.$$

5. The `printf` statement prints these values as integers:

$$\text{Output: } (int)(q - p) = 1, \quad (int)(*q - *p) = 5.$$

Final Answer:

(B)

Quick Tip

Pointer arithmetic on arrays calculates the difference in positions of elements, while dereferencing pointers gives the values stored at those positions.

37. Consider a single processor system with four processes A, B, C, and D, represented as given below, where for each process the first value is its arrival time, and the second value is its CPU burst time.

$$A(0, 10), B(2, 6), C(4, 3), D(6, 7).$$

Which one of the following options gives the average waiting times when preemptive Shortest Remaining Time First (SRTF) and Non-Preemptive Shortest Job First (NP-SJF) CPU scheduling algorithms are applied to the processes?

- (A) SRTF = 6, NP-SJF = 7
- (B) SRTF = 6, NP-SJF = 7.5
- (C) SRTF = 7, NP-SJF = 7.5
- (D) SRTF = 7, NP-SJF = 8.5

Correct Answer: (B)

Solution: Non-Preemptive SJF Scheduling:

Process	Arrival Time	Burst Time	Completion Time	Turnaround Time	Waiting Time
A	0	10	10	10	0
B	2	6	19	17	11
C	4	3	13	9	6
D	6	7	26	20	13

The overall average waiting time for NP-SJF is calculated as:

$$\text{Average Waiting Time} = \frac{0 + 11 + 6 + 13}{4} = \frac{30}{4} = 7.5.$$

Preemptive SRTF Scheduling:

Process	Arrival Time	Burst Time	Completion Time	Turnaround Time	Waiting Time
A	0	10	26	26	16
B	2	6	11	9	3
C	4	3	7	3	0
D	6	7	18	12	5

The overall average waiting time for SRTF is:

$$\text{Average Waiting Time} = \frac{16 + 3 + 0 + 5}{4} = \frac{24}{4} = 6.$$

Note:

- **Turnaround Time** is computed as: Completion Time – Arrival Time.
- **Waiting Time** is derived from: Turnaround Time – Burst Time.

Final Answer:

$$(2) \text{ SRTF} = 6, \text{ NP-SJF} = 7.5$$

Quick Tip

For CPU scheduling problems, use tables to determine Completion Time, Turnaround Time, and Waiting Time for each process. Be sure to apply the correct algorithm rules, such as preemption in SRTF and non-preemption in NP-SJF.

38. Which one of the following CIDR prefixes exactly represents the range of IP addresses 10.12.2.0 to 10.12.3.255?

- (A) 10.12.2.0/23
- (B) 10.12.2.0/24
- (C) 10.12.0.0/22
- (D) 10.12.2.0/22

Correct Answer: (A)

Solution: The given IP range is from 10.12.2.0 to 10.12.3.255.

The CIDR notation 10.12.2.0/23 corresponds to a subnet mask of 255.255.254.0, which includes the IP range:

10.12.2.0 (starting IP) to 10.12.3.255 (ending IP).

Let's examine the other options:

(B) 10.12.2.0/24: This only covers the range from 10.12.2.0 to 10.12.2.255, which is not the correct range.

(C) 10.12.0.0/22: This covers a broader range from 10.12.0.0 to 10.12.3.255, which exceeds the given range.

(D) 10.12.2.0/22: This option is incorrect as it overlaps with the IP range 10.12.4.x, which is outside the required range.

Final Answer:

(A)

Quick Tip

To determine the correct CIDR notation, calculate the subnet mask and check the range of IP addresses it encompasses.

39. You are given a set V of distinct integers. A binary search tree T is created by inserting all elements of V one by one, starting with an empty tree. The tree T follows the convention that, at each node, all values stored in the left subtree of the node are smaller than the value stored at the node. You are not aware of the sequence in which these values were inserted into T , and you do not have access to T .

Which one of the following statements is TRUE?

- (A) Inorder traversal of T can be determined from V
- (B) Root node of T can be determined from V
- (C) Preorder traversal of T can be determined from V
- (D) Postorder traversal of T can be determined from V

Correct Answer: (A)

Solution: In a Binary Search Tree (BST), performing an inorder traversal visits the nodes in ascending order of their values, irrespective of the order in which they were inserted. Since V is a set that contains all elements of T , sorting V will yield the inorder traversal of T . Hence, **Option (A)** is correct.

Option (B): The root node of T depends on the insertion sequence and cannot be determined solely from V . Hence, **Option (B)** is incorrect.

Option (C): The preorder traversal depends on the structure of T , which in turn is influenced by the sequence of insertions. Thus, it cannot be determined from V . Hence, **Option (C)** is incorrect.

Option (D): The postorder traversal also depends on the structure of T and the insertion order, so it cannot be derived from V . Hence, **Option (D)** is incorrect.

Final Answer:

(A)

Quick Tip

In a Binary Search Tree (BST), the inorder traversal always produces the elements in sorted order, regardless of the insertion sequence.

40. Consider the following context-free grammar where the start symbol is S and the set of terminals is $\{a, b, c, d\}$:

$$S \rightarrow AaAb \mid BbBa \quad A \rightarrow cS \mid \epsilon \quad B \rightarrow dS \mid \epsilon$$

The following is a partially-filled LL(1) parsing table.

	a	b	c	d	$\$$
S	$S \rightarrow AaAb$	$S \rightarrow BbBa$	(1)	(2)	
A	$A \rightarrow \epsilon$	(3)	$A \rightarrow cS$		
B	(4)	$B \rightarrow \epsilon$		$B \rightarrow dS$	

Which one of the following options represents the CORRECT combination for the numbered cells in the parsing table?

Note: In the options, "blank" denotes that the corresponding cell is empty.

- (A) (1) $S \rightarrow AaAb$, (2) $S \rightarrow BbBa$, (3) $A \rightarrow \epsilon$, (4) $B \rightarrow \epsilon$
 (B) (1) $S \rightarrow BbBa$, (2) $S \rightarrow AaAb$, (3) $A \rightarrow \epsilon$, (4) $B \rightarrow \epsilon$
 (C) (1) $S \rightarrow AaAb$, (2) $S \rightarrow BbBa$, (3) blank, (4) blank
 (D) (1) $S \rightarrow BbBa$, (2) $S \rightarrow AaAb$, (3) blank, (4) blank

Correct Answer: (A)

Solution: Given the grammar and the corresponding parsing table:

Cell (1): This cell corresponds to the production of S when the input begins with c . According to the grammar, $A \rightarrow cS$ and $S \rightarrow AaAb$, which leads to $S \rightarrow AaAb$ in this cell. Hence, in cell (1), we have $S \rightarrow AaAb$.

Cell (2): This cell corresponds to S when the input starts with d . From $B \rightarrow dS$ and $S \rightarrow BbBa$, this cell contains $S \rightarrow BbBa$. Therefore, in cell (2), we have $S \rightarrow BbBa$.

Cell (3): This cell corresponds to A when the input starts with b . According to the grammar, $A \rightarrow \epsilon$, so this cell contains $A \rightarrow \epsilon$. Hence, in cell (3), we have $A \rightarrow \epsilon$.

Cell (4): This cell corresponds to B when the input starts with a . From the grammar, $B \rightarrow \epsilon$, so this cell contains $B \rightarrow \epsilon$. Hence, in cell (4), we have $B \rightarrow \epsilon$.

The correct configuration of the cells is:

$$(1)S \rightarrow AaAb, (2)S \rightarrow BbBa, (3)A \rightarrow \epsilon, (4)B \rightarrow \epsilon.$$

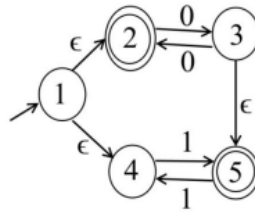
Final Answer:

(A)

Quick Tip

To populate an LL(1) parsing table, use the FIRST and FOLLOW sets of the grammar rules for each non-terminal.

41. Let M be the 5-state NFA with ϵ -transitions shown in the diagram below.



Which one of the following regular expressions represents the language accepted by M ?

- (A) $(00)^* + 1(11)^*$
- (B) $0^* + (1 + 0(00)^*)(11)^*$
- (C) $(00)^* + (1 + 0(00)^*)(11)^*$
- (D) $0^+ + 1(11)^* + 0(11)^*$

Correct Answer: (B)

Solution: The NFA M consists of five states (1, 2, 3, 4, 5) and includes ϵ -transitions. To determine the language accepted by M , we examine the transitions:

1. From state 1 to 2: The ϵ -transition allows for an immediate move to state 2 without consuming any input symbol.
2. From state 2 to 3: The transition on input 0 can loop on state 2 (repeated 0's), or transition to state 3 on a single 0. This represents $(00)^*$, i.e., any number of 0's.
3. From state 1 to 4: An ϵ -transition allows the move from state 1 directly to state 4 without consuming any input.
4. From state 4 to 5: The input 1 can be repeated (loop on 5) after consuming a single 1. This part corresponds to $(11)^*$, which represents any number of 1's.
5. Combining all possible paths: The final language encompasses:
 - 0^* , representing any sequence of 0's from state 2.
 - $(1 + 0(00)^*)(11)^*$, accounting for sequences starting with 1 or with 0 followed by $(00)^*$, and concluding with a sequence of $(11)^*$.

The regular expression for the language accepted by the NFA is:

$$0^* + (1 + 0(00)^*)(11)^*.$$

Final Answer:

(B)

Quick Tip

When dealing with NFAs containing ϵ -transitions, trace all possible paths from the start state to the accepting states, combining loops and concatenations to form the regular expression.

42. Consider an array X that contains n positive integers. A subarray of X is defined to be a sequence of array locations with consecutive indices.

The C code snippet given below has been written to compute the length of the longest subarray of X that contains at most two distinct integers. The code has two missing expressions labelled (P) and (Q) .

```
int first=0, second=0, len1=0, len2=0, maxlen=0;
for (int i=0; i < n; i++) {
    if (X[i] == first) {
        len2++; len1++;
    } else if (X[i] == second) {
        len2++;
        len1 = (P);
        second = first;
    } else {
        len2 = (Q);
        len1 = 1; second = first;
    }
    if (len2 > maxlen) {
        maxlen = len2;
    }
    first = X[i];
}
```

Which one of the following options gives the CORRECT missing expressions?

(A) $(P) \text{ len1} + 1, (Q) \text{ len2} + 1$

- (B) $(P) = 1, (Q) = \text{len1} + 1$
 (C) $(P) = 1, (Q) = \text{len2} + 1$
 (D) $(P) = \text{len2} + 1, (Q) = \text{len1} + 1$

Correct Answer: (B)

Solution: The objective is to keep track of the length of the longest subarray ending at index i that contains at most two distinct integers. The two variables are defined as follows:

1. len1 : Represents the length of the subarray ending at $X[i]$ where all elements are equal to $X[i]$.
2. len2 : Represents the length of the subarray ending at $X[i]$ with at most two distinct integers.

Analysis of (P): When $X[i] == \text{second}$, it indicates the start of a new sequence of identical integers. As a result, len1 is reset to 1. Therefore, $(P) = 1$.

Analysis of (Q): If $X[i]$ is neither equal to first nor second , we update len2 to reflect the subarray with the new two distinct integers. Since len1 corresponds to the length of the subarray with first , the updated value of len2 is $\text{len1} + 1$, where the new element $X[i]$ is added. Therefore, $(Q) = \text{len1} + 1$.

The correct values for (P) and (Q) are:

$$(P) = 1, \quad (Q) = \text{len1} + 1.$$

Final Answer:

(B)

Quick Tip

Make sure to track the roles of len1 and len2 accurately. len1 captures the contiguous subarray of identical elements, while len2 tracks the subarrays containing at most two distinct integers.

43. Consider the following expression: $x[i] = (p + r) * -s[i] + u/w$. **The following sequence shows the list of triples representing the given expression, with entries missing for triples (1), (3), and (6).**

(0)	+	p	r
(1)			
(2)	uminus	(1)	
(3)			
(4)	/	u	w
(5)	+	(3)	(4)
(6)			
(7)	=	(6)	$x[i]$

Which one of the following options fills in the missing entries **CORRECTLY**?

- (A) (1) = $\llbracket s[i] \rrbracket$, (3) = (0)(2), (6) = $\llbracket x[i] \rrbracket$
 (B) (1) = $\llbracket s[i] \rrbracket$, (3) = (0)(2), (6) = $\llbracket x(5) \rrbracket$
 (C) (1) = $\llbracket s[i] \rrbracket$, (3) = (0)(2), (6) = $\llbracket x(5) \rrbracket$
 (D) (1) = $s[i]$, (3) = (0)(2), (6) = $x[i]$

Correct Answer: (A)

Solution: To break down the expression $x[i] = (p + r) * -s[i] + u/w$ into intermediate triples:

1. Triple (0): The addition operation between p and r . This is already given as:

$$(0) : + pr.$$

2. Triple (1): The assignment of the value $-s[i]$, achieved using the unary minus (uminus) operator:

$$(1) := \llbracket s[i] \rrbracket.$$

3. Triple (2): Applying the unary minus to the result of (1). This is already given as:

$$(2) : \text{uminus } (1).$$

4. Triple (3): The multiplication of the results from (0) and (2):

$$(3) : * (0) (2).$$

5. Triple (4): The division operation between u and w , which is already given as:

$$(4) : / u w.$$

6. Triple (5): Adding the results from (3) and (4):

$$(5) : + (3) (4).$$

7. Triple (6): Assigning the result from (5) to $x[i]$:

$$(6) := [] x[i].$$

The correct missing entries are:

$$(1) := [] s[i], (3) : * (0) (2), (6) := [] x[i].$$

Final Answer:

(A)

Quick Tip

When generating three-address code, carefully follow the operator precedence to determine the order of intermediate steps and generate corresponding triples.

44. Let x and y be random variables, not necessarily independent, that take real values in the interval $[0, 1]$. Let $z = xy$ and let the mean values of x, y, z be $\bar{x}, \bar{y}, \bar{z}$, respectively. Which one of the following statements is TRUE?

(A) $\bar{z} = \bar{x}\bar{y}$

(B) $\bar{z} \leq \bar{x}\bar{y}$

(C) $\bar{z} \geq \bar{x}\bar{y}$

(D) $\bar{z} \leq \bar{x}$

Correct Answer: (D)

Solution: The random variables x and y lie within the interval $[0, 1]$, and we are given that $z = xy$. For any pair of real values $x, y \in [0, 1]$, the product xy will always be less than or equal to x since $y \leq 1$. Consequently, the expected value of z , denoted as $\bar{z}$, satisfies:

$$\bar{z} \leq \bar{x}.$$

This inequality holds universally, irrespective of whether x and y are dependent or independent.

Evaluating the other options:

Option (A): The equation $\bar{z} = \bar{x}\bar{y}$ holds only when x and y are independent, which is not guaranteed in this case. Therefore, (A) is incorrect.

Option (B): The inequality $\bar{z} \leq \bar{x}\bar{y}$ follows from the general result $E[xy] \leq E[x]E[y]$, which is weaker than the stricter condition $\bar{z} \leq \bar{x}$. Thus, (B) is not the strongest statement.

Option (C): The inequality $\bar{z} \geq \bar{x}\bar{y}$ is false, as $E[xy] \leq E[x]E[y]$ holds, contradicting this assertion.

Option (D): The inequality $\bar{z} \leq \bar{x}$ is always valid and represents the most restrictive and accurate statement. Thus, (D) is correct.

Final Answer:

(D)

Quick Tip

When dealing with random variables confined to bounded intervals, remember that their products are always limited by their individual ranges.

45. The relation schema, Person(pid, city), describes the city of residence for every person uniquely identified by pid. The following relational algebra operators are available: selection, projection, cross product, and rename.

To find the list of cities where at least 3 persons reside, using the above operators, the minimum number of cross product operations that must be used is

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

Correct Answer: (B)

Solution: The task is to identify cities where at least 3 people reside. The process can be broken down as follows:

1. Apply `selection` and `projection` operations to extract the relevant city values.
2. Use the `cross product` operation to generate all possible combinations of rows for each city, ensuring there are at least 3 distinct entries for each.

To confirm that each city has at least 3 people, a minimum of 2 `cross product` operations is needed.

Final Answer:

(B)

Quick Tip

The `cross product` operation in relational algebra combines rows from two relations, which can be useful for verifying cardinality constraints in queries.

46. Consider a multi-threaded program with two threads T_1 and T_2 . The threads share two semaphores: s_1 (initialized to 1) and s_2 (initialized to 0). The threads also share a global variable x (initialized to 0). The threads execute the code shown below.

<code>// code of T1</code>	<code>// code of T2</code>
<code>wait(s1);</code>	<code>wait(s1);</code>
<code>x = x+1;</code>	<code>x = x+1;</code>
<code>print(x);</code>	<code>print(x);</code>
<code>wait(s2);</code>	<code>signal(s2);</code>
<code>signal(s1);</code>	<code>signal(s1);</code>

Which of the following outcomes is/are possible when threads $T1$ and $T2$ execute concurrently?

- (A) $T1$ runs first and prints 1, $T2$ runs next and prints 2
- (B) $T2$ runs first and prints 1, $T1$ runs next and prints 2
- (C) $T1$ runs first and prints 1, $T2$ does not print anything (deadlock)
- (D) $T2$ runs first and prints 1, $T1$ does not print anything (deadlock)

Correct Answer: (B), (C)

Solution: Outcome (A): This scenario is not feasible because $T2$ cannot proceed after $T1$ without acquiring $s1$, but $s1$ is released only when $T1$ finishes. Hence, (A) is not possible.

Outcome (B): $T2$ may run first, acquire $s1$, increment x , and print 1. Then $T1$ acquires $s1$, increments x , and prints 2. This outcome is valid.

Outcome (C): $T1$ acquires $s1$, increments x , and prints 1. However, it waits on $s2$, which is never signaled by $T2$, leading to a deadlock. This scenario is possible.

Outcome (D): This scenario is impossible because $T2$ signals $s2$ after printing, which allows $T1$ to proceed. Therefore, (D) is not feasible.

Final Answer:

(B), (C)

Quick Tip

When analyzing semaphore synchronization, always check the order of signaling and acquiring resources to detect potential deadlocks and ensure correct execution flow.

47. Let A be an $n \times n$ matrix over the set of all real numbers $\mathbb{R}$. Let B be a matrix obtained from A by swapping two rows. Which of the following statements is/are TRUE?

- (A) The determinant of B is the negative of the determinant of A
- (B) If A is invertible, then B is also invertible
- (C) If A is symmetric, then B is also symmetric
- (D) If the trace of A is zero, then the trace of B is also zero

Correct Answer: (A), (B)

Solution: Option (A): Swapping two rows of a matrix results in a new matrix whose determinant is the negative of the original matrix's determinant. Therefore, (A) is correct.

Option (B): Swapping two rows in an invertible matrix does not affect its invertibility because the determinant remains non-zero, ensuring the matrix remains invertible. Hence, (B) is correct.

Option (C): Swapping rows of a symmetric matrix does not necessarily preserve symmetry, as the condition $A[i][j] = A[j][i]$ could be violated. Thus, (C) is incorrect.

Option (D): The trace of a matrix, being the sum of its diagonal elements, remains unaffected by row swaps. If the trace of A is zero, it will remain zero for matrix B , making (D) incorrect.

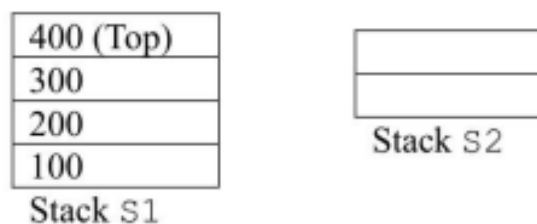
Final Answer:

(A), (B)

Quick Tip

Always consider how row operations influence key matrix properties like determinant, symmetry, and trace when analyzing matrix transformations.

48. Let $S1$ and $S2$ be two stacks. $S1$ has a capacity of 4 elements, and $S2$ has a capacity of 2 elements. $S1$ already has 4 elements: 100, 200, 300, and 400, whereas $S2$ is empty, as shown below.



Only the following three operations are available:

PushToS2 : Pop the top element from $S1$ and push it on $S2$.

PushToS1 : Pop the top element from $S2$ and push it on $S1$.

GenerateOutput : Pop the top element from $S1$ and output it to the user.

Note: - The pop operation is not allowed on an empty stack.

- The push operation is not allowed on a full stack.

Which of the following output sequences can be generated by using the above operations?

(A) 100, 200, 400, 300

(B) 200, 300, 400, 100

(C) 400, 200, 100, 300

(D) 300, 200, 400, 100

Correct Answer: (B), (C), (D)

Solution: The given operations offer flexibility for rearranging elements between $S1$ and $S2$.

Analyzing the sequences:

1. Sequence (A): The sequence 100, 200, 400, 300 is not possible because 400 must appear after 300 in $S1$. Reordering them would violate the sequence order, so (A) is not feasible.

2. Sequence (B): The sequence 200, 300, 400, 100 can be generated as follows:

- Output 200 and 300 from $S1$.
- Push 400 to $S2$ and output 100.
- Finally, retrieve 400 from $S2$ and output it. Thus, (B) is feasible.

3. Sequence (C): The sequence 400, 200, 100, 300 can be generated through:

- Output 400 from $S1$.
- Push 300 and 200 to $S2$.
- Output 100 from $S1$.
- Retrieve 300 from $S2$ and output it. Therefore, (C) is feasible.

4. Sequence (D): The sequence 300, 200, 400, 100 can be generated through:

- Push 400 to $S2$.
- Output 300 and 200 from $S1$.
- Retrieve 400 from $S2$ and output it.
- Lastly, output 100. Thus, (D) is feasible.

Final Answer:

(B), (C), (D)

Quick Tip

For stack-based problems, simulate the process step by step to verify the feasibility of each sequence and check the order of operations.

49. Which of the following is/are EQUAL to 224 in radix-5 (i.e., base-5) notation?

- (A) 64 in radix-10
- (B) 100 in radix-8
- (C) 50 in radix-16
- (D) 121 in radix-7

Correct Answer: (A), (B), (D)

Solution: The number 224 in base-5 can be expressed in decimal (base-10) by the following calculation:

$$224_5 = 2 \cdot 5^2 + 2 \cdot 5^1 + 4 \cdot 5^0 = 50 + 10 + 4 = 64.$$

Now, let's verify each option:

Option (A): The decimal equivalent of 224_5 is 64, which matches the given value. Therefore, (A) is correct.

Option (B): Convert 100_8 (base-8) to decimal:

$$100_8 = 1 \cdot 8^2 + 0 \cdot 8^1 + 0 \cdot 8^0 = 64.$$

This also equals 64, matching 224_5 . Thus, (B) is correct.

Option (C): Convert 50_{16} (base-16) to decimal:

$$50_{16} = 5 \cdot 16^1 + 0 \cdot 16^0 = 80.$$

This result differs from 64, so (C) is incorrect.

Option (D): Convert 121_7 (base-7) to decimal:

$$121_7 = 1 \cdot 7^2 + 2 \cdot 7^1 + 1 \cdot 7^0 = 49 + 14 + 1 = 64.$$

This matches 224_5 . Hence, **(D)** is correct.

Final Answer:

(A), (B), (D)

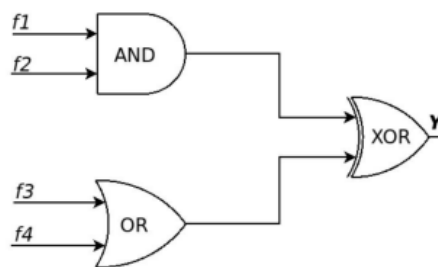
Quick Tip

To convert a number from any base to decimal, multiply each digit by the base raised to the power of its position, and sum the results.

50. Consider 4-variable functions f_1, f_2, f_3, f_4 expressed in sum-of-minterms form as given below.

$$f_1 = \Sigma(0, 2, 3, 5, 7, 8, 11, 13) \quad f_2 = \Sigma(1, 3, 5, 7, 11, 13, 15) \quad f_3 = \Sigma(0, 1, 4, 11) \quad f_4 = \Sigma(0, 2, 6, 13)$$

The circuit is represented as:



With respect to the circuit given above, which of the following options is/are CORRECT?

- (A) $Y = \Sigma(0, 1, 2, 11, 13)$
- (B) $Y = \Pi(3, 4, 5, 6, 7, 8, 9, 10, 12, 14, 15)$
- (C) $Y = \Sigma(0, 1, 2, 3, 4, 5, 6, 7)$
- (D) $Y = \Pi(8, 9, 10, 11, 12, 13, 14, 15)$

Correct Answer: (C), (D)

Solution: Step 1: Determine the result of $f_1 \wedge f_2$:

$$f_1 \wedge f_2 = \Sigma(3, 5, 7, 11, 13).$$

Step 2: Determine the result of $f_3 \vee f_4$:

$$f_3 \vee f_4 = \Sigma(0, 1, 2, 4, 6, 11, 13).$$

Step 3: Perform the XOR operation on the results of $f_1 \wedge f_2$ and $f_3 \vee f_4$: The XOR of two functions gives the minterms that are unique to one function but not shared between them:

$$Y = \Sigma(0, 1, 2, 3, 4, 5, 6, 7).$$

Analysis of the Options:

- Option (A): Incorrect, as Y does not match $\Sigma(0, 1, 2, 11, 13)$.
- Option (B): Incorrect, as $\Pi(3, 4, 5, 6, 7, 8, 9, 10, 12, 14, 15)$ does not align with Y .
- Option (C): Correct, because $Y = \Sigma(0, 1, 2, 3, 4, 5, 6, 7)$.
- Option (D): Correct, since $Y = \Sigma(0, 1, 2, 3, 4, 5, 6, 7)$, its complement is:

$$Y = \Pi(8, 9, 10, 11, 12, 13, 14, 15).$$

Final Answer:

(C), (D)

Quick Tip

When analyzing logical circuits, systematically evaluate each step of the operation and apply XOR to identify the unique minterms between the two functions.

51. Let G be an undirected connected graph in which every edge has a positive integer weight. Suppose that every spanning tree in G has even weight. Which of the following statements is/are TRUE for every such graph G ?

- (A) All edges in G have even weight
- (B) All edges in G have even weight **OR** all edges in G have odd weight
- (C) In each cycle C in G , all edges in C have even weight
- (D) In each cycle C in G , either all edges in C have even weight **OR** all edges in C have odd weight

Correct Answer: (D)

Solution: If every spanning tree in G has an even weight, this imposes specific constraints on the structure of the graph G . Let's analyze the options:

Option (A): Incorrect. Not every edge in G must have even weight; some edges can have odd weights, as long as the overall weight of the spanning trees remains even.

Option (B): Incorrect. G can include both even and odd weighted edges.

Option (C): Incorrect. It is not required that all edges in a cycle C have even weights.

Option (D): Correct. For every cycle C in G , the parity of the edges (whether even or odd) must be consistent across the cycle to ensure that the weight of every spanning tree remains even.

Final Answer:

(D)

Quick Tip

When studying spanning trees, consider how the parity of edge weights in cycles affects the overall weight of spanning trees.

52. Consider a context-free grammar G with the following 3 rules:

$$S \rightarrow aS, \quad S \rightarrow aSbS, \quad S \rightarrow c$$

Let $w \in L(G)$. Let $n_a(w), n_b(w), n_c(w)$ denote the number of times a, b, c occur in w , respectively. Which of the following statements is/are TRUE?

(A) $n_a(w) > n_b(w)$

(B) $n_a(w) > n_c(w) - 2$

(C) $n_c(w) = n_b(w) + 1$

(D) $n_c(w) = n_b(w) * 2$

Correct Answer: (B), (C)

Solution: The grammar G produces strings of the form:

$$w = a^p cb^q,$$

where each derivation of c introduces one b and one a . Therefore, the relationships between $n_a(w)$, $n_b(w)$, and $n_c(w)$ can be examined:

Option (A): Incorrect. Although $n_a(w)$ is typically greater than $n_b(w)$, there is no certainty, as both $n_a(w)$ and $n_b(w)$ depend on the specific derivation.

Option (B): Correct. Since the derivation adds at least two a 's for every c , the inequality $n_a(w) > n_c(w) - 2$ holds true.

Option (C): Correct. Each derivation of $S \rightarrow aSbS$ introduces one b for each c , meaning $n_c(w) = n_b(w) + 1$.

Option (D): Incorrect. The number of c 's is not necessarily twice the number of b 's.

Final Answer:

(B), (C)

Quick Tip

When working with context-free grammars, explore how terminal symbols relate to one another based on derivation structures.

53. Consider a disk with the following specifications: rotation speed of 6000 RPM, average seek time of 5 milliseconds, 500 sectors/track, 512-byte sectors. A file has content stored in 3000 sectors located randomly on the disk. Assuming average rotational latency, the total time (in seconds, rounded off to 2 decimal places) to read the entire file from the disk is ____.

Correct Answer: 29.5 to 30.5

Solution: The total time required to read the file includes:

- ****Seek time**:** The average seek time is 5 ms.
- ****Rotational latency**:** Given the rotation speed of 6000 RPM, the time for one complete rotation is:

$$\text{Time per rotation} = \frac{60 \text{ sec}}{6000} = 0.01 \text{ sec.}$$

The average rotational latency is half the time of one rotation:

$$\text{Average rotational latency} = \frac{0.01}{2} = 0.005 \text{ sec} = 5 \text{ ms.}$$

- ****Data transfer time per sector****: Each track contains 500 sectors, and the time to transfer one sector is:

$$\text{Time per sector} = \frac{\text{Time per rotation}}{\text{Sectors per track}} = \frac{0.01}{500} = 0.00002 \text{ sec} = 20 \mu\text{s}.$$

- ****Total time per sector****: The total time to access one sector is the sum of seek time, rotational latency, and transfer time:

$$\text{Time per sector} = 5 \text{ ms} + 5 \text{ ms} + 20 \mu\text{s} = 10.02 \text{ ms}.$$

- ****Total time for 3000 sectors****:

$$\text{Total time} = 3000 \times 10.02 \text{ ms} = 30060 \text{ ms} = 30.06 \text{ sec}.$$

Final Answer:

29.5 to 30.5 sec

Quick Tip

To determine the disk access time, add the seek time, rotational latency, and data transfer time.

54. Consider a TCP connection operating at a point of time with the congestion window of size 12 MSS (Maximum Segment Size), when a timeout occurs due to packet loss. Assuming that all the segments transmitted in the next two RTTs (Round Trip Time) are acknowledged correctly, the congestion window size (in MSS) during the third RTT will be ____.

Correct Answer: 4

Solution: Upon a timeout in TCP, the congestion control mechanism reduces the congestion window to 1 MSS and transitions into the slow-start phase. In this phase:

During the first Round-Trip Time (RTT), the window size doubles from 1 → 2 MSS.

In the second RTT, the window size doubles again from 2 → 4 MSS.

Thus, at the beginning of the third RTT, the congestion window size will be 4 MSS.

Final Answer:

4 MSS

Quick Tip

In TCP, after a timeout, the congestion window grows exponentially during the slow-start phase.

55. Consider an Ethernet segment with a transmission speed of 10^8 bits/sec and a maximum segment length of 500 meters. If the speed of propagation of the signal in the medium is 2×10^8 meters/sec, then the minimum frame size (in bits) required for collision detection is ____.

Correct Answer: 500

Solution: The round-trip time for a signal to propagate across the segment is calculated as:

$$\text{Time} = 2 \times \frac{\text{Distance}}{\text{Propagation speed}} = 2 \times \frac{500}{2 \times 10^8} = 5 \mu\text{s}.$$

In order to detect collisions, the frame transmission time must be at least equal to the round-trip time. The frame transmission time is given by:

$$\text{Frame transmission time} = \frac{\text{Frame size}}{\text{Transmission speed}}.$$

Setting the frame transmission time equal to the round-trip time:

$$\frac{\text{Frame size}}{10^8} = 5 \mu\text{s} \implies \text{Frame size} = 10^8 \times 5 \times 10^{-6} = 500 \text{ bits}.$$

Final Answer:

500 bits

Quick Tip

For Ethernet, collision detection requires the frame size to be sufficient to account for the round-trip propagation delay.

56. A functional dependency $F : X \rightarrow Y$ is termed as a useful functional dependency if and only if it satisfies all the following three conditions:

- X is not the empty set.
- Y is not the empty set.
- Intersection of X and Y is the empty set.

For a relation R with 4 attributes, the total number of possible useful functional dependencies is ____.

Correct Answer: 50

Solution: Consider a relation R with four attributes $\{A, B, C, D\}$. We are tasked with calculating the total number of valid functional dependencies $X \rightarrow Y$, where:

- X represents a non-empty subset of the attributes.
- Y is a non-empty subset of the attributes.
- X and Y must be disjoint ($X \cap Y = \emptyset$).

Step 1: Total number of subsets of R The set $R = \{A, B, C, D\}$ has $2^4 = 16$ subsets, which include:

$$\{\emptyset, \{A\}, \{B\}, \{C\}, \{D\}, \{A, B\}, \{A, C\}, \dots, \{A, B, C, D\}\}.$$

Excluding the empty set, we have 15 non-empty subsets.

Step 2: Counting valid combinations of X and Y Since $X \cap Y = \emptyset$, the set Y must be selected from the remaining attributes in R that are not in X . For each fixed subset X , the number of available attributes for Y is $4 - |X|$, where $|X|$ is the size of subset X .

For every non-empty subset X (15 possible choices), the number of non-empty subsets of the remaining attributes (for Y) is:

$$2^{4-|X|} - 1.$$

Step 3: Summing over all valid combinations To find the total number of useful functional dependencies, we sum over all non-empty subsets X :

$$\text{Total number of functional dependencies} = \sum_{k=1}^4 \binom{4}{k} \cdot (2^{4-k} - 1),$$

where $\binom{4}{k}$ represents the number of ways to choose k attributes for X .

Step 4: Calculate each term For $k = 1, 2, 3, 4$, we calculate the following:

- For $k = 1$:

$$\binom{4}{1} \cdot (2^{4-1} - 1) = 4 \cdot (8 - 1) = 4 \cdot 7 = 28.$$

- For $k = 2$:

$$\binom{4}{2} \cdot (2^{4-2} - 1) = 6 \cdot (4 - 1) = 6 \cdot 3 = 18.$$

- For $k = 3$:

$$\binom{4}{3} \cdot (2^{4-3} - 1) = 4 \cdot (2 - 1) = 4 \cdot 1 = 4.$$

- For $k = 4$:

$$\binom{4}{4} \cdot (2^{4-4} - 1) = 1 \cdot (1 - 1) = 0.$$

Step 5: Summing the results The total number of useful functional dependencies is:

$$28 + 18 + 4 + 0 = 50.$$

Final Answer:

50

Quick Tip

To count useful functional dependencies, ensure X and Y are disjoint, non-empty subsets of the attributes.

57. A processor with 16 general-purpose registers uses a 32-bit instruction format. The instruction format consists of an opcode field, an addressing mode field, two register operand fields, and a 16-bit scalar field. If 8 addressing modes are to be supported, the maximum number of unique opcodes possible for every addressing mode is ____.

Correct Answer: 32

Solution: The 32-bit instruction format is organized as follows:

1. ****16 general-purpose registers****: Each register requires 4 bits ($2^4 = 16$).
2. ****Two register operand fields****: Each operand field is 4 bits, so the total is 8 bits.

3. **Addressing mode field**: 3 bits are needed to represent 8 addressing modes ($2^3 = 8$).
4. **Scalar field**: This requires 16 bits.

The remaining bits are allocated to the **opcode field**:

$$\text{Opcode bits} = 32 - (8 + 3 + 16) = 5 \text{ bits.}$$

The total number of unique opcodes for each addressing mode is:

$$2^{\text{Opcode bits}} = 2^5 = 32.$$

Final Answer:

32

Quick Tip

To calculate the opcode size, subtract the bit requirements of all other fields from the total instruction size.

58. A non-pipelined instruction execution unit operating at 2 GHz takes an average of 6 cycles to execute an instruction of a program P . The unit is then redesigned to operate on a 5-stage pipeline at 2 GHz. Assume that the ideal throughput of the pipelined unit is 1 instruction per cycle. In the execution of program P , 20% instructions incur an average of 2 cycles stall due to data hazards and 20% instructions incur an average of 3 cycles stall due to control hazards. The speedup (rounded off to one decimal place) obtained by the pipelined design over the non-pipelined design is ____.

Correct Answer: 2.9 to 3.1

Solution: Step 1: Define throughput for a pipelined processor.

Throughput for a pipelined processor refers to the number of instructions that are processed by the pipeline per unit time. In this scenario, the unit of time is one clock cycle, which lasts 0.5 ns. Each stage of the pipeline completes its operation in one clock cycle (ignoring register delays and clock skew).

Step 2: Define CPI for both pipelined and non-pipelined processors.

CPI (Clock Cycles Per Instruction) represents the average number of clock cycles required to

process one instruction. For a non-pipelined processor, the instruction takes 6 clock cycles to complete. In contrast, for a pipelined processor, the average number of clock cycles per instruction is reduced to 1. Therefore, we have:

$$\text{CPI (non-pipelined)} = 6, \quad \text{CPI (pipelined)} = 1.$$

Step 3: Define the speedup of the pipelined processor.

The speedup of the pipelined processor relative to the non-pipelined processor is determined by the formula:

$$\text{Speedup} = \frac{\text{CPI (non-pipelined)}}{\text{Ideal CPI (pipelined)} + \text{Pipeline stall clock cycles}}.$$

The pipeline stall clock cycles are included to account for delays caused by pipeline stalls.

Step 4: Calculate the speedup.

In the case of a program with a total instruction count of IC , 20% of the instructions experience stalls for 2 cycles, while another 20% of the instructions experience stalls for 3 cycles. Therefore, the effective CPI for the pipelined processor is:

$$\text{Effective CPI (pipelined)} = 1 + (20\% \times 2) + (20\% \times 3).$$

Substituting the values, we get:

$$\text{Effective CPI (pipelined)} = 1 + 0.4 + 0.6 = 2.$$

Now, the speedup is calculated as:

$$\text{Speedup} = \frac{\text{CPI (non-pipelined)}}{\text{Effective CPI (pipelined)}} = \frac{6}{2} = 3.$$

Step 5: Conclude the result.

The pipelined processor is 3.0 times faster than the non-pipelined processor.

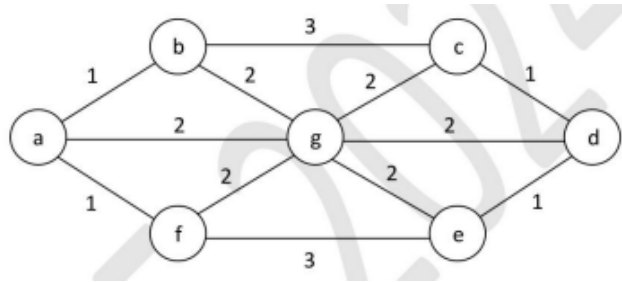
Final Answer:

The pipelined processor is 3.0 times faster than the non-pipelined processor.

Quick Tip

When calculating speedup, ensure to account for hazards and stalls that occur in pipelined systems.

59. The number of distinct minimum-weight spanning trees of the following graph is



Correct Answer: 9

Solution: Step 1: Analyze the graph and find the minimum weight.

The graph consists of 7 edges with the following weights: 2, 2, 2, 3, 3, 3, 1. To form a minimum spanning tree (MST), we need to minimize the sum of the edge weights.

Step 2: Apply Kruskal's algorithm.

Kruskal's algorithm works as follows:

- Select the edge with weight 1 (since it's the smallest and the only one with this weight).
- Choose three edges of weight 2. These edges form a cycle, meaning multiple selections are possible.
- Choose one edge of weight 3 to complete the MST.

Step 3: Count the distinct combinations.

- There are $\binom{3}{2} = 3$ ways to select two edges of weight 2.
- For each of these choices, there are 3 ways to select one edge of weight 3.

Thus, the total number of distinct MSTs is:

$$3 \times 3 = 9.$$

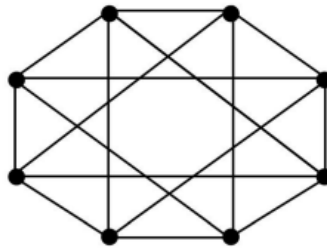
Final Answer:

9

Quick Tip

When determining the number of distinct MSTs, consider the different ways of selecting edges of the same weight while ensuring the MST properties are maintained.

60. The chromatic number of a graph is the minimum number of colours used in a proper colouring of the graph. The chromatic number of the following graph is



Correct Answer: 2

Solution: Step 1: Understanding the chromatic number.

The chromatic number of a graph is the minimum number of colours needed to assign to its vertices such that no two adjacent vertices share the same colour.

Step 2: Examine the structure of the graph.

The graph is bipartite, meaning its vertices can be split into two distinct sets with the following properties:

- Every edge in the graph connects a vertex from one set to a vertex from the other set.
- No edges exist between vertices within the same set.

Bipartite graphs always require just two colours for vertex colouring.

Step 3: Confirm the bipartite nature of the graph.

The graph consists of two disjoint vertex sets:

Set 1: Outer vertices. Set 2: Inner vertices.

Each edge links a vertex from Set 1 to Set 2, confirming the bipartite structure. Therefore, the chromatic number of the graph is 2.

Final Answer:

2

Quick Tip

The chromatic number of any bipartite graph is always 2, as it can be divided into two independent sets.

61. A processor uses a 32-bit instruction format and supports byte-addressable memory access. The ISA of the processor has 150 distinct instructions. The instructions are equally divided into two types, namely R-type and I-type, whose formats are shown below.

R-type Instruction Format:

OPCODE	UNUSED	DST Register	SRC Register1	SRC Register2
--------	--------	--------------	---------------	---------------

I-type Instruction Format:

OPCODE	DST Register	SRC Register	# Immediate Value/Address
--------	--------------	--------------	---------------------------

In the OPCODE, 1 bit is used to distinguish between I-type and R-type instructions, and the remaining bits indicate the operation. The processor has 50 architectural registers, and all register fields in the instructions are of equal size.

Let X be the number of bits used to encode the UNUSED field, Y be the number of bits used to encode the OPCODE field, and Z be the number of bits used to encode the immediate value/address field. The value of $X + 2Y + Z$ is ____.

Correct Answer: 34

Solution: Step 1: Analyze the instruction format.

Each instruction is 32 bits in total. The field breakdown is as follows:

• **R-type instruction fields:**

$$\text{OPCODE} + \text{UNUSED} + \text{DST Register} + \text{SRC Register1} + \text{SRC Register2} = 32 \text{ bits.}$$

- **I-type instruction fields:**

$$\text{OPCODE} + \text{DST Register} + \text{SRC Register} + \text{Immediate Value/Address} = 32 \text{ bits.}$$

Step 2: Calculate the bits required for registers.

The processor has 50 registers, and the number of bits needed to uniquely identify each register is:

$$\text{Register bits} = \lceil \log_2(50) \rceil = 6 \text{ bits.}$$

Therefore, each register field (DST Register, SRC Register1, SRC Register2) requires 6 bits.

Step 3: Determine the number of bits for the OPCODE field.

The instruction set has 150 distinct instructions, divided equally into R-type and I-type. With 1 bit used to distinguish instruction types, the remaining bits represent the operation:

$$\text{OPCODE bits} = \lceil \log_2(150) \rceil = 8 \text{ bits.}$$

Step 4: Calculate the unused field size (R-type) and immediate field size (I-type).

- **R-type:** The UNUSED field in the R-type instruction is:

$$X = 32 - (\text{OPCODE} + 3 \times \text{Register bits}) = 32 - (8 + 3 \times 6) = 6 \text{ bits.}$$

- **I-type:** The Immediate Value/Address field in the I-type instruction is:

$$Z = 32 - (\text{OPCODE} + 2 \times \text{Register bits}) = 32 - (8 + 2 \times 6) = 12 \text{ bits.}$$

Step 5: Compute $X + 2Y + Z$.

Substituting the values for X , Y , and Z :

$$X + 2Y + Z = 6 + 2 \times 8 + 12 = 34.$$

Final Answer:

34

Quick Tip

Ensure that the size of the OPCODE field aligns with the total number of instructions and the instruction format fits within the 32-bit limit.

62. Let L_1 be the language represented by the regular expression $b^*ab^*(ab^*ab^*)^*$ and $L_2 = \{w \in (a+b)^* \mid |w| \leq 4\}$, where $|w|$ denotes the length of string w . The number of strings in L_2 which are also in L_1 is ____.

Correct Answer: 15

Solution: Step 1: Define the set of symbols in the language.

Let the language alphabet be:

$$\Sigma = \{a, b\} \implies |\Sigma| = 2.$$

Step 2: Define the regular language L_1 .

The regular language L_1 consists of strings where the number of a 's is odd.

Step 3: Examine strings with an odd number of a 's.

For any string of length n over the alphabet $\{a, b\}$, half of the total strings will have an odd number of a 's. This can be expressed as:

$$\frac{2^n}{2},$$

where 2^n is the total number of possible strings of length n , based on the two symbols a and b .

Step 4: Calculate the total number of strings in L_1 for lengths up to 4.

We calculate half the total number of strings for all lengths from 0 to 4:

$$\frac{2^0}{2} + \frac{2^1}{2} + \frac{2^2}{2} + \frac{2^3}{2} + \frac{2^4}{2}.$$

Simplify each term:

$$0 + 1 + 2 + 4 + 8 = 15.$$

Final Answer: The total number of strings in L_1 (up to length 4) is:

$$\boxed{15}.$$

Quick Tip
To count strings with specific properties (e.g., an odd number of a certain symbol), calculate the total possible strings and use combinatorics to isolate the ones that satisfy the condition.

63. Let Z_n be the group of integers $\{0, 1, 2, \dots, n-1\}$ with addition modulo n as the group operation. The number of elements in the group $Z_2 \times Z_3 \times Z_4$ that are their own inverses is ____.

Correct Answer: 4

Solution: Step 1: Characterizing an element as its own inverse.

An element x in a group is said to be its own inverse if:

$$x + x \equiv 0 \pmod{n}.$$

This can be rewritten as:

$$2x \equiv 0 \pmod{n}.$$

Step 2: Evaluate each group.

- Z_2 : The group consists of $\{0, 1\}$.

$$2x \equiv 0 \pmod{2} \implies x = 0, 1.$$

- Z_3 : The group consists of $\{0, 1, 2\}$.

$$2x \equiv 0 \pmod{3} \implies x = 0.$$

- Z_4 : The group consists of $\{0, 1, 2, 3\}$.

$$2x \equiv 0 \pmod{4} \implies x = 0, 2.$$

Step 3: Calculate the total number of elements in $Z_2 \times Z_3 \times Z_4$. The size of $Z_2 \times Z_3 \times Z_4$ is:

$$|Z_2 \times Z_3 \times Z_4| = 2 \times 3 \times 4 = 24.$$

For an element $(x_2, x_3, x_4) \in Z_2 \times Z_3 \times Z_4$ to be its own inverse, we have:

- $x_2 \in \{0, 1\}$ (2 possibilities).
- $x_3 \in \{0\}$ (1 possibility).
- $x_4 \in \{0, 2\}$ (2 possibilities).

Thus, the total number of such elements is:

$$2 \times 1 \times 2 = 4.$$

Final Answer:

4

Quick Tip

To identify elements that are their own inverses, solve the congruence $2x \equiv 0 \pmod{\text{ulo}}$ the group order for each component.

64. Consider a 32-bit system with a 4 KB page size and page table entries of size 4 bytes each. Assume $1 \text{ KB} = 2^{10}$ bytes. The OS uses a 2-level page table for memory management, with the page table containing an outer page directory and an inner page table. The OS allocates a page for the outer page directory upon process creation. The OS uses demand paging when allocating memory for the inner page table, i.e., a page of the inner page table is allocated only if it contains at least one valid page table entry. An active process in this system accesses 2000 unique pages during its execution, and none of the pages are swapped out to disk. After it completes the page accesses, let X denote the minimum and Y denote the maximum number of pages across the two levels of the page table of the process. The value of $X + Y$ is ____.

Correct Answer: 1028

Solution: Given:

- Virtual address space (VA) = 32 bits
- Size of a page (PS) = 4 KB
- Number of pages allocated to the process (N) = 2000 pages
- Size of a page table entry (PTE) = 4 B

The problem states that the operating system assigns one page for the outer page directory during the creation of the process.

Step 1: Determine the number of bits required for the outer page table.

To represent the entries in the outer page table, the required number of bits is:

$$\log \left(\frac{PS}{PTE} \right) = \log \left(\frac{2^{12}}{2^2} \right) = \log(2^{10}) = 10 \text{ bits.}$$

Step 2: Determine the number of bits needed for the inner page table.

The bits needed to represent entries in the inner page table are:

$$32 \text{ (total VA bits)} - 10 \text{ (outer page table bits)} - 12 \text{ (page offset bits)} = 10 \text{ bits.}$$

Step 3: Calculate the value of X .

Each inner page table can hold 2^{10} entries. The minimum number of page tables (PT) required to store 2000 pages is:

$$\lceil \frac{2000}{2^{10}} \rceil = 2 \text{ } PT.$$

Thus:

$$X = 2 + 1 \text{ (for the outer page table)} = 3.$$

Step 4: Determine the value of Y .

To utilize the maximum capacity of a 2-level page table, at least one page table entry (PTE) for the 2000 pages must be placed in every page of the inner page table. This leads to:

$$Y = 2^{10} + 1 \text{ (for the outer page table)} = 1025.$$

Step 5: Calculate $X + Y$.

The sum of $X + Y$ is:

$$X + Y = 3 + 1025 = 1028.$$

Final Answer:

$$\boxed{1028}.$$

Quick Tip

When working with multi-level page tables, calculate the number of entries per inner table and consider both the minimum and maximum page allocation scenarios.

65. Consider the following augmented grammar, which is to be parsed with an SLR parser. The set of terminals is $\{a, b, c, d, \#, @\}$.

$$S' \rightarrow S \quad S \rightarrow SS \mid Aa \mid bAc \mid BC \mid bBa \quad A \rightarrow d\# \quad B \rightarrow @$$

Let $I_0 = \text{CLOSURE}(\{S' \rightarrow \cdot S\})$. The number of items in the set $\text{GOTO}(I_0, S)$ is -----.

Correct Answer: 9

Solution: Step 1: Examine the provided DFA diagram.

The DFA illustrates the elements present in the closure of I_0 along with the transitions defined by the GOTO function for the symbol S .

Step 2: Determine the items in the closure of I_0 .

The closure of I_0 consists of the following items:

- $S \rightarrow \cdot S$
- $S \rightarrow \cdot SS$
- $S \rightarrow \cdot Aa$
- $S \rightarrow \cdot bAc$
- $S \rightarrow \cdot Bc$
- $S \rightarrow \cdot bAb$
- $A \rightarrow \cdot d\#$
- $B \rightarrow \cdot @$

Step 3: Apply the GOTO operation for the symbol S .

Performing the GOTO function on S results in the following set of transitions:

- $S \rightarrow S \cdot$
- $S \rightarrow S \cdot S$
- $S \rightarrow SS \cdot$
- $S \rightarrow Aa \cdot$

- $S \rightarrow bAc$.
- $S \rightarrow Bc$.
- $S \rightarrow bAb$.
- $A \rightarrow d\#$.
- $B \rightarrow @$.

Step 4: Count the total number of items.

The GOTO function applied to the closure of I_0 for the symbol S results in 9 distinct items:

$S \rightarrow S$, $S \rightarrow S \cdot S$, $S \rightarrow SS$, $S \rightarrow Aa$, $S \rightarrow bAc$, $S \rightarrow Bc$, $S \rightarrow bAb$, $A \rightarrow d\#$, $B \rightarrow @$.

Final Answer: From the analysis of the DFA, it is concluded that:

GOTO(closure(I_0), S) consists of 9 items.

Quick Tip

To determine the number of items in a GOTO set, always evaluate the closure operation and ensure that all relevant production rules for the corresponding grammar symbol are included.