

AP ECET 2025 Mining Engineering Question Paper with Solutions

Time Allowed :3 Hour	Maximum Marks :200	Total Questions :200
----------------------	--------------------	----------------------

General Instructions

Read the following instructions very carefully and strictly follow them:

1. The test is of 3 hour duration.
2. The question paper consists of 200 questions. The maximum marks are 200.
3. The question paper is divided into four sections - Mathematics, Physics, Chemistry and Mining Engineering.

Mathematics

1. Order of the matrix $\begin{bmatrix} 1 & 6 \\ 2 & 0 \\ 7 & -1 \end{bmatrix}$ is

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

Correct Answer: (B) 3×2

Solution:

Step 1: Understanding the Question:

The question asks for the "order" of a given matrix.

The order of a matrix is a way to describe its dimensions, given as the number of rows by the number of columns.

Step 2: Detailed Explanation:

First, we count the number of horizontal lines, which are the rows of the matrix.

The given matrix is:

$$\begin{bmatrix} 1 & 6 \\ 2 & 0 \\ 7 & -1 \end{bmatrix}$$

Row 1 is $[1 \ 6]$.

Row 2 is $[2 \ 0]$.

Row 3 is $[7 \ -1]$.

So, there are 3 rows.

Next, we count the number of vertical lines, which are the columns of the matrix.

Column 1 is $\begin{bmatrix} 1 \\ 2 \\ 7 \end{bmatrix}$.

Column 2 is $\begin{bmatrix} 6 \\ 0 \\ -1 \end{bmatrix}$.

So, there are 2 columns.

The order is always written in the format (number of rows) $\times$ (number of columns).

Therefore, the order of this matrix is 3×2 .

Step 3: Final Answer:

The matrix has 3 rows and 2 columns, so its order is 3×2 . This corresponds to option (B).

💡 Quick Tip

To easily remember the format for the order of a matrix, think of "RC" - Rows first, then Columns. This simple mnemonic can help avoid confusion during exams.

2. If two rows (or columns) of a determinant of order 3 are identical then the value of determinant is

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

Correct Answer: (A) 0

Solution:

Step 1: Understanding the Question:

The question asks for the value of a determinant when two of its rows or two of its columns are identical. This is a question about a fundamental property of determinants.

Step 2: Detailed Explanation:

One of the key properties of determinants is as follows:

Property: If any two rows or any two columns of a determinant are identical, then the value of the determinant is zero.

Reasoning:

Let's consider a determinant Δ .

Another property of determinants states that if we interchange any two rows (or columns), the sign of the determinant changes.

Suppose rows R_i and R_j of the determinant are identical.

If we interchange these two identical rows, the determinant itself does not change.

So, the new determinant is still Δ .

However, according to the interchange property, the new determinant should be $-\Delta$.

Therefore, we have the equation:

$$\Delta = -\Delta$$

$$2\Delta = 0$$

$$\Delta = 0$$

This proves that the value of the determinant must be zero.

Step 3: Final Answer:

Based on the properties of determinants, if two rows or columns are identical, the determinant's value is 0. This corresponds to option (A).

💡 Quick Tip

This is a crucial property to memorize. In an exam, if you are asked to evaluate a determinant, first quickly scan the rows and columns. If you spot any identical or proportional rows/columns, you can immediately state the answer is 0 without performing any calculations.

3. Co-factor of -4 in $\begin{vmatrix} 1 & 2 & 3 \\ -4 & 3 & 6 \\ 2 & -7 & 9 \end{vmatrix}$ is

- (A) 3
- (B) 11
- (C) 39
- (D) -39

Correct Answer: (D) -39

Solution:

Step 1: Understanding the Question:

The question asks for the cofactor of a specific element, -4, in a given 3x3 determinant.

Step 2: Key Formula or Approach:

The cofactor of an element a_{ij} (located in the i-th row and j-th column) is given by the formula:

$$C_{ij} = (-1)^{i+j} M_{ij}$$

where M_{ij} is the minor of the element a_{ij} . The minor is the determinant of the submatrix formed by removing the i -th row and j -th column.

Step 3: Detailed Explanation:

First, we locate the element -4 in the determinant. It is in the 2nd row and 1st column. So, $i = 2$ and $j = 1$.

The element is $a_{21} = -4$.

Next, we calculate the minor M_{21} . We do this by removing the 2nd row and 1st column from the original matrix:

$$\text{Original determinant: } \begin{vmatrix} 1 & 2 & 3 \\ -4 & 3 & 6 \\ 2 & -7 & 9 \end{vmatrix}$$

$$\text{Submatrix for } M_{21}: \begin{vmatrix} 2 & 3 \\ -7 & 9 \end{vmatrix}$$

Now, we calculate the determinant of this 2x2 submatrix:

$$M_{21} = (2)(9) - (3)(-7) = 18 - (-21) = 18 + 21 = 39$$

Finally, we calculate the cofactor C_{21} using the formula:

$$C_{21} = (-1)^{2+1} M_{21}$$

$$C_{21} = (-1)^3 \times 39$$

$$C_{21} = -1 \times 39 = -39$$

Step 4: Final Answer:

The cofactor of the element -4 is -39. This corresponds to option (D).

💡 Quick Tip

The sign of the cofactor, given by $(-1)^{i+j}$, follows a checkerboard pattern:

$\begin{pmatrix} + & - & + \\ - & + & - \\ + & - & + \end{pmatrix}$. For the element at position (2,1), the sign is negative. You can quickly determine the sign and then just calculate the minor.

4. The Matrix $\begin{bmatrix} a & h & g \\ h & b & f \\ g & f & c \end{bmatrix}$ is

- (A) skew symmetric
- (B) symmetric
- (C) symmetric if a=b
- (D) Skew symmetric if b=c

Correct Answer: (B) symmetric

Solution:

Step 1: Understanding the Question:

The question asks to classify the given matrix based on its properties. The options suggest we need to check for symmetry or skew-symmetry.

Step 2: Key Formula or Approach:

A square matrix A is called **symmetric** if it is equal to its transpose, i.e., $A = A^T$. This means that the element in the i-th row and j-th column is equal to the element in the j-th row and i-th column, i.e., $a_{ij} = a_{ji}$ for all i and j.

A square matrix A is called **skew-symmetric** if $A = -A^T$, which means $a_{ij} = -a_{ji}$. This implies all diagonal elements must be zero.

Step 3: Detailed Explanation:

Let the given matrix be A:

$$A = \begin{bmatrix} a & h & g \\ h & b & f \\ g & f & c \end{bmatrix}$$

Let's check the condition for a symmetric matrix, $a_{ij} = a_{ji}$.

- Element at (1,2) is $a_{12} = h$. Element at (2,1) is $a_{21} = h$. So, $a_{12} = a_{21}$.

- Element at (1,3) is $a_{13} = g$. Element at (3,1) is $a_{31} = g$. So, $a_{13} = a_{31}$.

- Element at (2,3) is $a_{23} = f$. Element at (3,2) is $a_{32} = f$. So, $a_{23} = a_{32}$.

Since the condition $a_{ij} = a_{ji}$ holds for all elements, the matrix is symmetric.

Alternatively, we can find the transpose of A, A^T , by interchanging rows and columns:

$$A^T = \begin{bmatrix} a & h & g \\ h & b & f \\ g & f & c \end{bmatrix}$$

We can see that $A = A^T$. Therefore, the matrix is symmetric.

Step 4: Final Answer:

The given matrix satisfies the condition for a symmetric matrix. This corresponds to option (B).

 Quick Tip

A symmetric matrix is visually symmetric about its main diagonal (from top-left to bottom-right). You can quickly check if the elements are mirrored across this diagonal. This matrix is a standard representation of a symmetric matrix used in various fields like physics and geometry.

5. If $A = \begin{bmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{bmatrix}$, then $(A^{-1}) =$

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

Correct Answer: (A) A

Solution:

Step 1: Understanding the Question:

The question asks for the inverse of the given matrix A.

Step 2: Key Formula or Approach:

The inverse of a matrix A, denoted A^{-1} , is a matrix such that $AA^{-1} = A^{-1}A = I$, where I is the identity matrix.

A simple way to check if a matrix is its own inverse is to multiply the matrix by itself. If the result is the identity matrix ($A^2 = I$), then it means $A = A^{-1}$. Such a matrix is called an involutory matrix.

Step 3: Detailed Explanation:

Let's compute the product $A \times A$:

$$A^2 = A \times A = \begin{bmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{bmatrix}$$

We perform the matrix multiplication row by column:

The element in the 1st row, 1st column is $(0)(0) + (0)(0) + (1)(1) = 1$.

The element in the 1st row, 2nd column is $(0)(0) + (0)(1) + (1)(0) = 0$.

The element in the 1st row, 3rd column is $(0)(1) + (0)(0) + (1)(0) = 0$.

The element in the 2nd row, 1st column is $(0)(0) + (1)(0) + (0)(1) = 0$.

The element in the 2nd row, 2nd column is $(0)(0) + (1)(1) + (0)(0) = 1$.

The element in the 2nd row, 3rd column is $(0)(1) + (1)(0) + (0)(0) = 0$.

The element in the 3rd row, 1st column is $(1)(0) + (0)(0) + (0)(1) = 0$.

The element in the 3rd row, 2nd column is $(1)(0) + (0)(1) + (0)(0) = 0$.

The element in the 3rd row, 3rd column is $(1)(1) + (0)(0) + (0)(0) = 1$.

So, the resulting matrix is:

$$A^2 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = I$$

Since $A^2 = I$, we can pre-multiply by A^{-1} :

$$A^{-1}(AA) = A^{-1}I$$

$$(A^{-1}A)A = A^{-1}$$

$$IA = A^{-1}$$

$$A = A^{-1}$$

Step 4: Final Answer:

The inverse of matrix A is the matrix A itself. This corresponds to option (A).

💡 Quick Tip

The given matrix is a permutation matrix. Specifically, it swaps the first and third elements. Applying the same swap twice returns the original arrangement. Recognizing special types of matrices like permutation matrices, identity matrices, or zero matrices can often lead to a much faster solution.

6. If $\deg f(x) \geq \deg g(x)$, then the rational fraction $f(x)/g(x)$ is called

- (A) Polynomial
- (B) Proper fraction
- (C) Improper fraction
- (D) irrational fraction

Correct Answer: (C) Improper fraction

Solution:

Step 1: Understanding the Question:

This is a definition-based question from algebra concerning rational fractions (expressions involving polynomials). We need to identify the correct term for a rational fraction where the degree of the numerator is greater than or equal to the degree of the denominator.

Step 2: Detailed Explanation:

A **rational fraction** is an expression of the form $\frac{P(x)}{Q(x)}$, where $P(x)$ and $Q(x)$ are polynomials

and $Q(x) \neq 0$.

There are two main types of rational fractions based on the degrees of the polynomials:

1. **Proper Fraction:** A rational fraction $\frac{f(x)}{g(x)}$ is called a proper fraction if the degree of the numerator polynomial $f(x)$ is **less than** the degree of the denominator polynomial $g(x)$.

Example: $\frac{x+1}{x^2+2}$. (Degree of numerator is 1, degree of denominator is 2).

2. **Improper Fraction:** A rational fraction $\frac{f(x)}{g(x)}$ is called an improper fraction if the degree of the numerator polynomial $f(x)$ is **greater than or equal to** the degree of the denominator polynomial $g(x)$.

Example: $\frac{x^3+1}{x^2+2}$ or $\frac{x^2+1}{x^2+2}$.

The question states that $\deg f(x) \geq \deg g(x)$, which directly matches the definition of an improper fraction.

Step 3: Final Answer:

The rational fraction is called an improper fraction. This corresponds to option (C).

💡 Quick Tip

This terminology is analogous to numerical fractions. A fraction like $5/3$ is "improper" because the numerator is larger than the denominator. A fraction like $2/3$ is "proper". The same concept applies to the degrees of polynomials in rational expressions.

7. If $\frac{3x}{x^2+x-2} = \frac{A}{x+2} + \frac{B}{x-1}$ then the ordered pair (A, B) is

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

Correct Answer: (D) (2, 1)

Solution:

Step 1: Understanding the Question:

The question asks us to find the values of the constants A and B in the partial fraction decomposition of the given rational expression.

Step 2: Key Formula or Approach:

First, factor the denominator of the left-hand side: $x^2 + x - 2 = (x + 2)(x - 1)$. The setup is

correct.

The equation is:

$$\frac{3x}{(x+2)(x-1)} = \frac{A}{x+2} + \frac{B}{x-1}$$

To solve for A and B, we can combine the right side and equate the numerators.

$$\frac{3x}{(x+2)(x-1)} = \frac{A(x-1) + B(x+2)}{(x+2)(x-1)}$$

This gives the identity:

$$3x = A(x-1) + B(x+2)$$

We can find A and B by substituting strategic values for x (the roots of the denominator).

Step 3: Detailed Explanation:

To find B:

Let $x = 1$. This will make the term with A zero.

Substitute $x = 1$ into the identity:

$$3(1) = A(1-1) + B(1+2)$$

$$3 = A(0) + B(3)$$

$$3 = 3B$$

$$B = 1$$

To find A:

Let $x = -2$. This will make the term with B zero.

Substitute $x = -2$ into the identity:

$$3(-2) = A(-2-1) + B(-2+2)$$

$$-6 = A(-3) + B(0)$$

$$-6 = -3A$$

$$A = \frac{-6}{-3} = 2$$

So, we have $A = 2$ and $B = 1$.

Step 4: Final Answer:

The ordered pair (A, B) is (2, 1). This corresponds to option (D).

💡 Quick Tip

This method of substituting the roots of the denominator is called the "Heaviside cover-up method" and is the fastest way to find coefficients for distinct linear factors. To find A, "cover up" the (x+2) term on the left and substitute $x = -2$ into the rest:

$$\frac{3x}{x-1} \rightarrow \frac{3(-2)}{-2-1} = \frac{-6}{-3} = 2.$$

8. If $\tan A = \frac{4}{3}$ then the value of $\cos 2A$ is

- (A) $-\frac{7}{25}$
(B) $\frac{7}{24}$
(C) $-\frac{24}{7}$
(D) $\frac{7}{25}$

Correct Answer: (A) $-\frac{7}{25}$

Solution:

Step 1: Understanding the Question:

Given the value of $\tan A$, we need to find the value of $\cos 2A$. This requires using a double-angle trigonometric identity.

Step 2: Key Formula or Approach:

There are several formulas for $\cos 2A$, but the most direct one to use when $\tan A$ is given is:

$$\cos 2A = \frac{1 - \tan^2 A}{1 + \tan^2 A}$$

Step 3: Detailed Explanation:

We are given $\tan A = \frac{4}{3}$.

Substitute this value into the formula:

$$\cos 2A = \frac{1 - \left(\frac{4}{3}\right)^2}{1 + \left(\frac{4}{3}\right)^2}$$

First, calculate $\tan^2 A$:

$$\left(\frac{4}{3}\right)^2 = \frac{16}{9}$$

Now, substitute this back into the expression:

$$\cos 2A = \frac{1 - \frac{16}{9}}{1 + \frac{16}{9}}$$

Simplify the numerator and the denominator by finding a common denominator (which is 9):

$$\cos 2A = \frac{\frac{9}{9} - \frac{16}{9}}{\frac{9}{9} + \frac{16}{9}} = \frac{\frac{9-16}{9}}{\frac{9+16}{9}}$$

$$\cos 2A = \frac{\frac{-7}{9}}{\frac{25}{9}}$$

We can cancel the 9 in the denominator of both the numerator and the main denominator:

$$\cos 2A = -\frac{7}{25}$$

Step 4: Final Answer:

The value of $\cos 2A$ is $-\frac{7}{25}$. This corresponds to option (A).

💡 Quick Tip

Alternatively, you can visualize a right-angled triangle where $\tan A = \frac{\text{Opposite}}{\text{Adjacent}} = \frac{4}{3}$. By Pythagoras' theorem, the hypotenuse is $\sqrt{3^2 + 4^2} = \sqrt{9 + 16} = \sqrt{25} = 5$. From this triangle, $\cos A = \frac{3}{5}$. Then use the identity $\cos 2A = 2\cos^2 A - 1 = 2\left(\frac{3}{5}\right)^2 - 1 = 2\left(\frac{9}{25}\right) - 1 = \frac{18}{25} - \frac{25}{25} = -\frac{7}{25}$. Choose the method you find quickest.

9. If $-1 \leq x \leq 1$, then $\cos^{-1} x + \sin^{-1} x =$

- (A) $-\frac{\pi}{2}$
- (B) $\frac{\pi}{4}$
- (C) $\frac{\pi}{2}$
- (D) $-\frac{\pi}{16}$

Correct Answer: (C) $\frac{\pi}{2}$

Solution:

Step 1: Understanding the Question:

The question asks for the value of the sum of $\cos^{-1} x$ and $\sin^{-1} x$. This is a standard identity in inverse trigonometric functions.

Step 2: Key Formula or Approach:

The key identity for inverse sine and cosine functions is:

$$\sin^{-1} x + \cos^{-1} x = \frac{\pi}{2}$$

This identity is valid for all x in the domain $[-1, 1]$.

Step 3: Detailed Explanation:

This is a direct application of a fundamental property of inverse trigonometric functions.

Proof of the identity:

Let $\theta = \sin^{-1} x$. By definition, this means $\sin \theta = x$ and $-\frac{\pi}{2} \leq \theta \leq \frac{\pi}{2}$.

We know the trigonometric co-function identity: $\cos\left(\frac{\pi}{2} - \theta\right) = \sin \theta$.

Substituting $\sin \theta = x$, we get $\cos\left(\frac{\pi}{2} - \theta\right) = x$.

Now, we can take the inverse cosine of both sides:

$$\cos^{-1}\left(\cos\left(\frac{\pi}{2} - \theta\right)\right) = \cos^{-1} x$$

$$\frac{\pi}{2} - \theta = \cos^{-1} x$$

(This step is valid because $0 \leq \frac{\pi}{2} - \theta \leq \pi$, which is the range of $\cos^{-1}$).
Finally, substitute back $\theta = \sin^{-1} x$:

$$\frac{\pi}{2} - \sin^{-1} x = \cos^{-1} x$$

Rearranging the terms gives the identity:

$$\sin^{-1} x + \cos^{-1} x = \frac{\pi}{2}$$

Step 4: Final Answer:

The value of $\cos^{-1} x + \sin^{-1} x$ is $\frac{\pi}{2}$. This corresponds to option (C).

 Quick Tip

This is one of three important "complementary" identities for inverse trig functions that you should memorize: 1. $\sin^{-1} x + \cos^{-1} x = \frac{\pi}{2}$ 2. $\tan^{-1} x + \cot^{-1} x = \frac{\pi}{2}$ 3. $\sec^{-1} x + \csc^{-1} x = \frac{\pi}{2}$ They are frequently tested.

10. $\sin 15^\circ =$

(A) $\frac{\sqrt{3}-1}{\sqrt{3}+2}$

(B) $\frac{\sqrt{6}-\sqrt{2}}{4}$

(C) $\sqrt{6} \pm 1$

(D) $\frac{\sqrt{6}+\sqrt{2}}{4}$

Correct Answer: (B) $\frac{\sqrt{6}-\sqrt{2}}{4}$

Solution:

Step 1: Understanding the Question:

The question asks for the exact value of $\sin 15^\circ$. We can find this by expressing 15° as a difference of two standard angles (like 45° , 30° , 60°).

Step 2: Key Formula or Approach:

We will use the angle subtraction formula for sine:

$$\sin(A - B) = \sin A \cos B - \cos A \sin B$$

We can write 15° as $45^\circ - 30^\circ$.

Step 3: Detailed Explanation:

Let $A = 45^\circ$ and $B = 30^\circ$.

Applying the formula:

$$\sin(15^\circ) = \sin(45^\circ - 30^\circ) = \sin(45^\circ) \cos(30^\circ) - \cos(45^\circ) \sin(30^\circ)$$

We know the standard values for these angles:

$$\sin(45^\circ) = \frac{1}{\sqrt{2}}, \quad \cos(45^\circ) = \frac{1}{\sqrt{2}}$$

$$\sin(30^\circ) = \frac{1}{2}, \quad \cos(30^\circ) = \frac{\sqrt{3}}{2}$$

Substitute these values into the equation:

$$\sin(15^\circ) = \left(\frac{1}{\sqrt{2}}\right) \left(\frac{\sqrt{3}}{2}\right) - \left(\frac{1}{\sqrt{2}}\right) \left(\frac{1}{2}\right)$$

Combine the terms:

$$\sin(15^\circ) = \frac{\sqrt{3}}{2\sqrt{2}} - \frac{1}{2\sqrt{2}} = \frac{\sqrt{3} - 1}{2\sqrt{2}}$$

To match the format of the options, we rationalize the denominator by multiplying the numerator and denominator by $\sqrt{2}$:

$$\sin(15^\circ) = \frac{(\sqrt{3} - 1) \times \sqrt{2}}{(2\sqrt{2}) \times \sqrt{2}} = \frac{\sqrt{3}\sqrt{2} - 1\sqrt{2}}{2 \times 2} = \frac{\sqrt{6} - \sqrt{2}}{4}$$

Step 4: Final Answer:

The value of $\sin 15^\circ$ is $\frac{\sqrt{6}-\sqrt{2}}{4}$. This corresponds to option (B).

Quick Tip

The values for 15° and 75° are very common in competitive exams. It's highly beneficial to memorize them directly: - $\sin(15^\circ) = \cos(75^\circ) = \frac{\sqrt{6}-\sqrt{2}}{4}$ - $\cos(15^\circ) = \sin(75^\circ) = \frac{\sqrt{6}+\sqrt{2}}{4}$ - $\tan(15^\circ) = \cot(75^\circ) = 2 - \sqrt{3}$

11. If $2 \cos \theta = x + \frac{1}{x}$, then $2 \cos 3\theta =$

- (A) $x^3 - \frac{1}{x^3}$
- (B) $-x^3 + \frac{1}{x^3}$
- (C) $x^3 + \frac{1}{x^3}$
- (D) $x^2 + \frac{1}{x^2}$

Correct Answer: (C) $x^3 + \frac{1}{x^3}$

Solution:

Step 1: Understanding the Question:

We are given a relationship between $\cos \theta$ and a variable x . We need to find the corresponding relationship for $\cos 3\theta$.

Step 2: Key Formula or Approach:

This problem has a standard result derived from De Moivre's theorem. If $x = \cos \theta + i \sin \theta$, then $\frac{1}{x} = \cos \theta - i \sin \theta$. This leads to:

$$x + \frac{1}{x} = 2 \cos \theta$$

$$x^n + \frac{1}{x^n} = 2 \cos(n\theta)$$

We can use this general result directly.

Step 3: Detailed Explanation:**Method 1: Using the Standard Result**

We are given $x + \frac{1}{x} = 2 \cos \theta$.

Using the general identity $x^n + \frac{1}{x^n} = 2 \cos(n\theta)$, we can set $n = 3$.

This immediately gives:

$$x^3 + \frac{1}{x^3} = 2 \cos(3\theta)$$

Method 2: Algebraic Manipulation

We are given $2 \cos \theta = x + \frac{1}{x}$.

Let's cube both sides of the equation:

$$(2 \cos \theta)^3 = \left(x + \frac{1}{x}\right)^3$$

$$8 \cos^3 \theta = x^3 + 3 \cdot x^2 \cdot \left(\frac{1}{x}\right) + 3 \cdot x \cdot \left(\frac{1}{x}\right)^2 + \left(\frac{1}{x}\right)^3$$

$$8 \cos^3 \theta = x^3 + 3x + \frac{3}{x} + \frac{1}{x^3}$$

Group the terms:

$$8 \cos^3 \theta = \left(x^3 + \frac{1}{x^3}\right) + 3 \left(x + \frac{1}{x}\right)$$

We know that $x + \frac{1}{x} = 2 \cos \theta$, so we substitute this back in:

$$8 \cos^3 \theta = \left(x^3 + \frac{1}{x^3}\right) + 3(2 \cos \theta)$$

$$8 \cos^3 \theta = \left(x^3 + \frac{1}{x^3}\right) + 6 \cos \theta$$

Now, rearrange to solve for $x^3 + \frac{1}{x^3}$:

$$x^3 + \frac{1}{x^3} = 8 \cos^3 \theta - 6 \cos \theta$$

Factor out a 2 from the right-hand side:

$$x^3 + \frac{1}{x^3} = 2(4 \cos^3 \theta - 3 \cos \theta)$$

Using the triple angle identity for cosine, $\cos 3\theta = 4 \cos^3 \theta - 3 \cos \theta$, we get:

$$x^3 + \frac{1}{x^3} = 2 \cos 3\theta$$

Step 4: Final Answer:

The expression for $2 \cos 3\theta$ is $x^3 + \frac{1}{x^3}$. This corresponds to option (C).

💡 Quick Tip

Recognizing the connection to complex numbers is the key to solving this type of problem quickly. The relations $x^n + \frac{1}{x^n} = 2 \cos(n\theta)$ and $x^n - \frac{1}{x^n} = 2i \sin(n\theta)$ are powerful shortcuts.

12. In any $\triangle ABC$, $\tan\left(\frac{B+C}{2}\right) =$

- (A) $c \cot \frac{A}{2}$
- (B) $\cot \frac{A}{2}$
- (C) $\tan \frac{A}{2}$
- (D) $\tan \frac{C}{2}$

Correct Answer: (B) $\cot \frac{A}{2}$

Solution:

Step 1: Understanding the Question:

The question asks to express $\tan\left(\frac{B+C}{2}\right)$ in terms of an angle involving A, given that A, B, and C are angles of a triangle.

Step 2: Key Formula or Approach:

The fundamental property of a triangle is that the sum of its interior angles is 180° (or π radians).

$$A + B + C = 180^\circ$$

We will use this relation along with the trigonometric co-function identity $\tan(90^\circ - \theta) = \cot \theta$.

Step 3: Detailed Explanation:

Start with the angle sum property of a triangle:

$$A + B + C = 180^\circ$$

We need an expression for $B + C$, so we isolate these terms:

$$B + C = 180^\circ - A$$

The expression in the question involves $\frac{B+C}{2}$, so we divide the entire equation by 2:

$$\frac{B + C}{2} = \frac{180^\circ - A}{2}$$

$$\frac{B+C}{2} = \frac{180^\circ}{2} - \frac{A}{2}$$

$$\frac{B+C}{2} = 90^\circ - \frac{A}{2}$$

Now, we take the tangent of both sides:

$$\tan\left(\frac{B+C}{2}\right) = \tan\left(90^\circ - \frac{A}{2}\right)$$

Using the co-function identity $\tan(90^\circ - \theta) = \cot \theta$, with $\theta = \frac{A}{2}$:

$$\tan\left(\frac{B+C}{2}\right) = \cot\left(\frac{A}{2}\right)$$

Step 4: Final Answer:

The expression $\tan\left(\frac{B+C}{2}\right)$ is equal to $\cot \frac{A}{2}$. This corresponds to option (B).

💡 Quick Tip

For any problem involving the angles of a triangle, the first step is almost always to use the property $A + B + C = 180^\circ$. This allows you to relate the sum of any two angles to the third angle.

13. In a triangle $\triangle ABC$, the value of $\cos\left(\frac{B+C}{2}\right)$ in terms of angle A is

- (A) $\sqrt{\sin \frac{A}{2}}$
- (B) $\sqrt{A/2}$
- (C) $\sin \frac{A}{2}$
- (D) $\sqrt{2A}$

Correct Answer: (C) $\sin \frac{A}{2}$

Solution:

Step 1: Understanding the Question:

This question is similar to the previous one. It asks to express $\cos\left(\frac{B+C}{2}\right)$ in terms of a function of angle A.

Step 2: Key Formula or Approach:

We will again use the angle sum property of a triangle, $A + B + C = 180^\circ$. We will also use the trigonometric co-function identity $\cos(90^\circ - \theta) = \sin \theta$.

Step 3: Detailed Explanation:

From the angle sum property of a triangle:

$$A + B + C = 180^\circ$$

Isolate the term $B + C$:

$$B + C = 180^\circ - A$$

Divide by 2 to match the expression in the question:

$$\frac{B + C}{2} = \frac{180^\circ - A}{2} = 90^\circ - \frac{A}{2}$$

Now, take the cosine of both sides:

$$\cos\left(\frac{B + C}{2}\right) = \cos\left(90^\circ - \frac{A}{2}\right)$$

Using the co-function identity $\cos(90^\circ - \theta) = \sin \theta$, with $\theta = \frac{A}{2}$:

$$\cos\left(\frac{B + C}{2}\right) = \sin\left(\frac{A}{2}\right)$$

Step 4: Final Answer:

The value of $\cos\left(\frac{B+C}{2}\right)$ is $\sin \frac{A}{2}$. This corresponds to option (C).

 Quick Tip

Remember the set of co-function identities for $90^\circ - \theta$: $\sin(90^\circ - \theta) = \cos \theta$

$\cos(90^\circ - \theta) = \sin \theta$

$\tan(90^\circ - \theta) = \cot \theta$

These are essential for solving problems involving angles in a triangle.

14. The value of $\sin 45^\circ$ is

- (A) $\sqrt{2}$
- (B) 1
- (C) 0
- (D) $1/\sqrt{2}$

Correct Answer: (D) $1/\sqrt{2}$

Solution:

Step 1: Understanding the Question:

The question asks for the value of the sine function for the standard angle of 45° .

Step 2: Detailed Explanation:

We can determine this value by considering an isosceles right-angled triangle.

In such a triangle, the two angles other than the right angle are equal, so each is 45° .

Let the two equal sides (adjacent and opposite to the 45° angles) have a length of 1 unit. By the Pythagorean theorem, the hypotenuse h is:

$$h^2 = 1^2 + 1^2 = 1 + 1 = 2$$

$$h = \sqrt{2}$$

The definition of the sine of an angle in a right-angled triangle is:

$$\sin(\theta) = \frac{\text{Length of the Opposite Side}}{\text{Length of the Hypotenuse}}$$

For $\theta = 45^\circ$:

$$\sin(45^\circ) = \frac{1}{\sqrt{2}}$$

This can also be written by rationalizing the denominator as $\frac{\sqrt{2}}{2}$.

Step 3: Final Answer:

The value of $\sin 45^\circ$ is $1/\sqrt{2}$. This corresponds to option (D).

💡 Quick Tip

The trigonometric values for standard angles (0° , 30° , 45° , 60° , 90°) are fundamental and must be memorized for any competitive exam. Creating a table and practicing it can be very helpful.

15. In a $\triangle ABC$, if $a = 13$, $b = 14$ and $c = 15$ then the value of $\tan\left(\frac{A}{2}\right)$ is

- (A) $1/4$
- (B) $3/4$
- (C) $1/2$
- (D) $1/6$

Correct Answer: (C) $1/2$

Solution:

Step 1: Understanding the Question:

Given the lengths of the three sides of a triangle, we need to find the value of the tangent of a half-angle, specifically $\tan(A/2)$.

Step 2: Key Formula or Approach:

The half-angle formula for $\tan(A/2)$ in terms of the sides of a triangle is:

$$\tan\left(\frac{A}{2}\right) = \sqrt{\frac{(s-b)(s-c)}{s(s-a)}}$$

where s is the semi-perimeter of the triangle, calculated as $s = \frac{a+b+c}{2}$.

Step 3: Detailed Explanation:

First, calculate the semi-perimeter s :

$$s = \frac{13 + 14 + 15}{2} = \frac{42}{2} = 21$$

Next, calculate the values of $s - a$, $s - b$, and $s - c$:

$$s - a = 21 - 13 = 8$$

$$s - b = 21 - 14 = 7$$

$$s - c = 21 - 15 = 6$$

Now, substitute these values into the half-angle formula:

$$\tan\left(\frac{A}{2}\right) = \sqrt{\frac{(7)(6)}{(21)(8)}}$$

Simplify the expression inside the square root:

$$\tan\left(\frac{A}{2}\right) = \sqrt{\frac{42}{168}}$$

The fraction $\frac{42}{168}$ simplifies to $\frac{1}{4}$ (since $4 \times 42 = 168$).

$$\tan\left(\frac{A}{2}\right) = \sqrt{\frac{1}{4}}$$

$$\tan\left(\frac{A}{2}\right) = \frac{1}{2}$$

Step 4: Final Answer:

The value of $\tan\left(\frac{A}{2}\right)$ is $1/2$. This corresponds to option (C).

💡 Quick Tip

Another useful formula is $\tan(A/2) = \frac{\Delta}{s(s-a)}$, where Δ is the area of the triangle. Using Heron's formula, $\Delta = \sqrt{s(s-a)(s-b)(s-c)} = \sqrt{21 \cdot 8 \cdot 7 \cdot 6} = \sqrt{7056} = 84$. Then $\tan(A/2) = \frac{84}{21(8)} = \frac{84}{168} = \frac{1}{2}$. Both methods work well.

16. In a $\triangle ABC$, $\sum a^3 \cos(B - C) =$

- (A) $4abc$
- (B) $3abc$
- (C) $4a+b+c$

(D) abc

Correct Answer: (B) $3abc$

Solution:

Step 1: Understanding the Question:

The question asks for the simplified value of the cyclic sum $a^3 \cos(B - C) + b^3 \cos(C - A) + c^3 \cos(A - B)$ in a triangle ABC. This is a standard identity related to the properties of triangles.

Step 2: Key Formula or Approach:

This is a known, albeit complex, identity in trigonometry. The most straightforward approach in an exam setting is to recognize the identity or verify it with a special case.

Let's consider the term $a \cos(B - C)$.

Using the sine rule, $a = 2R \sin A$, and the angle sum property $A = 180^\circ - (B + C)$, so $\sin A = \sin(B + C)$.

$$a \cos(B - C) = 2R \sin(B + C) \cos(B - C)$$

Using the identity $2 \sin X \cos Y = \sin(X + Y) + \sin(X - Y)$:

$$a \cos(B - C) = R[\sin(2B) + \sin(2C)]$$

The full expression $a^3 \cos(B - C) = a^2 \cdot [a \cos(B - C)] = (2R \sin A)^2 \cdot R(\sin 2B + \sin 2C)$. Summing this up is very complicated.

Therefore, we treat this as a standard result.

Step 3: Detailed Explanation:

The expression is a standard result in the study of properties of triangles:

$$\sum a^3 \cos(B - C) = 3abc$$

We can verify this result for a simple case, like an equilateral triangle.

Verification for an equilateral triangle:

Let $a = b = c$ and $A = B = C = 60^\circ$.

The sum becomes:

$$a^3 \cos(60^\circ - 60^\circ) + b^3 \cos(60^\circ - 60^\circ) + c^3 \cos(60^\circ - 60^\circ)$$

Since $a = b = c$, this is:

$$a^3 \cos(0^\circ) + a^3 \cos(0^\circ) + a^3 \cos(0^\circ)$$

As $\cos(0^\circ) = 1$:

$$a^3(1) + a^3(1) + a^3(1) = 3a^3$$

Now, let's check the right-hand side from the options. The correct option is $3abc$.

$$3abc = 3(a)(a)(a) = 3a^3$$

Since the Left Hand Side (LHS) equals the Right Hand Side (RHS), the identity holds for an equilateral triangle, giving us confidence in the result.

Step 4: Final Answer:

The value of the given summation is $3abc$. This corresponds to option (B).

💡 Quick Tip

For complex cyclic sum identities in trigonometry, it's often not feasible to derive them during a timed exam. It's better to be familiar with common results. If you are unsure, testing a simple case like an equilateral triangle ($a=b=c$, $A=B=C=60^\circ$) or an isosceles right triangle can help you confirm the correct option.

17. Principle value of $\cot^{-1}(-1)$ is

- (A) $\frac{2\pi}{3}$
- (B) $-\frac{2\pi}{3}$
- (C) π
- (D) $\frac{3\pi}{4}$

Correct Answer: (D) $\frac{3\pi}{4}$

Solution:**Step 1: Understanding the Question:**

We need to find the principal value of the inverse cotangent of -1. The principal value is the value that lies within the defined principal value branch (range) of the function.

Step 2: Key Formula or Approach:

Let $y = \cot^{-1}(-1)$. By definition, this means $\cot(y) = -1$.

The principal value range for the inverse cotangent function, $y = \cot^{-1}(x)$, is $0 < y < \pi$.

This means the angle y must be in the first or second quadrant.

Step 3: Detailed Explanation:

First, we find the reference angle α for which the cotangent value is positive 1.

$$\cot(\alpha) = 1$$

This gives $\alpha = \frac{\pi}{4}$ (or 45°).

Since $\cot(y) = -1$ is negative, the angle y must lie in a quadrant where cotangent is negative. Within the range $(0, \pi)$, this is the second quadrant.

To find the angle in the second quadrant that has a reference angle of $\alpha = \frac{\pi}{4}$, we use the relation:

$$y = \pi - \alpha$$

$$y = \pi - \frac{\pi}{4}$$

$$y = \frac{4\pi - \pi}{4} = \frac{3\pi}{4}$$

This value, $\frac{3\pi}{4}$, lies within the principal value range $(0, \pi)$.

Step 4: Final Answer:

The principle value of $\cot^{-1}(-1)$ is $\frac{3\pi}{4}$. This corresponds to option (D).

 Quick Tip

Remember the principal value ranges for all inverse trigonometric functions. For negative inputs: - $\sin^{-1}, \tan^{-1}, \csc^{-1}$ give angles in $(-\pi/2, 0)$. - $\cos^{-1}, \cot^{-1}, \sec^{-1}$ give angles in $(\pi/2, \pi)$. For $\cot^{-1}(-x)$, a useful identity is $\cot^{-1}(-x) = \pi - \cot^{-1}(x)$. So, $\cot^{-1}(-1) = \pi - \cot^{-1}(1) = \pi - \frac{\pi}{4} = \frac{3\pi}{4}$.

18. $(-1 + 2i) + (\frac{1}{2} - i) =$

- (A) $\frac{1}{2} + i$
- (B) $-\frac{1}{2} - i$
- (C) $-\frac{1}{2} + i$
- (D) $\frac{1}{2} \pm i$

Correct Answer: (C) $-\frac{1}{2} + i$

Solution:

Step 1: Understanding the Question:

The question asks us to perform the addition of two complex numbers.

Step 2: Key Formula or Approach:

To add two complex numbers of the form $(a + bi)$ and $(c + di)$, we add the real parts together and the imaginary parts together:

$$(a + bi) + (c + di) = (a + c) + (b + d)i$$

Step 3: Detailed Explanation:

The two complex numbers are $-1 + 2i$ and $\frac{1}{2} - i$.

First, identify and group the real parts: -1 and $\frac{1}{2}$.

Add the real parts:

$$-1 + \frac{1}{2} = -\frac{2}{2} + \frac{1}{2} = -\frac{1}{2}$$

Next, identify and group the imaginary parts: $2i$ and $-i$.

Add the imaginary parts:

$$2i - i = (2 - 1)i = 1i = i$$

Combine the new real and imaginary parts to get the result:

$$-\frac{1}{2} + i$$

Step 4: Final Answer:

The sum of the complex numbers is $-\frac{1}{2} + i$. This corresponds to option (C).

💡 Quick Tip

Adding and subtracting complex numbers is just like combining like terms in algebra. Treat the imaginary unit 'i' as a variable, and simply combine the constant terms (real parts) and the 'i' terms (imaginary parts) separately.

19. For any real θ , $(\cos \theta + i \sin \theta)(\cos \theta - i \sin \theta) =$

- (A) 1
- (B) -1
- (C) 0
- (D) 4i

Correct Answer: (A) 1

Solution:

Step 1: Understanding the Question:

The question asks for the product of two complex numbers. We should recognize that the two numbers are complex conjugates of each other.

Step 2: Key Formula or Approach:

The product of a complex number $z = a + bi$ and its conjugate $\bar{z} = a - bi$ is given by the formula:

$$z\bar{z} = (a + bi)(a - bi) = a^2 + b^2$$

Alternatively, we can use the algebraic identity $(x + y)(x - y) = x^2 - y^2$.

Step 3: Detailed Explanation:

Method 1: Using the Conjugate Property

Let $z = \cos \theta + i \sin \theta$. Then its conjugate is $\bar{z} = \cos \theta - i \sin \theta$.

Here, $a = \cos \theta$ and $b = \sin \theta$.

Using the formula $z\bar{z} = a^2 + b^2$:

$$(\cos \theta + i \sin \theta)(\cos \theta - i \sin \theta) = (\cos \theta)^2 + (\sin \theta)^2 = \cos^2 \theta + \sin^2 \theta$$

Using the fundamental Pythagorean identity in trigonometry, $\cos^2 \theta + \sin^2 \theta = 1$.

Method 2: Using Algebraic Expansion

We expand the product using the difference of squares formula, $(x + y)(x - y) = x^2 - y^2$. Let $x = \cos \theta$ and $y = i \sin \theta$.

$$\begin{aligned} & (\cos \theta)^2 - (i \sin \theta)^2 \\ &= \cos^2 \theta - (i^2 \sin^2 \theta) \end{aligned}$$

Since $i^2 = -1$:

$$\begin{aligned} &= \cos^2 \theta - (-1 \cdot \sin^2 \theta) \\ &= \cos^2 \theta + \sin^2 \theta = 1 \end{aligned}$$

Method 3: Using Euler's Formula

Euler's formula states $e^{i\theta} = \cos \theta + i \sin \theta$.

This means $\cos \theta - i \sin \theta = \cos(-\theta) + i \sin(-\theta) = e^{-i\theta}$.

The product is:

$$e^{i\theta} \cdot e^{-i\theta} = e^{i\theta - i\theta} = e^0 = 1$$

Step 4: Final Answer:

The result of the multiplication is 1. This corresponds to option (A).

💡 Quick Tip

The product of a complex number and its conjugate always results in a real number equal to the square of the modulus of the complex number ($|z|^2$). For $z = \cos \theta + i \sin \theta$, the modulus $|z| = \sqrt{\cos^2 \theta + \sin^2 \theta} = \sqrt{1} = 1$. Therefore, $z\bar{z} = |z|^2 = 1^2 = 1$.

20. The centre and radius of the circle $x^2 + y^2 - 4x - 8y - 41 = 0$ are

- (A) (1, -2), 5
- (B) (2, 1), 3
- (C) (2, 4), $\sqrt{61}$
- (D) (1, -2), $\sqrt{51}$

Correct Answer: (C) (2, 4), $\sqrt{61}$

Solution:

Step 1: Understanding the Question:

We are given the general equation of a circle and asked to find its center and radius.

Step 2: Key Formula or Approach:

The general equation of a circle is given by $x^2 + y^2 + 2gx + 2fy + c = 0$.

The center of the circle is at the point $(-g, -f)$.

The radius of the circle is given by the formula $r = \sqrt{g^2 + f^2 - c}$.

Step 3: Detailed Explanation:

First, we compare the given equation $x^2 + y^2 - 4x - 8y - 41 = 0$ with the general form $x^2 + y^2 + 2gx + 2fy + c = 0$.

Comparing the coefficients of x:

$$2g = -4 \implies g = -2$$

Comparing the coefficients of y:

$$2f = -8 \implies f = -4$$

Comparing the constant term:

$$c = -41$$

Now, we can find the center of the circle using the formula $(-g, -f)$:

$$\text{Center} = (-(-2), -(-4)) = (2, 4).$$

Next, we calculate the radius using the formula $r = \sqrt{g^2 + f^2 - c}$:

$$r = \sqrt{(-2)^2 + (-4)^2 - (-41)}$$

$$r = \sqrt{4 + 16 + 41}$$

$$r = \sqrt{20 + 41}$$

$$r = \sqrt{61}$$

Step 4: Final Answer:

The center of the circle is $(2, 4)$ and the radius is $\sqrt{61}$. This corresponds to option (C).

💡 Quick Tip

To find the center from the general equation, simply take half the coefficients of x and y and change their signs. For the radius, square these halved coefficients, subtract the constant term 'c', and then take the square root. Be careful with the sign of 'c'.

21. The number of common tangents to the circles $x^2 + y^2 - x = 0$ and $x^2 + y^2 + x = 0$ is

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

Correct Answer: (D) 3

Solution:

Step 1: Understanding the Question:

We need to find the number of common tangents to two given circles. This depends on the relative positions of the two circles.

Step 2: Key Formula or Approach:

First, find the center and radius of each circle. Let the centers be C_1, C_2 and radii be r_1, r_2 . Then, calculate the distance d between the centers, $d = |C_1C_2|$.

Compare d with $r_1 + r_2$ and $|r_1 - r_2|$ to determine the relative positions:

- If $d > r_1 + r_2$, circles are separate (4 common tangents).
- If $d = r_1 + r_2$, circles touch externally (3 common tangents).
- If $|r_1 - r_2| < d < r_1 + r_2$, circles intersect (2 common tangents).
- If $d = |r_1 - r_2|$, circles touch internally (1 common tangent).
- If $d < |r_1 - r_2|$, one circle is inside another (0 common tangents).

Step 3: Detailed Explanation:

For the first circle, C_1 : $x^2 + y^2 - x = 0$

Comparing with $x^2 + y^2 + 2gx + 2fy + c = 0$:

$$2g_1 = -1 \implies g_1 = -1/2.$$

$$2f_1 = 0 \implies f_1 = 0.$$

$$c_1 = 0.$$

$$\text{Center } C_1 = (-g_1, -f_1) = (1/2, 0).$$

$$\text{Radius } r_1 = \sqrt{g_1^2 + f_1^2 - c_1} = \sqrt{(-1/2)^2 + 0^2 - 0} = \sqrt{1/4} = 1/2.$$

For the second circle, C_2 : $x^2 + y^2 + x = 0$

$$2g_2 = 1 \implies g_2 = 1/2.$$

$$2f_2 = 0 \implies f_2 = 0.$$

$$c_2 = 0.$$

$$\text{Center } C_2 = (-g_2, -f_2) = (-1/2, 0).$$

$$\text{Radius } r_2 = \sqrt{g_2^2 + f_2^2 - c_2} = \sqrt{(1/2)^2 + 0^2 - 0} = \sqrt{1/4} = 1/2.$$

Distance between centers:

The distance d between $C_1(1/2, 0)$ and $C_2(-1/2, 0)$ is:

$$d = \sqrt{\left(\frac{1}{2} - \left(-\frac{1}{2}\right)\right)^2 + (0 - 0)^2} = \sqrt{\left(\frac{1}{2} + \frac{1}{2}\right)^2} = \sqrt{1^2} = 1$$

Compare distance with sum of radii:

Sum of radii: $r_1 + r_2 = 1/2 + 1/2 = 1$.

Since $d = r_1 + r_2$ (as $1 = 1$), the two circles touch each other externally.

Step 4: Final Answer:

When two circles touch each other externally, they have exactly 3 common tangents (two direct and one transverse). This corresponds to option (D).

 Quick Tip

Before doing any calculations, notice that the second circle's equation is obtained by replacing 'x' with '-x' in the first. This implies the second circle is a reflection of the first across the y-axis. They both pass through the origin (0,0), which must be their point of contact. This visual understanding confirms they touch externally.

22. Equation of the circle with centre (-3, 2) and radius 4 is

- (A) $(x + 3)^2 + (y + 2)^2 = 4^2$
- (B) $(x - 3)^2 + (y + 2)^2 = 16$
- (C) $(x + 3)^2 + (y - 2)^2 = 16$
- (D) $(x - 2)^2 + (y + 3)^2 = 4^2$

Correct Answer: (C) $(x + 3)^2 + (y - 2)^2 = 16$

Solution:

Step 1: Understanding the Question:

We are given the center and radius of a circle and need to find its standard equation.

Step 2: Key Formula or Approach:

The standard equation of a circle with center at (h, k) and radius r is:

$$(x - h)^2 + (y - k)^2 = r^2$$

Step 3: Detailed Explanation:

We are given:

Center $(h, k) = (-3, 2)$.

Radius $r = 4$.

Substitute these values into the standard formula:

$$(x - (-3))^2 + (y - 2)^2 = 4^2$$

Simplify the equation:

$$(x + 3)^2 + (y - 2)^2 = 16$$

Step 4: Final Answer:

The equation of the circle is $(x + 3)^2 + (y - 2)^2 = 16$. This corresponds to option (C).

 Quick Tip

Be very careful with the signs when substituting the center coordinates (h, k) into the formula $(x - h)^2 + (y - k)^2 = r^2$. The signs inside the brackets will be the opposite of the signs of the coordinates of the center.

23. The length of the latus rectum of the parabola $y^2 = 12x$ and the focal distance of the point (3, -6) is

- (A) 3, 4
- (B) 2, 6
- (C) -12, 6
- (D) 12, 6

Correct Answer: (D) 12, 6

Solution:

Step 1: Understanding the Question:

This is a two-part question. We need to find:

1. The length of the latus rectum for the parabola $y^2 = 12x$.
2. The focal distance of a specific point on that parabola.

Step 2: Key Formula or Approach:

For a parabola in the standard form $y^2 = 4ax$:

- The length of the latus rectum is $4a$.
- The focal distance of any point (x_1, y_1) on the parabola is given by the formula $x_1 + a$.

Step 3: Detailed Explanation:

Part 1: Length of the Latus Rectum

First, compare the given equation $y^2 = 12x$ with the standard form $y^2 = 4ax$.

$$\begin{aligned} 4a &= 12 \\ a &= \frac{12}{4} = 3 \end{aligned}$$

The length of the latus rectum is equal to $4a$.

$$\text{Length of latus rectum} = 12$$

Part 2: Focal Distance

We need to find the focal distance of the point $(x_1, y_1) = (3, -6)$.

(First, let's verify that the point lies on the parabola: $(-6)^2 = 36$ and $12(3) = 36$. The point is on the parabola).

The formula for the focal distance is $x_1 + a$.

We have $x_1 = 3$ and we found $a = 3$.

$$\text{Focal distance} = 3 + 3 = 6$$

Step 4: Final Answer:

The length of the latus rectum is 12, and the focal distance of the point (3, -6) is 6. This corresponds to the pair (12, 6), which is option (D).

 Quick Tip

The length of the latus rectum is simply the coefficient of x (or y) in the standard form $y^2 = 4ax$ or $x^2 = 4ay$. You can read it directly from the equation. The focal distance is the distance from a point on the parabola to the focus, which by definition is also its distance to the directrix. The distance to the directrix ($x = -a$) is $|x_1 - (-a)| = x_1 + a$.

24. The equation of the Parabola, whose focus is (0,-2) and the vertex is (0,0), is

- (A) $y^2 = 32x$
- (B) $x^2 = -8y$
- (C) $x^2 = 4y$
- (D) $y^2 = -32x$

Correct Answer: (B) $x^2 = -8y$

Solution:

Step 1: Understanding the Question:

We are given the focus and vertex of a parabola and need to determine its equation.

Step 2: Key Formula or Approach:

1. Identify the axis of symmetry: Since the x-coordinates of the vertex (0,0) and focus (0,-2) are the same, the axis of symmetry is the y-axis.
2. Determine the direction: The vertex is at the origin, and the focus is below the vertex (at $y=-2$). This means the parabola opens downwards.
3. Use the standard equation: The standard equation for a parabola with vertex at the origin opening downwards is $x^2 = -4ay$, where a is the distance from the vertex to the focus. The focus for this type of parabola is at $(0, -a)$.

Step 3: Detailed Explanation:

The vertex is $V(0, 0)$ and the focus is $S(0, -2)$.

The distance from the vertex to the focus is a :

$$a = \sqrt{(0 - 0)^2 + (-2 - 0)^2} = \sqrt{(-2)^2} = 2$$

Alternatively, by comparing the focus $(0, -2)$ with the standard form $(0, -a)$, we directly get $a = 2$.

Now, substitute the value of $a = 2$ into the standard equation for a downward-opening parabola, $x^2 = -4ay$:

$$\begin{aligned}x^2 &= -4(2)y \\x^2 &= -8y\end{aligned}$$

Step 4: Final Answer:

The equation of the parabola is $x^2 = -8y$. This corresponds to option (B).

 Quick Tip

Quickly determine the orientation and direction. Vertex at origin and focus on an axis means a standard parabola. If the focus is on the y-axis, the equation is of the form $x^2 = \dots$. If the y-coordinate of the focus is negative, it opens down, so use a minus sign: $x^2 = -4ay$.

25. The eccentricity of $x^2 + 2y^2 = 3$ is

- (A) $\frac{1}{\sqrt{2}}$
- (B) $\sqrt{2}$
- (C) $\pm\sqrt{2}$
- (D) $\frac{\sqrt{3}}{2}$

Correct Answer: (A) $\frac{1}{\sqrt{2}}$

Solution:

Step 1: Understanding the Question:

We need to find the eccentricity of the conic section represented by the given equation. The form of the equation suggests it is an ellipse.

Step 2: Key Formula or Approach:

First, convert the given equation to the standard form of an ellipse, $\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1$.

The formula for eccentricity e is $e = \sqrt{1 - \frac{\text{minor axis squared}}{\text{major axis squared}}}$.

- If $a^2 > b^2$, the formula is $e = \sqrt{1 - \frac{b^2}{a^2}}$.

- If $b^2 > a^2$, the formula is $e = \sqrt{1 - \frac{a^2}{b^2}}$.

Step 3: Detailed Explanation:

The given equation is $x^2 + 2y^2 = 3$.

To get it into standard form, we divide the entire equation by 3:

$$\frac{x^2}{3} + \frac{2y^2}{3} = 1$$

$$\frac{x^2}{3} + \frac{y^2}{3/2} = 1$$

Now, we compare this with the standard form $\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1$.

We have $a^2 = 3$ and $b^2 = 3/2$.

Since $3 > 3/2$, we have $a^2 > b^2$. This means the major axis is horizontal.

Now we use the eccentricity formula for a horizontal ellipse:

$$e = \sqrt{1 - \frac{b^2}{a^2}}$$

$$e = \sqrt{1 - \frac{3/2}{3}}$$

$$e = \sqrt{1 - \frac{1}{2}}$$

$$e = \sqrt{\frac{1}{2}} = \frac{1}{\sqrt{2}}$$

Note on Answer Key:

The provided image indicates option (D) is correct, while our calculation yields option (A). There is a high probability of a typo in the question or the answer key. For the answer to be $\frac{\sqrt{3}}{2}$, the equation would need to be $x^2 + 4y^2 = \text{const}$ or $3x^2 + 4y^2 = \text{const}$. Based on the given equation, the calculation is correct. We will proceed with the mathematically derived answer.

Step 4: Final Answer:

The eccentricity of the ellipse $x^2 + 2y^2 = 3$ is $\frac{1}{\sqrt{2}}$. This corresponds to option (A).

 Quick Tip

Eccentricity 'e' of an ellipse is always between 0 and 1. This helps eliminate options like (B) and (C) immediately. To find e^2 , you only need the ratio of the denominators in the standard form: $e^2 = 1 - (\text{smaller denominator})/(\text{larger denominator})$.

26. $\frac{d}{dx}[e^x(x^2 + 1)] =$

- (A) $e^x(2x + x^2 + 1)$
- (B) $e^x(2x - x^2 + 1)$
- (C) $e^x(2x + x^3 + 1)$
- (D) $e^{-x}(2x + x^2 + 1)$

Correct Answer: (A) $e^x(2x + x^2 + 1)$

Solution:

Step 1: Understanding the Question:

We are asked to find the derivative of a product of two functions: e^x and $x^2 + 1$.

Step 2: Key Formula or Approach:

We must use the Product Rule for differentiation, which states:

$$\frac{d}{dx}(uv) = u \frac{dv}{dx} + v \frac{du}{dx}$$

Step 3: Detailed Explanation:

Let $u = e^x$ and $v = x^2 + 1$.

First, find the derivatives of u and v separately:

$$\frac{du}{dx} = \frac{d}{dx}(e^x) = e^x$$

$$\frac{dv}{dx} = \frac{d}{dx}(x^2 + 1) = 2x$$

Now, apply the product rule formula:

$$\frac{d}{dx}[e^x(x^2 + 1)] = (e^x)(2x) + (x^2 + 1)(e^x)$$

Factor out the common term e^x :

$$\begin{aligned} &= e^x[2x + (x^2 + 1)] \\ &= e^x(x^2 + 2x + 1) \end{aligned}$$

The expression in the parenthesis can also be written as $(x + 1)^2$. Rearranging the terms to match the option:

$$= e^x(2x + x^2 + 1)$$

Step 4: Final Answer:

The derivative is $e^x(2x + x^2 + 1)$. This corresponds to option (A).

💡 Quick Tip

A useful shortcut for the product rule when one function is e^x is: $\frac{d}{dx}(e^x f(x)) = e^x(f(x) + f'(x))$. Here, $f(x) = x^2 + 1$ and $f'(x) = 2x$. The result is $e^x((x^2 + 1) + 2x)$, which is the same.

27. When $a > 0$, $\lim_{x \rightarrow 0} \frac{a^x - 1}{x} =$

- (A) $\log a$
- (B) 0
- (C) $\log(x-1)$
- (D) $\log(x-a)$

Correct Answer: (A) $\log a$

Solution:

Step 1: Understanding the Question:

We need to evaluate a standard limit involving an exponential function.

Step 2: Key Formula or Approach:

This is a standard limit formula. However, we can derive it using L'Hôpital's Rule because direct substitution leads to an indeterminate form.

When we substitute $x = 0$: Numerator: $a^0 - 1 = 1 - 1 = 0$.

Denominator: 0.

Since we have the indeterminate form $\frac{0}{0}$, we can apply L'Hôpital's Rule, which states that $\lim_{x \rightarrow c} \frac{f(x)}{g(x)} = \lim_{x \rightarrow c} \frac{f'(x)}{g'(x)}$.

Step 3: Detailed Explanation:

Let $f(x) = a^x - 1$ and $g(x) = x$.

Find the derivatives of the numerator and the denominator:

$$f'(x) = \frac{d}{dx}(a^x - 1) = a^x \ln(a)$$

$$g'(x) = \frac{d}{dx}(x) = 1$$

Now, apply L'Hôpital's Rule:

$$\lim_{x \rightarrow 0} \frac{a^x - 1}{x} = \lim_{x \rightarrow 0} \frac{a^x \ln(a)}{1}$$

Now, we can substitute $x = 0$ into the new expression:

$$= \frac{a^0 \ln(a)}{1} = \frac{1 \cdot \ln(a)}{1} = \ln(a)$$

Note: In many contexts, $\log a$ is understood to mean the natural logarithm, $\ln a$.

Step 4: Final Answer:

The value of the limit is $\log a$. This corresponds to option (A).

💡 Quick Tip

This is a fundamental limit that is very useful to memorize for competitive exams:
 $\lim_{x \rightarrow 0} \frac{a^x - 1}{x} = \ln a$. A special case of this is when $a = e$, which gives $\lim_{x \rightarrow 0} \frac{e^x - 1}{x} = \ln e = 1$.

28. $\frac{d}{dx}[\tan^{-1} x] =$

- (A) $\frac{1}{x^2+1}$
- (B) $\frac{1}{x^2-1}$
- (C) $\frac{2}{x^2+2}$
- (D) $-\frac{1}{x^2+1}$

Correct Answer: (A) $\frac{1}{x^2+1}$

Solution:

Step 1: Understanding the Question:

The question asks for the derivative of the inverse tangent function, $\tan^{-1}(x)$.

Step 2: Key Formula or Approach:

This is a standard derivative formula in calculus. The derivation involves implicit differentiation.

Let $y = \tan^{-1}(x)$. Then $\tan(y) = x$.

Differentiate both sides with respect to x :

$$\frac{d}{dx}(\tan(y)) = \frac{d}{dx}(x)$$

$$\sec^2(y) \cdot \frac{dy}{dx} = 1$$

$$\frac{dy}{dx} = \frac{1}{\sec^2(y)}$$

Using the trigonometric identity $\sec^2(y) = 1 + \tan^2(y)$:

$$\frac{dy}{dx} = \frac{1}{1 + \tan^2(y)}$$

Since $\tan(y) = x$, we substitute this back:

$$\frac{dy}{dx} = \frac{1}{1 + x^2}$$

Step 3: Final Answer:

The derivative of $\tan^{-1}(x)$ with respect to x is $\frac{1}{1+x^2}$. This corresponds to option (A).

 Quick Tip

Memorizing the derivatives of the six inverse trigonometric functions is essential for speed in calculus problems. - $\frac{d}{dx}(\sin^{-1} x) = \frac{1}{\sqrt{1-x^2}}$ - $\frac{d}{dx}(\cos^{-1} x) = -\frac{1}{\sqrt{1-x^2}}$ - $\frac{d}{dx}(\tan^{-1} x) = \frac{1}{1+x^2}$ Note the relationship between the derivatives of co-functions (e.g., $\sin/\cos$, $\tan/\cot$).

29. If $4x - 7y + 15 = 0$ then derivative of y with respect to x is

- (A) $-4/7$
- (B) 0
- (C) 4
- (D) $4/7$

Correct Answer: (D) $4/7$

Solution:

Step 1: Understanding the Question:

We are given a linear equation relating x and y and asked to find the derivative of y with respect to x , which is $\frac{dy}{dx}$.

Step 2: Key Formula or Approach:

We can use two methods: implicit differentiation or rearranging the equation to make y the subject.

Step 3: Detailed Explanation:

Method 1: Implicit Differentiation

Differentiate each term of the equation $4x - 7y + 15 = 0$ with respect to x :

$$\frac{d}{dx}(4x) - \frac{d}{dx}(7y) + \frac{d}{dx}(15) = \frac{d}{dx}(0)$$

$$4 - 7\frac{dy}{dx} + 0 = 0$$

Now, solve for $\frac{dy}{dx}$:

$$4 = 7\frac{dy}{dx}$$

$$\frac{dy}{dx} = \frac{4}{7}$$

Method 2: Rearranging the Equation

First, solve the equation for y :

$$4x - 7y + 15 = 0$$

$$7y = 4x + 15$$

$$y = \frac{4x + 15}{7}$$

$$y = \frac{4}{7}x + \frac{15}{7}$$

This equation is now in the slope-intercept form $y = mx + c$, where the slope m is the derivative $\frac{dy}{dx}$.

By comparing, we can see that $m = \frac{4}{7}$. Therefore, $\frac{dy}{dx} = \frac{4}{7}$.

Step 4: Final Answer:

The derivative of y with respect to x is $4/7$. This corresponds to option (D).

 Quick Tip

For any linear equation in the form $Ax + By + C = 0$, the derivative $\frac{dy}{dx}$ (which represents the slope) is always equal to $-A/B$. In this case, $A = 4$ and $B = -7$, so the slope is $-4/(-7) = 4/7$.

30. If $y = \cos x$ then $\frac{d^2y}{dx^2} =$

- (A) $-\cos x$
- (B) $-\sin x$
- (C) $\cos x$
- (D) $\sin x$

Correct Answer: (A) $-\cos x$

Solution:

Step 1: Understanding the Question:

We are asked to find the second derivative of the function $y = \cos x$. This means we need to differentiate the function twice.

Step 2: Key Formula or Approach:

We will use the standard derivatives of trigonometric functions.

$$\frac{d}{dx}(\cos x) = -\sin x$$

$$\frac{d}{dx}(\sin x) = \cos x$$

Step 3: Detailed Explanation:

First, find the first derivative of $y = \cos x$:

$$\frac{dy}{dx} = \frac{d}{dx}(\cos x) = -\sin x$$

Now, differentiate the first derivative to find the second derivative:

$$\begin{aligned}\frac{d^2y}{dx^2} &= \frac{d}{dx} \left(\frac{dy}{dx} \right) = \frac{d}{dx}(-\sin x) \\ &= -\frac{d}{dx}(\sin x) \\ &= -(\cos x) = -\cos x\end{aligned}$$

Step 4: Final Answer:

The second derivative of $y = \cos x$ is $-\cos x$. This corresponds to option (A).

 Quick Tip

The derivatives of sine and cosine follow a cycle of four: $\sin x \rightarrow \cos x \rightarrow -\sin x \rightarrow -\cos x \rightarrow \sin x$. This pattern can be useful for quickly finding higher-order derivatives. Since $y = \cos x$, the first derivative is $-\sin x$ and the second is $-\cos x$.

31. If $u = e^x \sin y$ then first partial derivative of u with respect to y is

- (A) $e^x \sin y$
- (B) $e^x \cos y$
- (C) $-e^x \cos y$
- (D) 0

Correct Answer: (B) $e^x \cos y$

Solution:

Step 1: Understanding the Question:

We are given a function u of two variables, x and y , and asked to find its partial derivative with respect to y . This is denoted as $\frac{\partial u}{\partial y}$.

Step 2: Key Formula or Approach:

To find the partial derivative of a function with respect to one variable (e.g., y), we treat all other variables (e.g., x) as constants and apply the standard rules of differentiation.

Step 3: Detailed Explanation:

The function is $u(x, y) = e^x \sin y$.

We need to find $\frac{\partial u}{\partial y}$.

In this process, we treat x and therefore e^x as a constant.

$$\frac{\partial u}{\partial y} = \frac{\partial}{\partial y}(e^x \sin y)$$

We can take the constant term e^x outside the derivative:

$$= e^x \frac{\partial}{\partial y}(\sin y)$$

The derivative of $\sin y$ with respect to y is $\cos y$.

$$= e^x \cos y$$

Step 4: Final Answer:

The first partial derivative of u with respect to y is $e^x \cos y$. This corresponds to option (B).

 Quick Tip

When performing partial differentiation, mentally "freeze" the variables you are not differentiating with respect to. Imagine e^x is just a number like 5. The problem then becomes finding the derivative of $5 \sin y$, which is simply $5 \cos y$. Then replace the 5 back with e^x .

32. $\frac{d}{dx}(e^{3\log x}) =$

- (A) $\log x$
- (B) $3x$
- (C) x^3
- (D) $3x^2$

Correct Answer: (D) $3x^2$

Solution:

Step 1: Understanding the Question:

We need to find the derivative of the function $e^{3\log x}$. The first step should always be to simplify the function if possible before differentiating.

Step 2: Key Formula or Approach:

We will use the following properties of logarithms and exponentials:

1. Logarithm power rule: $n \log a = \log(a^n)$.
2. Inverse property: $e^{\log a} = e^{\ln a} = a$ (assuming $\log$ denotes the natural logarithm).
3. Power rule for differentiation: $\frac{d}{dx}(x^n) = nx^{n-1}$.

Step 3: Detailed Explanation:

First, simplify the expression $e^{3\log x}$.

Using the logarithm power rule, we can rewrite the exponent:

$$3 \log x = \log(x^3)$$

So the original expression becomes:

$$e^{\log(x^3)}$$

Using the inverse property of e^x and $\log x$, this simplifies to:

$$x^3$$

Now, the problem is reduced to finding the derivative of x^3 .

Using the power rule for differentiation:

$$\frac{d}{dx}(x^3) = 3x^{3-1} = 3x^2$$

Step 4: Final Answer:

The derivative of $e^{3\log x}$ is $3x^2$. This corresponds to option (D).

 Quick Tip

Always look for ways to simplify a function before differentiating or integrating. Expressions involving logarithms and exponentials, like $e^{\log f(x)}$ or $\log(e^{f(x)})$, can almost always be simplified to just $f(x)$, making the calculus much easier.

33. If $u(x, y) = \sin^{-1} \frac{x}{y} + \tan^{-1} \frac{y}{x}$ **then** $xu_x + yu_y =$

- (A) $u'(x, y)$
- (B) 0
- (C) 1
- (D) $u(x, y)$

Correct Answer: (B) 0

Solution:

Step 1: Understanding the Question:

We are given a function $u(x, y)$ and asked to compute the value of the expression $x \frac{\partial u}{\partial x} + y \frac{\partial u}{\partial y}$. This expression is related to Euler's Theorem for homogeneous functions.

Step 2: Key Formula or Approach:

Euler's Theorem on Homogeneous Functions: If $u(x, y)$ is a homogeneous function of degree n , then $x \frac{\partial u}{\partial x} + y \frac{\partial u}{\partial y} = nu$.

A function $u(x, y)$ is homogeneous of degree n if $u(tx, ty) = t^n u(x, y)$ for any constant t .

Step 3: Detailed Explanation:

First, let's check if the given function $u(x, y)$ is homogeneous. We substitute tx for x and ty for y :

$$u(tx, ty) = \sin^{-1} \left(\frac{tx}{ty} \right) + \tan^{-1} \left(\frac{ty}{tx} \right)$$

The t in the numerator and denominator cancels out in both terms:

$$u(tx, ty) = \sin^{-1} \left(\frac{x}{y} \right) + \tan^{-1} \left(\frac{y}{x} \right)$$

This is the same as the original function $u(x, y)$.

$$u(tx, ty) = u(x, y)$$

We can write this as:

$$u(tx, ty) = t^0 u(x, y)$$

This shows that $u(x, y)$ is a homogeneous function of degree $n = 0$.

Now, we can apply Euler's Theorem:

$$x \frac{\partial u}{\partial x} + y \frac{\partial u}{\partial y} = n \cdot u$$

Substituting $n = 0$:

$$x \frac{\partial u}{\partial x} + y \frac{\partial u}{\partial y} = 0 \cdot u = 0$$

Step 4: Final Answer:

The value of the expression $xu_x + yu_y$ is 0. This corresponds to option (B).

💡 Quick Tip

Whenever you see the expression $xu_x + yu_y$, your first thought should be Euler's Theorem. Quickly check if the function is homogeneous by replacing x with tx and y with ty . If all the t 's cancel out, the degree is 0, and the answer is 0. This saves you from calculating the partial derivatives explicitly.

34. If $S = 12t - 3t^2$ then $\frac{ds}{dt} =$

- (A) $12 - 6t$
- (B) $12t - 6$
- (C) $12 - 3t$
- (D) $12 - 6t^2$

Correct Answer: (A) $12 - 6t$

Solution:

Step 1: Understanding the Question:

We are given an equation for S in terms of t , and we need to find its derivative with respect to t , $\frac{ds}{dt}$.

Step 2: Key Formula or Approach:

We will use the power rule of differentiation, which states $\frac{d}{dt}(t^n) = nt^{n-1}$. We also use the sum/difference rule and the constant multiple rule.

Step 3: Detailed Explanation:

The function is $S = 12t - 3t^2$.

We differentiate term by term:

$$\frac{ds}{dt} = \frac{d}{dt}(12t) - \frac{d}{dt}(3t^2)$$

For the first term, $\frac{d}{dt}(12t) = 12 \cdot \frac{d}{dt}(t^1) = 12 \cdot 1 \cdot t^0 = 12$.

For the second term, $\frac{d}{dt}(3t^2) = 3 \cdot \frac{d}{dt}(t^2) = 3 \cdot (2t^1) = 6t$.

Combining the results:

$$\frac{ds}{dt} = 12 - 6t$$

Step 4: Final Answer:

The derivative of S with respect to t is $12 - 6t$. This corresponds to option (A).

 Quick Tip

This is a standard polynomial differentiation. Remember that the derivative of a term at^n is ant^{n-1} . The derivative of a linear term at is just the constant a , and the derivative of a constant is zero.

35. $\int \cot^2 x \, dx =$

- (A) $-\cot x + x + c$
- (B) $\cot x - x + c$
- (C) $\cot^2 x - x + c$
- (D) $-\cot x - x + c$

Correct Answer: (D) $-\cot x - x + c$

Solution:

Step 1: Understanding the Question:

We need to find the indefinite integral of $\cot^2(x)$.

Step 2: Key Formula or Approach:

There is no direct integration formula for $\cot^2(x)$. We must first use a trigonometric identity to rewrite it in a form that can be integrated.

The relevant Pythagorean identity is: $1 + \cot^2(x) = \csc^2(x)$.

From this, we get $\cot^2(x) = \csc^2(x) - 1$.

We know the standard integral: $\int \csc^2(x) \, dx = -\cot(x) + C$.

Step 3: Detailed Explanation:

Substitute the identity into the integral:

$$\int \cot^2(x) \, dx = \int (\csc^2(x) - 1) \, dx$$

Split the integral into two parts:

$$= \int \csc^2(x) \, dx - \int 1 \, dx$$

Now, integrate each part:

$$\begin{aligned} \int \csc^2(x) \, dx &= -\cot(x) \\ \int 1 \, dx &= x \end{aligned}$$

Combine the results and add the constant of integration, C :

$$= -\cot(x) - x + C$$

Step 4: Final Answer:

The integral of $\cot^2 x$ is $-\cot x - x + c$. This corresponds to option (D).

💡 Quick Tip

Similarly, to integrate $\tan^2(x)$, use the identity $\tan^2(x) = \sec^2(x) - 1$. This gives $\int \tan^2(x) dx = \int (\sec^2(x) - 1) dx = \tan(x) - x + C$. These two integrals are common exam questions.

36. $\int \frac{1}{\sqrt{a^2 - x^2}} dx =$

- (A) $\log |x + \sqrt{x^2 + a^2}| + c$
- (B) $\log |x + \sqrt{x^2 - a^2}| + c$
- (C) $\sin^{-1} \frac{x}{a} + c$
- (D) $\sin^{-1} x$

Correct Answer: (C) $\sin^{-1} \frac{x}{a} + c$

Solution:

Step 1: Understanding the Question:

We are asked to find the indefinite integral of the function $\frac{1}{\sqrt{a^2 - x^2}}$.

Step 2: Key Formula or Approach:

This is a standard integration formula that results in an inverse trigonometric function. The formula is:

$$\int \frac{1}{\sqrt{a^2 - x^2}} dx = \sin^{-1} \left(\frac{x}{a} \right) + C$$

This formula is derived using a trigonometric substitution $x = a \sin \theta$.

Step 3: Detailed Explanation:

This question requires the direct application of the standard integral formula.

By recalling the standard forms of integrals:

- $\int \frac{1}{x^2 + a^2} dx = \frac{1}{a} \tan^{-1} \left(\frac{x}{a} \right) + C$
- $\int \frac{1}{\sqrt{a^2 - x^2}} dx = \sin^{-1} \left(\frac{x}{a} \right) + C$
- $\int \frac{1}{\sqrt{x^2 + a^2}} dx = \ln |x + \sqrt{x^2 + a^2}| + C$
- $\int \frac{1}{\sqrt{x^2 - a^2}} dx = \ln |x + \sqrt{x^2 - a^2}| + C$

The given integral matches the second formula in this list.

Step 4: Final Answer:

The integral of $\frac{1}{\sqrt{a^2 - x^2}}$ is $\sin^{-1} \frac{x}{a} + c$. This corresponds to option (C).

 Quick Tip

It is crucial to memorize the standard integral formulas, especially those involving square roots of quadratic expressions. Pay close attention to the signs: $\sqrt{a^2 - x^2}$ leads to $\sin^{-1}$, while $\sqrt{x^2 + a^2}$ and $\sqrt{x^2 - a^2}$ lead to logarithmic forms.

37. $\int e^x \cos x \, dx =$

(A) $\frac{1}{2}e^x(\cos x + \sin x) + c$

(B) $\frac{1}{2}e^x \cos x$

(C) $\frac{1}{2}e^x(\cos x + \csc x) + c$

(D) $\frac{1}{2}e^x(\cos x - \sin x) + c$

Correct Answer: (A) $\frac{1}{2}e^x(\cos x + \sin x) + c$

Solution:

Step 1: Understanding the Question:

We need to find the indefinite integral of the product of e^x and $\cos x$.

Step 2: Key Formula or Approach:

This integral is typically solved using integration by parts, often twice. The formula for integration by parts is $\int u \, dv = uv - \int v \, du$.

There is also a standard formula for integrals of this type:

$$\int e^{ax} \cos(bx) \, dx = \frac{e^{ax}}{a^2 + b^2} (a \cos(bx) + b \sin(bx)) + C$$

Step 3: Detailed Explanation:

Let $I = \int e^x \cos x \, dx$.

We use integration by parts. Let $u = \cos x$ and $dv = e^x \, dx$.

Then $du = -\sin x \, dx$ and $v = e^x$.

Applying the formula:

$$I = (\cos x)(e^x) - \int (e^x)(-\sin x \, dx)$$

$$I = e^x \cos x + \int e^x \sin x \, dx$$

Now we must apply integration by parts to the new integral, $\int e^x \sin x \, dx$.

Let $u_2 = \sin x$ and $dv_2 = e^x \, dx$.

Then $du_2 = \cos x \, dx$ and $v_2 = e^x$.

$$\int e^x \sin x \, dx = (\sin x)(e^x) - \int (e^x)(\cos x \, dx) = e^x \sin x - I$$

Now, substitute this back into the equation for I :

$$I = e^x \cos x + (e^x \sin x - I)$$

$$I = e^x(\cos x + \sin x) - I$$

Now, solve for I :

$$2I = e^x(\cos x + \sin x)$$

$$I = \frac{1}{2}e^x(\cos x + \sin x) + c$$

Step 4: Final Answer:

The integral is $\frac{1}{2}e^x(\cos x + \sin x) + c$. This corresponds to option (A).

 Quick Tip

For exams, memorizing the direct formulas for $\int e^{ax} \cos(bx) dx$ and $\int e^{ax} \sin(bx) dx$ can save a significant amount of time compared to performing integration by parts twice.

38. $\int \frac{dx}{\sqrt{x}} =$

(A) $-2\sqrt{x} + c$

(B) $\sqrt{x} + c$

(C) $2\sqrt{x} + c$

(D) $x + c$

Correct Answer: (C) $2\sqrt{x} + c$

Solution:

Step 1: Understanding the Question:

We need to find the indefinite integral of $\frac{1}{\sqrt{x}}$.

Step 2: Key Formula or Approach:

First, rewrite the integrand using exponent notation: $\frac{1}{\sqrt{x}} = \frac{1}{x^{1/2}} = x^{-1/2}$.

Then, use the power rule for integration:

$$\int x^n dx = \frac{x^{n+1}}{n+1} + C \quad (\text{for } n \neq -1)$$

Step 3: Detailed Explanation:

Rewrite the integral:

$$\int \frac{dx}{\sqrt{x}} = \int x^{-1/2} dx$$

Apply the power rule with $n = -1/2$:

$$\begin{aligned} &= \frac{x^{-1/2+1}}{-1/2+1} + C \\ &= \frac{x^{1/2}}{1/2} + C \end{aligned}$$

Dividing by $1/2$ is the same as multiplying by 2:

$$= 2x^{1/2} + C$$

Rewrite the result using radical notation:

$$= 2\sqrt{x} + C$$

Step 4: Final Answer:

The integral is $2\sqrt{x} + c$. This corresponds to option (C).

💡 Quick Tip

The integral of $1/\sqrt{x}$ is a very common one. It's helpful to remember it directly. Differentiating the answer, $\frac{d}{dx}(2\sqrt{x}) = 2 \cdot \frac{1}{2}x^{-1/2} = \frac{1}{\sqrt{x}}$, confirms the result quickly.

39. $\int \sin \frac{y}{2} dy =$

- (A) $2 \cos \frac{y}{2} + c$
- (B) $2 \sin x/2 + c$
- (C) $2 \cos 2y + c$
- (D) $-2 \cos \frac{y}{2} + c$

Correct Answer: (D) $-2 \cos \frac{y}{2} + c$

Solution:

Step 1: Understanding the Question:

We need to find the indefinite integral of $\sin(y/2)$.

Step 2: Key Formula or Approach:

We will use the standard integral $\int \sin(u) du = -\cos(u) + C$ combined with the method of u-substitution. A useful shortcut for linear arguments is: $\int f(ay + b) dy = \frac{1}{a}F(ay + b) + C$, where F is the antiderivative of f .

Step 3: Detailed Explanation:

Let's use u-substitution. Let $u = \frac{y}{2}$.

Then, differentiate u with respect to y :

$$\frac{du}{dy} = \frac{1}{2} \implies dy = 2 du$$

Now, substitute u and dy into the original integral:

$$\begin{aligned} \int \sin\left(\frac{y}{2}\right) dy &= \int \sin(u) \cdot (2 du) \\ &= 2 \int \sin(u) du \end{aligned}$$

Integrate with respect to u :

$$= 2(-\cos(u)) + C = -2\cos(u) + C$$

Finally, substitute back $u = \frac{y}{2}$:

$$= -2\cos\left(\frac{y}{2}\right) + C$$

Step 4: Final Answer:

The integral is $-2\cos\frac{y}{2} + c$. This corresponds to option (D).

 Quick Tip

When integrating $\sin(ay)$ or $\cos(ay)$, the result is $-\frac{1}{a}\cos(ay)$ or $\frac{1}{a}\sin(ay)$ respectively. Just integrate the function as usual and then divide by the coefficient of the variable. Here the coefficient of y is $1/2$, so we divide by $1/2$ (which means multiplying by 2).

40. $\int_0^\pi dx =$

- (A) $\pi/2$
- (B) $-\pi/2$
- (C) π
- (D) $\pi/8$

Correct Answer: (C) π

Solution:

Step 1: Understanding the Question:

We need to evaluate the definite integral of the function $f(x) = 1$ over the interval from 0 to π .

Step 2: Key Formula or Approach:

The fundamental theorem of calculus states that if $F'(x) = f(x)$, then $\int_a^b f(x) dx = F(b) - F(a)$.

First, we find the antiderivative of the integrand. The integrand is implicitly 1.

$$\int 1 \, dx = x$$

Step 3: Detailed Explanation:

The antiderivative of the integrand (1) is $F(x) = x$.

Now, we evaluate this antiderivative at the upper and lower limits of integration, which are $b = \pi$ and $a = 0$.

$$\begin{aligned}\int_0^\pi 1 \, dx &= [x]_0^\pi \\ &= F(\pi) - F(0) \\ &= (\pi) - (0) \\ &= \pi\end{aligned}$$

Geometrically, this integral represents the area of a rectangle with height 1 and width π , so the area is $1 \times \pi = \pi$.

Step 4: Final Answer:

The value of the definite integral is π . This corresponds to option (C).

 Quick Tip

The definite integral $\int_a^b C \, dx$ of any constant C is simply $C \times (b - a)$. In this case, $C = 1$, $b = \pi$, and $a = 0$, so the result is $1 \times (\pi - 0) = \pi$.

41. If $f(x)$ is an even function, then $\int_{-a}^a f(x) \, dx =$

- (A) $\int_0^a f(x) \, dx$
- (B) $2 \int_0^a f(x) \, dx$
- (C) $2a$
- (D) 0

Correct Answer: (B) $2 \int_0^a f(x) \, dx$

Solution:

Step 1: Understanding the Question:

The question asks to simplify the definite integral of an even function over a symmetric interval $[-a, a]$.

Step 2: Key Formula or Approach:

An even function is defined by the property $f(-x) = f(x)$. Geometrically, its graph is symmetric with respect to the y-axis.

We use a property of definite integrals to split the interval:

$$\int_{-a}^a f(x)dx = \int_{-a}^0 f(x)dx + \int_0^a f(x)dx$$

Step 3: Detailed Explanation:

Let's evaluate the first part of the split integral, $\int_{-a}^0 f(x)dx$.

We use a substitution. Let $x = -u$. Then $dx = -du$.

We must also change the limits of integration: - When $x = -a$, $u = -(-a) = a$.

- When $x = 0$, $u = -(0) = 0$.

Substituting these into the integral gives:

$$\int_a^0 f(-u)(-du)$$

Since $f(x)$ is an even function, $f(-u) = f(u)$.

$$\int_a^0 f(u)(-du) = - \int_a^0 f(u)du$$

We can reverse the limits of integration by changing the sign of the integral:

$$= \int_0^a f(u)du$$

Since u is a dummy variable, this is the same as $\int_0^a f(x)dx$.

Now, substitute this result back into the original split equation:

$$\begin{aligned} \int_{-a}^a f(x)dx &= \left(\int_0^a f(x)dx \right) + \int_0^a f(x)dx \\ &= 2 \int_0^a f(x)dx \end{aligned}$$

Step 4: Final Answer:

For an even function, the integral over a symmetric interval is twice the integral over the positive half of the interval. This corresponds to option (B).

💡 Quick Tip

For an even function, the area under the curve from $-a$ to 0 is identical to the area from 0 to a due to y-axis symmetry. Thus, the total area is simply twice the area of one side. For an odd function ($f(-x) = -f(x)$), the areas cancel out, and $\int_{-a}^a f(x)dx = 0$.

42. The area under the curve $f(x) = \sin x$ in $[0, 2\pi]$ is

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

Correct Answer: (D) 4

Solution:

Step 1: Understanding the Question:

The question asks for the "area under the curve", which implies the total geometric area between the curve and the x-axis. This means we must treat any area below the x-axis as positive. The formula for area is $A = \int_a^b |f(x)| dx$.

Step 2: Key Formula or Approach:

We need to analyze the sign of $\sin x$ in the interval $[0, 2\pi]$.

- In the interval $[0, \pi]$, $\sin x \geq 0$, so $|\sin x| = \sin x$.
- In the interval $[\pi, 2\pi]$, $\sin x \leq 0$, so $|\sin x| = -\sin x$.

Therefore, we must split the integral:

$$A = \int_0^{2\pi} |\sin x| dx = \int_0^{\pi} \sin x dx + \int_{\pi}^{2\pi} (-\sin x) dx$$

Step 3: Detailed Explanation:

First, evaluate the integral for the interval $[0, \pi]$:

$$\begin{aligned} \int_0^{\pi} \sin x dx &= [-\cos x]_0^{\pi} \\ &= (-\cos(\pi)) - (-\cos(0)) = -(-1) - (-1) = 1 + 1 = 2 \end{aligned}$$

Next, evaluate the integral for the interval $[\pi, 2\pi]$:

$$\begin{aligned} \int_{\pi}^{2\pi} -\sin x dx &= [\cos x]_{\pi}^{2\pi} \\ &= (\cos(2\pi)) - (\cos(\pi)) = (1) - (-1) = 1 + 1 = 2 \end{aligned}$$

The total area is the sum of the areas from the two intervals:

$$A = 2 + 2 = 4$$

Step 4: Final Answer:

The total area under the curve $f(x) = \sin x$ in the interval $[0, 2\pi]$ is 4. This corresponds to option (D).

 Quick Tip

Be careful with questions asking for "area". If you calculate $\int_0^{2\pi} \sin x \, dx$ directly, you get 0, as the positive and negative areas cancel out. "Area under the curve" almost always implies the total absolute area. The area of one "hump" of the sine curve is 2. Over a full period, there are two humps, so the total area is 4.

43. When $a=b$ then $\int_a^b f(x)dx =$

- (A) b
- (B) 0
- (C) a
- (D) 2a

Correct Answer: (B) 0

Solution:

Step 1: Understanding the Question:

We are asked to evaluate a definite integral where the upper and lower limits of integration are the same.

Step 2: Key Formula or Approach:

According to the Fundamental Theorem of Calculus, if $F(x)$ is the antiderivative of $f(x)$, then:

$$\int_a^b f(x)dx = F(b) - F(a)$$

Step 3: Detailed Explanation:

The question states that $a = b$. We can substitute a for b in the formula:

$$\begin{aligned}\int_a^a f(x)dx &= F(a) - F(a) \\ &= 0\end{aligned}$$

Geometrically, a definite integral represents the area under a curve between two points. If the starting and ending points are the same, the width of the region is zero, and therefore its area is also zero.

Step 4: Final Answer:

When the limits of integration are equal, the value of the definite integral is 0. This corresponds to option (B).

 Quick Tip

This is a fundamental property of definite integrals that you should know by heart. Any integral of the form $\int_a^a f(x)dx$ is always zero, regardless of the function $f(x)$.

44. The Order of the differential equation $\left[\frac{d^2y}{dx^2} + \left(\frac{dy}{dx} \right)^3 \right]^{6/5} = 6y$ is

- (A) 3
- (B) 2
- (C) 6/5
- (D) 3

Correct Answer: (B) 2

Solution:

Step 1: Understanding the Question:

We need to find the 'order' of the given differential equation.

Step 2: Detailed Explanation:

The **order** of a differential equation is defined as the order of the highest derivative that appears in the equation.

Let's examine the derivatives present in the equation:

$$\left[\frac{d^2y}{dx^2} + \left(\frac{dy}{dx} \right)^3 \right]^{6/5} = 6y$$

The derivatives involved are: 1. $\frac{dy}{dx}$: This is the first derivative, so its order is 1.

2. $\frac{d^2y}{dx^2}$: This is the second derivative, so its order is 2.

The highest order among all the derivatives in the equation is 2.

The powers and fractional exponents on the derivatives or the terms do not affect the order of the equation. They are related to the 'degree' of the equation, which is a different concept.

Step 3: Final Answer:

The order of the highest derivative is 2, so the order of the differential equation is 2. This corresponds to option (B).

 Quick Tip

Don't confuse order with degree. - **Order**: The highest derivative (e.g., y'' , y''' , etc.). It's the most important classification. - **Degree**: The power of the highest-order derivative after the equation has been cleared of any radicals or fractional exponents involving the derivatives. In this case, raising both sides to the power of 5 gives a degree of 6.

45. The Integrating factor of $\frac{dy}{dx} + 3x = 2y$ is

- (A) e^{3x}
- (B) e^{-2x}
- (C) e^x
- (D) 0

Correct Answer: (B) e^{-2x}

Solution:

Step 1: Understanding the Question:

We need to find the integrating factor for the given differential equation. First, we must arrange the equation into the standard linear form.

Step 2: Key Formula or Approach:

A first-order linear differential equation has the standard form:

$$\frac{dy}{dx} + P(x)y = Q(x)$$

The integrating factor (I.F.) for this form is given by the formula:

$$\text{I.F.} = e^{\int P(x)dx}$$

Step 3: Detailed Explanation:

First, rearrange the given equation $\frac{dy}{dx} + 3x = 2y$ into the standard linear form. We need to move the term with 'y' to the left side.

$$\frac{dy}{dx} - 2y = -3x$$

Now, compare this with the standard form $\frac{dy}{dx} + P(x)y = Q(x)$. We can identify:

$$P(x) = -2$$

$$Q(x) = -3x$$

Next, we calculate the integrating factor using its formula:

$$\text{I.F.} = e^{\int P(x)dx} = e^{\int (-2)dx}$$

Performing the integration in the exponent:

$$\int -2 \, dx = -2x$$

So, the integrating factor is:

$$\text{I.F.} = e^{-2x}$$

Step 4: Final Answer:

The integrating factor of the differential equation is e^{-2x} . This corresponds to option (B).

 Quick Tip

The most common mistake in these problems is incorrect identification of $P(x)$. Always ensure the equation is in the exact standard form $y' + P(x)y = Q(x)$ before you determine $P(x)$. Pay close attention to the signs.

46. Transform $dx + xdy = e^{-y} \sec^2 y \, dy$ into linear form

(A) $\frac{dx}{dy} - x = e^{-y} \sec^2 y$

(B) $\frac{dx}{dy} = e^{-y} \sec^2 y$

(C) $\frac{dx}{dy} + x = e^{-y} \sec^2 y + c$

(D) $\frac{dx}{dy} + x = e^{-y} \sec^2 y$

Correct Answer: (D) $\frac{dx}{dy} + x = e^{-y} \sec^2 y$

Solution:

Step 1: Understanding the Question:

We are asked to rewrite the given differential equation into a standard linear form. The equation involves both dx and dy , so it could be linear in y (with x as the independent variable) or linear in x (with y as the independent variable).

Step 2: Key Formula or Approach:

The standard linear form with x as the dependent variable and y as the independent variable is:

$$\frac{dx}{dy} + P(y)x = Q(y)$$

We will try to arrange the given equation into this form.

Step 3: Detailed Explanation:

The given equation is:

$$dx + xdy = e^{-y} \sec^2 y \, dy$$

To get $\frac{dx}{dy}$, we should divide the entire equation by dy .

$$\frac{dx}{dy} + \frac{xdy}{dy} = \frac{e^{-y} \sec^2 y dy}{dy}$$

Simplifying each term:

$$\frac{dx}{dy} + x = e^{-y} \sec^2 y$$

This equation is now in the standard linear form $\frac{dx}{dy} + P(y)x = Q(y)$, where $P(y) = 1$ and $Q(y) = e^{-y} \sec^2 y$.

Step 4: Final Answer:

The linear form of the equation is $\frac{dx}{dy} + x = e^{-y} \sec^2 y$. This corresponds to option (D).

💡 Quick Tip

If a differential equation is hard to arrange into the standard $\frac{dy}{dx}$ linear form, always check if it can be arranged into the $\frac{dx}{dy}$ linear form. The clue is often that the terms multiplying dy are functions of y only.

47. The necessary and sufficient condition for the differential equation $Mdx + Ndy = 0$ to be exact is

(A) $\frac{\partial M}{\partial y} = \frac{\partial N}{\partial y}$

(B) $\frac{\partial M}{\partial y} = \frac{\partial N}{\partial x}$

(C) $\frac{\partial M}{\partial y} \neq \frac{\partial N}{\partial x}$

(D) $\frac{\partial M}{\partial x} = \frac{\partial N}{\partial x}$

Correct Answer: (B) $\frac{\partial M}{\partial y} = \frac{\partial N}{\partial x}$

Solution:

Step 1: Understanding the Question:

This is a theoretical question asking for the condition that defines an exact differential equation.

Step 2: Detailed Explanation:

A differential equation of the form $M(x, y)dx + N(x, y)dy = 0$ is said to be **exact** if its left-hand side is the total differential of some function $u(x, y)$.

The total differential of $u(x, y)$ is given by:

$$du = \frac{\partial u}{\partial x}dx + \frac{\partial u}{\partial y}dy$$

For the equation to be exact, we must have:

$$M = \frac{\partial u}{\partial x} \quad \text{and} \quad N = \frac{\partial u}{\partial y}$$

Now, we take the partial derivative of M with respect to y and N with respect to x :

$$\frac{\partial M}{\partial y} = \frac{\partial}{\partial y} \left(\frac{\partial u}{\partial x} \right) = \frac{\partial^2 u}{\partial y \partial x}$$

$$\frac{\partial N}{\partial x} = \frac{\partial}{\partial x} \left(\frac{\partial u}{\partial y} \right) = \frac{\partial^2 u}{\partial x \partial y}$$

According to Clairaut's theorem, if the second partial derivatives are continuous, then the order of differentiation does not matter, meaning $\frac{\partial^2 u}{\partial y \partial x} = \frac{\partial^2 u}{\partial x \partial y}$.

Therefore, the necessary and sufficient condition for the equation to be exact is:

$$\frac{\partial M}{\partial y} = \frac{\partial N}{\partial x}$$

Step 3: Final Answer:

The condition for an exact differential equation is $\frac{\partial M}{\partial y} = \frac{\partial N}{\partial x}$. This corresponds to option (B).

💡 Quick Tip

A simple way to remember the condition is to think of it as "cross-differentiation": The coefficient of dx (M) is differentiated with respect to the "other" variable (y), and the coefficient of dy (N) is differentiated with respect to its "other" variable (x).

48. Complementary function of the differential equation $(D^3 - 8)y = x$ is

- (A) $C_1 e^{2x} + e^x \{C_2 \cos(x\sqrt{3}) + C_3 \sin(x\sqrt{3})\}$
- (B) $C_1 e^{2x} + e^{-x} (\cos \sqrt{3} + \sin \sqrt{3})$
- (C) $C_1 e^{-2x} + C_2 e^{-x} + C_3 \cos x$
- (D) $C_1 e^{2x} + e^{-x} \{C_2 \cos(x\sqrt{3}) + C_3 \sin(x\sqrt{3})\}$

Correct Answer: (D) $C_1 e^{2x} + e^{-x} \{C_2 \cos(x\sqrt{3}) + C_3 \sin(x\sqrt{3})\}$

Solution:

Step 1: Understanding the Question:

We need to find the complementary function (CF) for the given linear non-homogeneous differential equation. The CF is the general solution to the corresponding homogeneous equation.

Step 2: Key Formula or Approach:

The homogeneous equation is $(D^3 - 8)y = 0$. To solve this, we first form the auxiliary equation by replacing the operator D with a variable m .

$$m^3 - 8 = 0$$

We then find the roots of this auxiliary equation. The form of the CF depends on the nature of these roots (real, repeated, complex).

Step 3: Detailed Explanation:

The auxiliary equation is $m^3 - 8 = 0$, or $m^3 = 8$.

We can factor this as a difference of cubes, $a^3 - b^3 = (a - b)(a^2 + ab + b^2)$:

$$(m - 2)(m^2 + 2m + 4) = 0$$

One root is obtained from the first factor: $m - 2 = 0 \implies m_1 = 2$. This is a real, distinct root.

The other two roots come from the quadratic factor $m^2 + 2m + 4 = 0$. We use the quadratic formula $m = \frac{-b \pm \sqrt{b^2 - 4ac}}{2a}$:

$$m = \frac{-2 \pm \sqrt{2^2 - 4(1)(4)}}{2(1)} = \frac{-2 \pm \sqrt{4 - 16}}{2} = \frac{-2 \pm \sqrt{-12}}{2}$$

$$m = \frac{-2 \pm \sqrt{4 \cdot 3 \cdot (-1)}}{2} = \frac{-2 \pm 2i\sqrt{3}}{2} = -1 \pm i\sqrt{3}$$

So, we have a pair of complex conjugate roots: $m_{2,3} = -1 \pm i\sqrt{3}$. This is of the form $\alpha \pm i\beta$ where $\alpha = -1$ and $\beta = \sqrt{3}$.

Now we construct the complementary function from the roots: - The real root $m_1 = 2$ gives the term $C_1 e^{2x}$.

- The complex pair $\alpha \pm i\beta = -1 \pm i\sqrt{3}$ gives the term $e^{\alpha x}(C_2 \cos(\beta x) + C_3 \sin(\beta x))$, which is $e^{-x}(C_2 \cos(x\sqrt{3}) + C_3 \sin(x\sqrt{3}))$.

Combining these parts gives the complete complementary function:

$$y_c = C_1 e^{2x} + e^{-x} \{C_2 \cos(x\sqrt{3}) + C_3 \sin(x\sqrt{3})\}$$

Step 4: Final Answer:

The complementary function is $C_1 e^{2x} + e^{-x} \{C_2 \cos(x\sqrt{3}) + C_3 \sin(x\sqrt{3})\}$. This corresponds to option (D).

💡 Quick Tip

Recognizing algebraic forms like the difference/sum of cubes is very helpful for solving auxiliary equations quickly. For $m^3 - a^3 = 0$, the roots are always a and $a(-\frac{1}{2} \pm i\frac{\sqrt{3}}{2})$. Here $a = 2$.

49. Bernoulli's equation is of the form

(A) $(\frac{dy}{dx}) + y = Qy$

(B) $(\frac{dy}{dx})^2 + y^n = Qy$

(C) $(\frac{dy}{dx}) + Py = Qy^n$

(D) $(\frac{d^2y}{dx^2}) + Py = Qy^n$

Correct Answer: (C) $(\frac{dy}{dx}) + Py = Qy^n$

Solution:

Step 1: Understanding the Question:

This is a definition-based question that asks for the standard form of Bernoulli's differential equation.

Step 2: Detailed Explanation:

A Bernoulli differential equation is a first-order ordinary differential equation of the form:

$$\frac{dy}{dx} + P(x)y = Q(x)y^n$$

where n is any real number.

Let's analyze the options: - (A) is a linear equation. - (B) is non-linear due to $(\frac{dy}{dx})^2$ and is not the Bernoulli form. - (C) matches the standard definition of a Bernoulli equation exactly, where P and Q are functions of x . - (D) is a second-order differential equation.

Step 3: Final Answer:

The correct form for Bernoulli's equation is $\frac{dy}{dx} + Py = Qy^n$. This corresponds to option (C).

 Quick Tip

Bernoulli's equation is a generalization of the linear equation. If $n = 0$ or $n = 1$, the equation becomes a standard first-order linear equation. The key feature is the y^n term on the right-hand side.

50. Particular integral of $f(D)y = \cos ax$ is

(A) $\frac{1}{f(-a^2)} \cos ax$ if $f(-a^2) \neq 0$

(B) $\frac{1}{f(a^2)} \cos ax$ if $f(-a^2) \neq 0$

(C) $\frac{1}{f(a)} \cos ax$ if $f(-a^2) \neq 0$

(D) $\frac{1}{2} \cos ax$ if $f(-a^2) \neq 0$

Correct Answer: (A) $\frac{1}{f(-a^2)} \cos ax$ if $f(-a^2) \neq 0$

Solution:

Step 1: Understanding the Question:

This question asks for the standard rule to find the particular integral (PI) for a linear differential equation with constant coefficients when the right-hand side (the forcing function) is $\cos(ax)$.

Step 2: Detailed Explanation:

For a differential equation of the form $f(D)y = X(x)$, where $f(D)$ is a polynomial in the differential operator $D = \frac{d}{dx}$, the particular integral is formally written as:

$$y_{PI} = \frac{1}{f(D)} X(x)$$

When the forcing function is $X(x) = \cos(ax)$ or $\sin(ax)$, there is a specific rule for evaluating the operator $\frac{1}{f(D)}$.

The rule states that you can replace every occurrence of D^2 in the polynomial $f(D)$ with $-a^2$. So, if $X(x) = \cos(ax)$, then:

$$y_{PI} = \frac{1}{f(D^2)} \cos(ax) = \frac{1}{f(-a^2)} \cos(ax)$$

This rule is only valid provided that the denominator does not become zero, i.e., $f(-a^2) \neq 0$. If $f(-a^2) = 0$, it is a "case of failure" and a different procedure must be followed.

Step 3: Final Answer:

The formula for the particular integral is $\frac{1}{f(-a^2)} \cos ax$, provided $f(-a^2) \neq 0$. This corresponds to option (A).

Quick Tip

The most crucial part of this rule is the substitution: replace D^2 with $-a^2$. Note the negative sign is outside the square. For example, if the function is $\cos(3x)$, you replace D^2 with $-(3^2) = -9$, not $(-3)^2 = 9$.

Physics

51. If the unit of mass is 1 Kg, the unit of length is 1m and the unit of time is 1 minute, the unit of pressure in Nm^{-2} is

- (A) $1/60$
- (B) 60
- (C) $1/3600$

(D) 3600

Correct Answer: (C) 1/3600

Solution:

Step 1: Understanding the Question:

We are asked to find the value of 1 unit of pressure in a new system of units and express it in the standard SI unit for pressure (Nm^{-2} or Pascal).

Step 2: Key Formula or Approach:

First, we need the dimensional formula for pressure.

$$\text{Pressure (P)} = \frac{\text{Force}}{\text{Area}} = \frac{\text{Mass} \times \text{Acceleration}}{\text{Area}}$$

The dimensional formula is:

$$[P] = \frac{[M][LT^{-2}]}{[L^2]} = [ML^{-1}T^{-2}]$$

Step 3: Detailed Explanation:

In the standard SI system: - Unit of Mass ($[M]$): 1 kg - Unit of Length ($[L]$): 1 m - Unit of Time ($[T]$): 1 s The SI unit of pressure is $1 \text{ kg} \cdot \text{m}^{-1} \cdot \text{s}^{-2}$, which is 1 Nm^{-2} .

In the new system: - Unit of Mass ($[M']$): 1 kg - Unit of Length ($[L']$): 1 m - Unit of Time ($[T']$): 1 minute = 60 s

Let's express 1 unit of pressure in the new system using its fundamental units:

$$1 \text{ unit of new pressure} = 1 [M'] [L']^{-1} [T']^{-2}$$

Now, we substitute the SI equivalents for these new units:

$$\begin{aligned} &= (1 \text{ kg})(1 \text{ m})^{-1}(1 \text{ minute})^{-2} \\ &= (1 \text{ kg})(1 \text{ m})^{-1}(60 \text{ s})^{-2} \\ &= 1 \text{ kg} \cdot \text{m}^{-1} \cdot \frac{1}{60^2} \text{ s}^{-2} \\ &= \frac{1}{3600} (\text{kg} \cdot \text{m}^{-1} \cdot \text{s}^{-2}) \end{aligned}$$

Since $1 \text{ kg} \cdot \text{m}^{-1} \cdot \text{s}^{-2} = 1 \text{ Nm}^{-2}$, we have:

$$= \frac{1}{3600} \text{ Nm}^{-2}$$

Step 4: Final Answer:

One unit of pressure in the new system is equal to $1/3600 \text{ Nm}^{-2}$. This corresponds to option (C).

 Quick Tip

When converting units, write down the dimensional formula. Then, substitute the conversion factors for each fundamental unit into the formula. Pay close attention to the powers on each dimension.

52. MLT^{-1} is the dimensional formula for

- (A) Speed
- (B) Acceleration
- (C) Impulse
- (D) Force

Correct Answer: (C) Impulse

Solution:

Step 1: Understanding the Question:

We need to identify which of the given physical quantities has the dimensional formula $[MLT^{-1}]$.

Step 2: Detailed Explanation:

Let's find the dimensional formula for each option:

(A) Speed:

Speed = $\frac{\text{Distance}}{\text{Time}}$. Dimensions are $\frac{[L]}{[T]} = [LT^{-1}]$. This is incorrect.

(B) Acceleration:

Acceleration = $\frac{\text{Velocity}}{\text{Time}}$. Dimensions are $\frac{[LT^{-1}]}{[T]} = [LT^{-2}]$. This is incorrect.

(C) Impulse:

Impulse = Force $\times$ Time.

First, find the dimensions of Force: Force = Mass $\times$ Acceleration = $[M] \times [LT^{-2}] = [MLT^{-2}]$.

Now, dimensions of Impulse = $[MLT^{-2}] \times [T] = [MLT^{-1}]$. This is correct.

(Alternatively, Impulse = Change in Momentum. Momentum = Mass $\times$ Velocity = $[M] \times [LT^{-1}] = [MLT^{-1}]$).

(D) Force:

As calculated above, the dimensions of Force are $[MLT^{-2}]$. This is incorrect.

Step 3: Final Answer:

The dimensional formula $[MLT^{-1}]$ corresponds to Impulse. This is option (C).

💡 Quick Tip

Remembering the dimensional formulas for Force (MLT^{-2}) and Energy/Work (ML^2T^{-2}) is very useful, as many other quantities can be quickly derived from them. For example, Impulse = Force $\times$ Time, so you just multiply the dimensions of Force by T.

53. If $|\vec{A} \times \vec{B}| = \sqrt{3}(\vec{A} \cdot \vec{B})$ then the value of $|\vec{A} + \vec{B}|$ is

- (A) $(A^2 + B^2 + AB)^{1/2}$
- (B) $(A^2 + B^2 + \frac{AB}{\sqrt{2}})^{1/2}$
- (C) $A + B$
- (D) $(A^2 + B^2 + \sqrt{3}AB)^{1/2}$

Correct Answer: (A) $(A^2 + B^2 + AB)^{1/2}$

Solution:

Step 1: Understanding the Question:

We are given a relationship between the magnitudes of the cross product and the dot product of two vectors. We need to find the magnitude of their sum.

Step 2: Key Formula or Approach:

1. Use the definitions of cross product and dot product to find the angle θ between the vectors.

$$|\vec{A} \times \vec{B}| = AB \sin \theta$$

$$\vec{A} \cdot \vec{B} = AB \cos \theta$$

2. Use the formula for the magnitude of the sum of two vectors.

$$|\vec{A} + \vec{B}| = \sqrt{A^2 + B^2 + 2AB \cos \theta}$$

(Here, A and B represent the magnitudes $|\vec{A}|$ and $|\vec{B}|$).

Step 3: Detailed Explanation:

Find the angle θ :

From the given condition:

$$AB \sin \theta = \sqrt{3}(AB \cos \theta)$$

Assuming $A \neq 0$ and $B \neq 0$, we can divide by $AB \cos \theta$:

$$\frac{\sin \theta}{\cos \theta} = \tan \theta = \sqrt{3}$$

This means the angle between the vectors is $\theta = 60^\circ$.

Find the magnitude of the sum:

Now substitute the value of $\cos \theta$ into the magnitude formula.

$$\cos(60^\circ) = \frac{1}{2}$$

$$|\vec{A} + \vec{B}|^2 = A^2 + B^2 + 2AB \cos(60^\circ)$$

$$|\vec{A} + \vec{B}|^2 = A^2 + B^2 + 2AB \left(\frac{1}{2}\right)$$

$$|\vec{A} + \vec{B}|^2 = A^2 + B^2 + AB$$

Taking the square root of both sides:

$$|\vec{A} + \vec{B}| = \sqrt{A^2 + B^2 + AB} = (A^2 + B^2 + AB)^{1/2}$$

Step 4: Final Answer:

The value of $|\vec{A} + \vec{B}|$ is $(A^2 + B^2 + AB)^{1/2}$. This corresponds to option (A).

💡 Quick Tip

The ratio $\frac{|\vec{A} \times \vec{B}|}{\vec{A} \cdot \vec{B}}$ is a direct way to find the tangent of the angle between two vectors. This is a common pattern in competitive exam questions.

54. Of the vectors given below, the parallel vectors are

$$\vec{A} = 6\hat{i} + 8\hat{j} \quad \vec{B} = 210\hat{i} + 280\hat{k} \quad \vec{C} = 5.1\hat{i} + 6.8\hat{j} \quad \vec{D} = 3.6\hat{i} + 8\hat{j} + 48\hat{k}$$

- (A) $\vec{A}$ and $\vec{C}$
- (B) $\vec{A}$ and $\vec{B}$
- (C) $\vec{A}$ and $\vec{D}$
- (D) $\vec{C}$ and $\vec{D}$

Correct Answer: (A) $\vec{A}$ and $\vec{C}$

Solution:

Step 1: Understanding the Question:

We need to identify which pair of vectors from the given list are parallel to each other.

Step 2: Key Formula or Approach:

Two non-zero vectors are parallel if one is a scalar multiple of the other. That is, $\vec{V}_1 = k\vec{V}_2$ for some scalar k .

This implies that the ratio of their corresponding components must be constant. For two vectors $\vec{V}_1 = a_1\hat{i} + b_1\hat{j} + c_1\hat{k}$ and $\vec{V}_2 = a_2\hat{i} + b_2\hat{j} + c_2\hat{k}$, they are parallel if $\frac{a_1}{a_2} = \frac{b_1}{b_2} = \frac{c_1}{c_2}$.

Step 3: Detailed Explanation:

Let's check the pairs:

$\vec{A}$ and $\vec{B}$: $\vec{A}$ has no $\hat{k}$ component, while $\vec{B}$ has no $\hat{j}$ component. They cannot be parallel.

$\vec{A}$ and $\vec{C}$: Both vectors only have $\hat{i}$ and $\hat{j}$ components. Let's check the ratio of their components.

$$\frac{C_x}{A_x} = \frac{5.1}{6} \quad \text{and} \quad \frac{C_y}{A_y} = \frac{6.8}{8}$$

Let's simplify the ratios:

$$\frac{5.1}{6} = \frac{51}{60} = \frac{17 \times 3}{20 \times 3} = \frac{17}{20} = 0.85$$
$$\frac{6.8}{8} = \frac{68}{80} = \frac{17 \times 4}{20 \times 4} = \frac{17}{20} = 0.85$$

Since the ratios are equal, $\vec{A}$ and $\vec{C}$ are parallel. Specifically, $\vec{C} = 0.85\vec{A}$.

$\vec{A}$ and $\vec{D}$: $\vec{A}$ has no $\hat{k}$ component, but $\vec{D}$ does. They cannot be parallel.

$\vec{C}$ and $\vec{D}$: $\vec{C}$ has no $\hat{k}$ component, but $\vec{D}$ does. They cannot be parallel.

Step 4: Final Answer:

The parallel vectors are $\vec{A}$ and $\vec{C}$. This corresponds to option (A).

💡 Quick Tip

To quickly check for parallel vectors, first look for vectors that exist in the same plane (e.g., both are in the xy-plane). Then, check if the ratios of their components are the same. You can often see the relationship by factoring out common terms. Here, $\vec{A} = 2(3\hat{i} + 4\hat{j})$ and $\vec{C} = 1.7(3\hat{i} + 4\hat{j})$, making the parallel relationship obvious.

55. The position x of a particle with respect to time 't' along x- axis is given by $x = 9t^2 - t^3$ where x is in metres and t in seconds. The position of this particle when it achieves maximum speed along the x direction is

- (A) 24 m
- (B) 32 m
- (C) 54 m
- (D) 81 m

Correct Answer: (C) 54 m

Solution:

Step 1: Understanding the Question:

We are given the position of a particle as a function of time. We need to find its position at the instant its speed is maximum. Since motion is along the x-axis, speed is the magnitude of velocity.

Step 2: Key Formula or Approach:

1. Find the velocity function, $v(t)$, by differentiating the position function $x(t)$ with respect to time.
2. Find the acceleration function, $a(t)$, by differentiating the velocity function $v(t)$.
3. To find the time of maximum velocity, set the acceleration $a(t)$ to zero and solve for t .
4. Substitute this value of t back into the original position function $x(t)$.

Step 3: Detailed Explanation:**1. Find velocity $v(t)$:**

$$x(t) = 9t^2 - t^3$$
$$v(t) = \frac{dx}{dt} = \frac{d}{dt}(9t^2 - t^3) = 18t - 3t^2$$

2. Find acceleration $a(t)$:

$$a(t) = \frac{dv}{dt} = \frac{d}{dt}(18t - 3t^2) = 18 - 6t$$

3. Find time of maximum velocity:

Set $a(t) = 0$:

$$18 - 6t = 0$$

$$6t = 18$$

$$t = 3 \text{ s}$$

(To verify it's a maximum, we can check the second derivative of velocity, which is $\frac{da}{dt} = -6$. Since it's negative, the velocity is indeed at a maximum).

4. Find position at $t = 3$ s:

Substitute $t = 3$ into the position equation:

$$x(3) = 9(3)^2 - (3)^3$$

$$x(3) = 9(9) - 27$$

$$x(3) = 81 - 27 = 54 \text{ m}$$

Step 4: Final Answer:

The position of the particle when it achieves maximum speed is 54 m. This corresponds to option (C).

💡 Quick Tip

This is a standard maxima/minima problem from calculus applied to physics. Remember the sequence: position $\xrightarrow{\text{differentiate}}$ velocity $\xrightarrow{\text{differentiate}}$ acceleration. To find the maximum/minimum of a quantity, differentiate it and set the result to zero.

56. A ball is projected vertically up with a velocity of 40 ms^{-1} from ground. At the same time another ball is dropped from a height of 100 m. The magnitudes of their velocities are equal after

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

Correct Answer: (B) 2 s

Solution:

Step 1: Understanding the Question:

One ball goes up and another comes down. We need to find the time when their speeds are equal. Let's take the upward direction as positive and use $g \approx 10 \text{ ms}^{-2}$ as the options are integers.

Step 2: Key Formula or Approach:

We will use the first equation of motion, $v = u + at$, for both balls. The acceleration for both is $a = -g$.

- Ball A (projected up): $u_A = +40 \text{ ms}^{-1}$
- Ball B (dropped): $u_B = 0 \text{ ms}^{-1}$

Step 3: Detailed Explanation:

Velocity of Ball A at time t:

$$v_A(t) = u_A + at = 40 - gt$$

Velocity of Ball B at time t:

$$v_B(t) = u_B + at = 0 - gt = -gt$$

The negative sign indicates it's moving downward.

We are interested in when the **magnitudes** of their velocities are equal. The magnitude of velocity is speed.

$$|v_A(t)| = |v_B(t)|$$

The magnitude of v_B is always $|-gt| = gt$.

The velocity of ball A is positive while it's going up and negative when it's coming down. It reaches its peak when $v_A = 0$, which occurs at $t = 40/g = 40/10 = 4 \text{ s}$. We should assume the time we are looking for is $t < 4 \text{ s}$.

In this case, $v_A(t) = 40 - gt$ is positive, so $|v_A(t)| = 40 - gt$.

Set the magnitudes equal:

$$40 - gt = gt$$

$$40 = 2gt$$

$$t = \frac{40}{2g} = \frac{20}{g}$$

Using $g = 10 \text{ ms}^{-2}$:

$$t = \frac{20}{10} = 2 \text{ s}$$

Since $t = 2 \text{ s}$ is less than 4 s , our assumption was correct.

Step 4: Final Answer:

The magnitudes of their velocities will be equal after 2 seconds. This corresponds to option (B).

💡 Quick Tip

In problems involving magnitudes of velocities under gravity, be mindful of the direction of motion. The speed of the upwardly projected body decreases, while the speed of the dropped body increases. They will meet at a point in time.

57. Two stones are projected with the same speed but making different angles with the horizontal. Their horizontal ranges are equal. The angle of projection of one is $\pi/3$ and the maximum height reached by it is 102 metres. Then the maximum height reached by the other in metres is

- (A) 336
- (B) 224
- (C) 56
- (D) 34

Correct Answer: (D) 34

Solution:

Step 1: Understanding the Question:

We have two projectiles with the same initial speed but different angles, resulting in the same horizontal range. We are given the angle and maximum height for one projectile and asked to find the maximum height of the other.

Step 2: Key Formula or Approach:

- For a given initial speed, the horizontal range is the same for two projection angles θ_1 and θ_2 if they are complementary, i.e., $\theta_1 + \theta_2 = 90^\circ$.
- The formula for maximum height is $H = \frac{u^2 \sin^2 \theta}{2g}$.

Step 3: Detailed Explanation:

- Find the angle of the second stone:**

The angle of the first stone is $\theta_1 = \pi/3$ radians, which is 60° .

Since the ranges are equal, the angle of the second stone must be complementary to the first.

$$\theta_2 = 90^\circ - \theta_1 = 90^\circ - 60^\circ = 30^\circ$$

2. Relate the maximum heights:

Let H_1 be the height for θ_1 and H_2 be the height for θ_2 .

$$H_1 = \frac{u^2 \sin^2 \theta_1}{2g} = 102 \text{ m}$$

$$H_2 = \frac{u^2 \sin^2 \theta_2}{2g}$$

To find H_2 , we can take a ratio of the two heights to eliminate the common term $\frac{u^2}{2g}$.

$$\frac{H_2}{H_1} = \frac{\sin^2 \theta_2}{\sin^2 \theta_1}$$

$$H_2 = H_1 \left(\frac{\sin \theta_2}{\sin \theta_1} \right)^2$$

Substitute the known values:

$$H_2 = 102 \left(\frac{\sin 30^\circ}{\sin 60^\circ} \right)^2$$

We know $\sin 30^\circ = 1/2$ and $\sin 60^\circ = \sqrt{3}/2$.

$$H_2 = 102 \left(\frac{1/2}{\sqrt{3}/2} \right)^2 = 102 \left(\frac{1}{\sqrt{3}} \right)^2$$

$$H_2 = 102 \times \frac{1}{3} = 34 \text{ m}$$

Step 4: Final Answer:

The maximum height reached by the other stone is 34 metres. This corresponds to option (D).

💡 Quick Tip

For two complementary projection angles, the ratio of maximum heights is $H_2/H_1 = \tan^2 \theta_2$. Since $\theta_2 = 30^\circ$, $\tan 30^\circ = 1/\sqrt{3}$. So, $H_2 = H_1 \times (1/\sqrt{3})^2 = H_1/3$. This is a useful shortcut.

58. A projectile is thrown into air with velocity u at an angle θ to the horizontal. The time at which its direction of motion is perpendicular to its initial direction is

(A) $\frac{u}{g \sin \theta}$

(B) $\frac{u}{g \cos \theta}$

(C) $\frac{u}{g \tan \theta}$

(D) $\frac{u}{g \cot \theta}$

Correct Answer: (A) $\frac{u}{g \sin \theta}$

Solution:

Step 1: Understanding the Question:

We need to find the time t when the velocity vector of a projectile is perpendicular to its initial velocity vector.

Step 2: Key Formula or Approach:

Two vectors are perpendicular if their dot product is zero. We will write the initial velocity vector and the velocity vector at time t , and then set their dot product to zero.

Step 3: Detailed Explanation:

The initial velocity vector is:

$$\vec{v}_0 = (u \cos \theta)\hat{i} + (u \sin \theta)\hat{j}$$

The velocity vector at time t is:

$$\vec{v}(t) = (u \cos \theta)\hat{i} + (u \sin \theta - gt)\hat{j}$$

For the vectors to be perpendicular, their dot product must be zero:

$$\vec{v}_0 \cdot \vec{v}(t) = 0$$

$$((u \cos \theta)\hat{i} + (u \sin \theta)\hat{j}) \cdot ((u \cos \theta)\hat{i} + (u \sin \theta - gt)\hat{j}) = 0$$

Calculate the dot product:

$$(u \cos \theta)(u \cos \theta) + (u \sin \theta)(u \sin \theta - gt) = 0$$

$$u^2 \cos^2 \theta + u^2 \sin^2 \theta - ugt \sin \theta = 0$$

Factor out u^2 and use the identity $\cos^2 \theta + \sin^2 \theta = 1$:

$$u^2(\cos^2 \theta + \sin^2 \theta) - ugt \sin \theta = 0$$

$$u^2(1) - ugt \sin \theta = 0$$

$$u^2 = ugt \sin \theta$$

Assuming $u \neq 0$, we can divide by u :

$$u = gt \sin \theta$$

Now, solve for t :

$$t = \frac{u}{g \sin \theta}$$

Step 4: Final Answer:

The time is $\frac{u}{g \sin \theta}$. This corresponds to option (A).

💡 Quick Tip

This problem is a classic application of the dot product to check for orthogonality (perpendicularity). Whenever a question involves angles between vectors (especially 90 degrees), the dot product is usually the most efficient tool to use.

59. When a bicycle is in motion and pedalled, the force of friction exerted by ground on the two wheels is such that it acts

- (A) In the backward direction on the front wheel and in the forward direction on the rear wheel
- (B) In the forward direction on the front wheel and in the backward direction on the rear wheel
- (C) In the backward direction on both the front and rear wheels
- (D) In the forward direction on both the front and rear wheels

Correct Answer: (A) In the backward direction on the front wheel and in the forward direction on the rear wheel

Solution:

Step 1: Understanding the Question:

We need to determine the direction of the force of static friction on the front and rear wheels of a bicycle that is being pedalled and moving forward.

Step 2: Detailed Explanation:

Rear Wheel (Driving Wheel):

The cyclist applies a torque via the pedals and chain to the rear wheel. This torque makes the bottom of the rear wheel try to slide or push the ground **backwards**.

According to Newton's third law, the ground exerts an equal and opposite force on the wheel. This reaction force is the force of static friction, and it points in the **forward** direction. This forward frictional force is what propels the bicycle.

Front Wheel (Driven Wheel):

The front wheel is not connected to the pedals; it rolls freely. The bicycle's frame pushes the axle of the front wheel forward. Due to this push, the bottom of the front wheel has a tendency to slide forward over the ground.

The force of static friction from the ground opposes this tendency of sliding. Therefore, the friction on the front wheel acts in the **backward** direction. This backward friction is a form of rolling resistance.

Step 3: Final Answer:

The frictional force on the rear wheel is forward, and on the front wheel, it is backward. This corresponds to option (A).

 Quick Tip

For any vehicle, the friction on a driven (powered) wheel is in the direction of motion, as it provides the propulsion. The friction on a non-driven (freely rolling) wheel opposes the motion and acts as a retarding force. The same logic applies to the front-wheel drive and rear-wheel drive cars.

60. Two blocks of masses 4 Kg and 2 Kg are connected by a heavy string and placed on rough horizontal plane. The 2 Kg block is pulled with a constant force F. The coefficient of friction between the blocks and the ground is 0.5. The value of F so that tension in the string is constant throughout during the motion of the blocks is

- (A) 40 N
- (B) 30 N
- (C) 50 N
- (D) 60 N

Correct Answer: (B) 30 N

Solution:

Step 1: Understanding the Question:

We have a system of two blocks connected by a string on a rough surface. A force F pulls the 2 kg block. We need to find the value of F that ensures the tension in the string remains constant during motion.

The phrase "tension in the string is constant throughout" implies two things: 1. The tension is uniform along the length of the string. This would be true for a massless string, or for a massive ("heavy") string if it is not accelerating. 2. The tension does not change with time. Since the string is described as "heavy", the only way for the tension to be uniform along its length during motion is if the acceleration of the system is zero. This means the blocks move at a constant velocity.

Step 2: Key Formula or Approach:

We will use the condition for motion with constant velocity ($a = 0$), which means the net force on the system is zero. The force of kinetic friction is given by $f_k = \mu N$, where N is the normal force. On a horizontal surface, $N = mg$. We will consider $g = 10 \text{ m/s}^2$.

Step 3: Detailed Explanation:

Let's consider the two blocks and the string as a single system. For the system to move at a constant velocity, the applied force F must be exactly equal to the total frictional force acting on the system.

First, calculate the frictional force on the 4 kg block (m_1):

$$f_1 = \mu \times N_1 = \mu \times m_1 g$$

$$f_1 = 0.5 \times 4 \text{ kg} \times 10 \text{ m/s}^2 = 20 \text{ N}$$

Next, calculate the frictional force on the 2 kg block (m_2):

$$f_2 = \mu \times N_2 = \mu \times m_2 g$$

$$f_2 = 0.5 \times 2 \text{ kg} \times 10 \text{ m/s}^2 = 10 \text{ N}$$

The total frictional force on the system is the sum of the individual frictional forces:

$$f_{total} = f_1 + f_2 = 20 \text{ N} + 10 \text{ N} = 30 \text{ N}$$

For the system to move with constant velocity ($a = 0$), the net horizontal force must be zero.

$$F_{net} = F_{applied} - f_{total} = 0$$

$$F - 30 \text{ N} = 0$$

$$F = 30 \text{ N}$$

Step 4: Final Answer:

The value of the force F required to move the blocks at a constant velocity, thus ensuring constant tension, is 30 N. This corresponds to option (B).

💡 Quick Tip

In problems with connected blocks, the condition "constant tension" usually implies constant acceleration. If the connecting string is described as "heavy" or "massive," constant tension along its length can only be achieved if the acceleration is zero (i.e., constant velocity).

61. In a hydroelectric power station, the height of the dam is 10 m. How many kilograms of water must fall per second on the blades of a turbine in order to generate 1 MW of electrical power? [$g = 10 \text{ m/s}^2$].

- (A) 10^3 Kg/s
- (B) 10^4 Kg/s
- (C) 10^5 Kg/s
- (D) 10^6 Kg/s

Correct Answer: (B) 10^4 Kg/s

Solution:

Step 1: Understanding the Question:

The question asks for the mass flow rate of water ($\frac{m}{t}$) required to generate a specific amount of power from a given height. This is a problem of energy conversion.

Step 2: Key Formula or Approach:

The potential energy (PE) of a mass m of water at a height h is given by $PE = mgh$.

Power (P) is the rate at which energy is converted, so $P = \frac{PE}{t}$.

Substituting the formula for PE, we get:

$$P = \frac{mgh}{t} = \left(\frac{m}{t}\right) gh$$

We need to solve for the mass flow rate, $\frac{m}{t}$.

Step 3: Detailed Explanation:

First, let's list the given values in SI units: - Power, $P = 1 \text{ MW} = 1 \times 10^6 \text{ Watts (or J/s)}$.

- Height, $h = 10 \text{ m}$.

- Acceleration due to gravity, $g = 10 \text{ m/s}^2$.

Rearrange the power formula to solve for the mass flow rate $\frac{m}{t}$:

$$\frac{m}{t} = \frac{P}{gh}$$

Substitute the given values into the formula:

$$\frac{m}{t} = \frac{1 \times 10^6}{10 \times 10}$$

$$\frac{m}{t} = \frac{10^6}{100} = \frac{10^6}{10^2}$$

$$\frac{m}{t} = 10^{6-2} = 10^4 \text{ Kg/s}$$

Step 4: Final Answer:

A mass of 10^4 kilograms of water must fall per second. This corresponds to option (B).

💡 Quick Tip

This calculation assumes 100% efficiency in converting potential energy to electrical energy. In real-world scenarios, the efficiency would be less than 100%, and a larger mass of water would be required. However, for exam problems, assume 100% efficiency unless stated otherwise.

62. The kinetic energy at the highest point of the trajectory of a projectile is 200 J. If the mass of the projectile is 1 Kg and the maximum height reached by it is 20 m, then velocity of the projectile from the ground is

- (A) 20 m/s
- (B) 10 m/s
- (C) $20\sqrt{2}$ m/s

(D) $10\sqrt{2}$ m/s

Correct Answer: (C) $20\sqrt{2}$ m/s

Solution:

Step 1: Understanding the Question:

We are given the kinetic energy at the highest point and the maximum height of a projectile. We need to find its initial launch velocity using the principle of conservation of energy.

Step 2: Key Formula or Approach:

According to the principle of conservation of mechanical energy, the total energy of the projectile remains constant throughout its flight (ignoring air resistance).

Total Energy at launch = Total Energy at highest point

$$E_{initial} = E_{top}$$

$$KE_{initial} + PE_{initial} = KE_{top} + PE_{top}$$

Step 3: Detailed Explanation:

Let's define the energy at each point: - At the launch point (ground): - Height is 0, so $PE_{initial} = 0$. - $KE_{initial} = \frac{1}{2}mu^2$, where u is the initial velocity.

- At the highest point: - We are given $KE_{top} = 200$ J. - The potential energy is $PE_{top} = mgh_{max}$.

Now, let's calculate PE_{top} : - Mass, $m = 1$ Kg. - Max height, $h_{max} = 20$ m. - Let's use $g = 10$ m/s².

$$PE_{top} = (1 \text{ Kg})(10 \text{ m/s}^2)(20 \text{ m}) = 200 \text{ J}$$

Now, apply the conservation of energy principle:

$$KE_{initial} + 0 = KE_{top} + PE_{top}$$

$$\frac{1}{2}mu^2 = 200 \text{ J} + 200 \text{ J} = 400 \text{ J}$$

Substitute the mass $m = 1$ Kg and solve for u :

$$\frac{1}{2}(1)u^2 = 400$$

$$u^2 = 800$$

$$u = \sqrt{800} = \sqrt{400 \times 2} = 20\sqrt{2} \text{ m/s}$$

Step 4: Final Answer:

The initial velocity of the projectile from the ground is $20\sqrt{2}$ m/s. This corresponds to option (C).

 Quick Tip

At the highest point, the projectile's velocity is not zero; only its vertical component is zero. The kinetic energy at the top corresponds purely to the horizontal component of velocity. The total energy is the sum of this kinetic energy and the potential energy at that height.

63. A force applied by an engine on train of mass 2.05×10^6 Kg changes its velocity from 5 m/s to 25 m/s in 5 minutes. The power of the engine is

- (A) 1.025 MW
- (B) 2.05 MW
- (C) 5 MW
- (D) 6 MW

Correct Answer: (B) 2.05 MW

Solution:

Step 1: Understanding the Question:

We need to calculate the power of an engine given the mass of the train, the change in its velocity, and the time taken for this change.

Step 2: Key Formula or Approach:

Power is the rate at which work is done ($P = W/t$).

According to the Work-Energy Theorem, the work done on an object is equal to the change in its kinetic energy ($W = \Delta KE$).

Combining these, we get:

$$P = \frac{\Delta KE}{t} = \frac{KE_{final} - KE_{initial}}{t} = \frac{\frac{1}{2}mv^2 - \frac{1}{2}mu^2}{t}$$

Step 3: Detailed Explanation:

First, list the given values in SI units: - Mass, $m = 2.05 \times 10^6$ Kg. - Initial velocity, $u = 5$ m/s.
- Final velocity, $v = 25$ m/s. - Time, $t = 5$ minutes = $5 \times 60 = 300$ s.

Calculate the change in kinetic energy (ΔKE):

$$\begin{aligned}\Delta KE &= \frac{1}{2}m(v^2 - u^2) \\ \Delta KE &= \frac{1}{2}(2.05 \times 10^6)((25)^2 - (5)^2) \\ \Delta KE &= \frac{1}{2}(2.05 \times 10^6)(625 - 25) \\ \Delta KE &= \frac{1}{2}(2.05 \times 10^6)(600)\end{aligned}$$

$$\Delta KE = 300 \times 2.05 \times 10^6 = 615 \times 10^6 \text{ J}$$

Now, calculate the power:

$$P = \frac{\Delta KE}{t} = \frac{615 \times 10^6 \text{ J}}{300 \text{ s}}$$

$$P = 2.05 \times 10^6 \text{ W}$$

Since $1 \text{ MW} = 10^6 \text{ W}$, the power is 2.05 MW.

Step 4: Final Answer:

The power of the engine is 2.05 MW. This corresponds to option (B).

💡 Quick Tip

This question asks for the average power over the 5-minute interval. Be careful to convert all units to standard SI units (mass in kg, velocity in m/s, time in seconds) before performing calculations.

64. Two identical wires have a fundamental frequency of 100 Hz when kept under the same tension. If the tension of one of the wires is increased by 21%, the number of beats produced is

- (A) 11
- (B) 10
- (C) 9
- (D) 8

Correct Answer: (B) 10

Solution:

Step 1: Understanding the Question:

Two identical wires have the same initial frequency. The tension in one is changed, which changes its frequency. We need to find the resulting beat frequency.

Step 2: Key Formula or Approach:

- The fundamental frequency (f) of a stretched string is given by $f = \frac{1}{2L} \sqrt{\frac{T}{\mu}}$, where L is length, T is tension, and μ is linear mass density. For identical wires, L and μ are constant, so $f \propto \sqrt{T}$.
- The beat frequency is the difference between the two frequencies: $f_{beat} = |f_2 - f_1|$.

Step 3: Detailed Explanation:

Let the initial frequency be $f_1 = 100 \text{ Hz}$ and the initial tension be T_1 .

The tension in the second wire is increased by 21%. The new tension T_2 is:

$$T_2 = T_1 + 0.21T_1 = 1.21T_1$$

Since $f \propto \sqrt{T}$, we can write the relationship between the new frequency f_2 and the old frequency f_1 as:

$$\frac{f_2}{f_1} = \sqrt{\frac{T_2}{T_1}}$$

Substitute the values:

$$\frac{f_2}{100} = \sqrt{\frac{1.21T_1}{T_1}} = \sqrt{1.21} = 1.1$$

Solve for the new frequency f_2 :

$$f_2 = 100 \times 1.1 = 110 \text{ Hz}$$

The frequency of the first wire remains $f_1 = 100 \text{ Hz}$.

Now, calculate the number of beats produced per second (beat frequency):

$$f_{\text{beat}} = |f_2 - f_1| = |110 - 100| = 10 \text{ Hz}$$

Step 4: Final Answer:

The number of beats produced is 10. This corresponds to option (B).

💡 Quick Tip

For small percentage changes in tension ($x\%$), the percentage change in frequency is approximately $(x/2)\%$. Here, the change is large (21%), so the approximation is not accurate. It's better to use the direct square root relationship $f' = f\sqrt{1 + \text{fractional change}}$.

65. A body executing S.H.M. has a maximum velocity of 1 ms^{-1} and a maximum acceleration of 4 ms^{-2} . Its amplitude in metres is:

- (A) 1
- (B) 0.75
- (C) 0.5
- (D) 0.25

Correct Answer: (D) 0.25

Solution:

Step 1: Understanding the Question:

We are given the maximum velocity and maximum acceleration of a body in Simple Harmonic

Motion (S.H.M.) and need to find its amplitude.

Step 2: Key Formula or Approach:

For a body in S.H.M. with amplitude A and angular frequency ω : - Maximum velocity is given by $v_{max} = A\omega$.

- Maximum acceleration is given by $a_{max} = A\omega^2$.

We can solve this system of two equations to find A .

Step 3: Detailed Explanation:

We are given: 1. $v_{max} = A\omega = 1$ 2. $a_{max} = A\omega^2 = 4$

We can express a_{max} in terms of v_{max} :

$$a_{max} = (A\omega)\omega = v_{max} \cdot \omega$$

Substitute the given values to find ω :

$$4 = (1) \cdot \omega$$

$$\omega = 4 \text{ rad/s}$$

Now that we have the angular frequency ω , we can use the equation for maximum velocity to find the amplitude A :

$$v_{max} = A\omega$$

$$1 = A \cdot (4)$$

$$A = \frac{1}{4} = 0.25 \text{ m}$$

Step 4: Final Answer:

The amplitude of the S.H.M. is 0.25 metres. This corresponds to option (D).

💡 Quick Tip

A very quick way to solve this is to take the ratio of the given quantities. $a_{max}/v_{max} = (A\omega^2)/(A\omega) = \omega$. And amplitude $A = v_{max}^2/a_{max} = (A\omega)^2/(A\omega^2) = A^2\omega^2/A\omega^2 = A$. So, $A = (1)^2/4 = 0.25$.

66. A simple pendulum of length l_1 has frequency $\frac{1}{4}$ Hz and another simple pendulum of length l_2 has frequency $\frac{1}{3}$ Hz. Then time period of pendulum of length $(l_1 - l_2)$ is

- (A) 5 s
- (B) 1 s
- (C) $\sqrt{7}$ s
- (D) $\sqrt{12}$ s

Correct Answer: (C) $\sqrt{7}$ s

Solution:

Step 1: Understanding the Question:

We are given the frequencies of two simple pendulums and asked to find the time period of a third pendulum whose length is the difference of the lengths of the first two.

Step 2: Key Formula or Approach:

The time period T of a simple pendulum of length l is given by $T = 2\pi\sqrt{\frac{l}{g}}$.

Frequency f is the reciprocal of the time period, $f = \frac{1}{T}$.

From these, we can relate length and time period: $T^2 = \frac{4\pi^2}{g}l$, which means $l \propto T^2$.

Step 3: Detailed Explanation:

For pendulum 1:

Frequency $f_1 = \frac{1}{4}$ Hz.

Time period $T_1 = \frac{1}{f_1} = 4$ s.

Since $l \propto T^2$, we can write $l_1 = kT_1^2 = k(4^2) = 16k$ for some constant $k = \frac{g}{4\pi^2}$.

For pendulum 2:

Frequency $f_2 = \frac{1}{3}$ Hz.

Time period $T_2 = \frac{1}{f_2} = 3$ s.

Similarly, $l_2 = kT_2^2 = k(3^2) = 9k$.

From this we see that $l_1 > l_2$, so the difference is positive.

For the new pendulum:

The length of the new pendulum is $l_{\text{new}} = l_1 - l_2$.

$$l_{\text{new}} = 16k - 9k = 7k$$

The time period of this new pendulum, T_{new} , is related to its length by $l_{\text{new}} = kT_{\text{new}}^2$.

$$7k = kT_{\text{new}}^2$$

$$T_{\text{new}}^2 = 7$$

$$T_{\text{new}} = \sqrt{7} \text{ s}$$

Step 4: Final Answer:

The time period of the pendulum of length $(l_1 - l_2)$ is $\sqrt{7}$ s. This corresponds to option (C).

💡 Quick Tip

For simple pendulums, the length is directly proportional to the square of the time period ($l \propto T^2$). This relationship is very useful for comparing pendulums without needing to calculate the constant $g/4\pi^2$.

67. A source of sound producing wavelength of 50 cm is moving away from stationary observer with $\frac{1}{5}$ th speed of sound. The wavelength of the sound heard by the observer is

- (A) 70 cm
- (B) 55 cm
- (C) 40 cm
- (D) 60 cm

Correct Answer: (D) 60 cm

Solution:

Step 1: Understanding the Question:

This is a problem about the Doppler effect for sound waves. A sound source is moving away from a stationary observer, and we need to find the apparent wavelength detected by the observer.

Step 2: Key Formula or Approach:

When a source moves away from a stationary observer, the observed frequency (f_o) is lower than the source frequency (f_s). The formula is:

$$f_o = f_s \left(\frac{v}{v + v_s} \right)$$

where v is the speed of sound and v_s is the speed of the source.

Wavelength is related to frequency and speed by $\lambda = v/f$. The source wavelength is $\lambda_s = v/f_s$. The observed wavelength is $\lambda_o = v/f_o$.

Step 3: Detailed Explanation:

We can derive a formula for the apparent wavelength. From the frequency formula, $\frac{f_s}{f_o} = \frac{v + v_s}{v}$. Substituting $f = v/\lambda$:

$$\begin{aligned} \frac{v/\lambda_s}{v/\lambda_o} &= \frac{\lambda_o}{\lambda_s} = \frac{v + v_s}{v} \\ \lambda_o &= \lambda_s \left(\frac{v + v_s}{v} \right) \end{aligned}$$

This shows that when the source moves away, the observed wavelength increases (a "stretch").

Now, substitute the given values: - Source wavelength, $\lambda_s = 50$ cm. - Source speed, $v_s = \frac{1}{5}v$.

$$\begin{aligned} \lambda_o &= 50 \left(\frac{v + \frac{1}{5}v}{v} \right) \\ \lambda_o &= 50 \left(\frac{\frac{6}{5}v}{v} \right) \\ \lambda_o &= 50 \left(\frac{6}{5} \right) = 10 \times 6 = 60 \text{ cm} \end{aligned}$$

Step 4: Final Answer:

The wavelength of the sound heard by the observer is 60 cm. This corresponds to option (D).

💡 Quick Tip

Remember the physical effect: when a source moves away, the waves are "stretched out" behind it, leading to a longer wavelength and a lower frequency (lower pitch). When it moves towards, the waves are "compressed," leading to a shorter wavelength and higher frequency.

68. To have a good sound effect inside a hall

- (A) the hall should not have any sound absorbing material
- (B) the reverberation time has to be maximum
- (C) the reverberation time has to be zero
- (D) the reverberation time has to be optimum

Correct Answer: (D) the reverberation time has to be optimum

Solution:

Step 1: Understanding the Question:

The question asks about the requirement for good acoustics in a hall, specifically regarding reverberation time.

Step 2: Detailed Explanation:

Reverberation is the persistence of sound in a particular space after the original sound is produced. Reverberation time is the time it takes for the sound to "fade away" or decay.

- **Option (A):** If there are no sound-absorbing materials (like carpets, curtains, or acoustic panels), the sound waves will reflect excessively off the hard surfaces (walls, ceiling, floor). This leads to a very long reverberation time, causing echoes and making speech or music sound muddled and unclear. This is undesirable.

- **Option (B):** A maximum reverberation time is the same issue as in option (A) and is not good for sound clarity.

- **Option (C):** A reverberation time of zero would mean that sound dies out instantly. This is characteristic of an anechoic chamber. Such an environment sounds unnatural and "dead," and is not suitable for performances like concerts, where some resonance is needed to enrich the sound.

- **Option (D):** An optimum reverberation time provides a balance. It allows the sound to persist long enough to blend and create a sense of fullness and space (liveness) but is short enough

to ensure that individual sounds and words remain clear and distinct. The ideal reverberation time varies depending on the purpose of the hall (e.g., longer for orchestral music, shorter for lectures).

Step 3: Final Answer:

For good acoustics, the reverberation time must be optimized for the intended use of the hall. This corresponds to option (D).

💡 Quick Tip

Think of the difference between singing in a tiled bathroom (long reverberation) versus a room full of soft furniture (short reverberation). Neither is perfect for a concert. Good concert hall acoustics are a carefully engineered compromise between these extremes.

69. If the pressure of an ideal gas contained in a closed vessel is increased by 0.5%, the increase in temperature is 2°C. The initial temperature of the gas is

- (A) 27°C
- (B) 127°C
- (C) 300°C
- (D) 400°C

Correct Answer: (B) 127°C

Solution:

Step 1: Understanding the Question:

For an ideal gas in a closed vessel (constant volume), a small percentage change in pressure results in a specific absolute change in temperature. We need to find the initial temperature.

Step 2: Key Formula or Approach:

For an ideal gas at constant volume (isochoric process), Gay-Lussac's Law applies:

$$\frac{P}{T} = \text{constant}$$

where T must be in Kelvin. This implies $\frac{P_1}{T_1} = \frac{P_2}{T_2}$.

For small changes, we can use the differential form: $\frac{dP}{P} = \frac{dT}{T}$. This can be approximated as $\frac{\Delta P}{P} \approx \frac{\Delta T}{T}$.

Step 3: Detailed Explanation:

Let the initial pressure be P_1 and initial temperature be T_1 .

The pressure is increased by 0.5%, so the change in pressure is $\Delta P = 0.005P_1$.

The fractional change in pressure is $\frac{\Delta P}{P_1} = 0.005$.

The increase in temperature is $\Delta T = 2^\circ\text{C}$. A change of 2°C is equal to a change of 2 Kelvin, so $\Delta T = 2\text{ K}$.

Using the approximation for small changes:

$$\frac{\Delta P}{P_1} = \frac{\Delta T}{T_1}$$
$$0.005 = \frac{2}{T_1}$$

Solve for the initial temperature T_1 in Kelvin:

$$T_1 = \frac{2}{0.005} = \frac{2}{5/1000} = \frac{2000}{5} = 400\text{ K}$$

The options are given in degrees Celsius. We need to convert the initial temperature from Kelvin to Celsius:

$$T(^{\circ}\text{C}) = T(\text{K}) - 273.15$$

$$T_1(^{\circ}\text{C}) = 400 - 273 = 127^{\circ}\text{C}$$

Step 4: Final Answer:

The initial temperature of the gas is 127°C . This corresponds to option (B).

 Quick Tip

Always remember to use absolute temperature (Kelvin) in gas law calculations. Also, a change in temperature (ΔT) has the same numerical value in Celsius and Kelvin, but the absolute temperatures are different.

70. During the free expansion of an ideal gas, which of the following physical quantity remains constant

- (A) Temperature
- (B) Pressure
- (C) Volume
- (D) Ratio of pressure to volume

Correct Answer: (A) Temperature

Solution:

Step 1: Understanding the Question:

We need to identify the physical quantity that is conserved during the free expansion of an

ideal gas.

Step 2: Key Formula or Approach:

Free expansion (also known as Joule expansion) is a process where a gas expands into a vacuum. The key characteristics are: 1. No work is done ($W = 0$) because the gas expands against zero external pressure.

2. The system is typically considered to be thermally isolated, so there is no heat exchange ($Q = 0$).

We analyze the consequences using the First Law of Thermodynamics: $\Delta U = Q - W$.

Step 3: Detailed Explanation:

Applying the First Law of Thermodynamics to the free expansion process:

$$\Delta U = 0 - 0 = 0$$

This means the change in the internal energy of the gas is zero. The internal energy remains constant.

For an **ideal gas**, the internal energy is a function of temperature only ($U = f(T)$).

Therefore, if $\Delta U = 0$, it implies that the change in temperature is also zero ($\Delta T = 0$).

So, the temperature of the ideal gas remains constant during free expansion.

- **Pressure** decreases because the gas occupies a larger volume. - **Volume** increases by definition of expansion. - The **ratio of pressure to volume** (P/V) also changes.

Step 4: Final Answer:

During the free expansion of an ideal gas, the temperature remains constant. This corresponds to option (A).

 Quick Tip

Remember the three key conditions for free expansion of an ideal gas: $W = 0$, $Q = 0$, and therefore $\Delta U = 0$. For an ideal gas, $\Delta U = 0$ directly implies $\Delta T = 0$.

71. The specific heat at constant volume for a monoatomic gas is 0.075 cal/kg/K and its gram molecular specific heat is 3 cal/mol/K . Then mass of one atom of that gas is

- (A) $6.67 \times 10^{-23} \text{ gm}$
- (B) $6.67 \times 10^{23} \text{ gm}$
- (C) $2 \times 10^{-23} \text{ gm}$
- (D) $2 \times 10^{23} \text{ gm}$

Correct Answer: (A) $6.67 \times 10^{-23} \text{ gm}$

Solution:**Step 1: Understanding the Question:**

We are given the specific heat per unit mass (c_V) and the specific heat per mole (C_V) for a gas. We need to find the mass of a single atom.

Step 2: Key Formula or Approach:

1. Relate molar specific heat (C_V) and specific heat (c_V) using the molar mass (M): $C_V = M \cdot c_V$.
2. The molar mass (M) is the mass of one mole of the substance.
3. One mole contains Avogadro's number ($N_A \approx 6.022 \times 10^{23}$) of atoms.
4. The mass of a single atom (m_{atom}) is the molar mass divided by Avogadro's number: $m_{atom} = \frac{M}{N_A}$.

Step 3: Detailed Explanation:

The given values are: - Molar specific heat, $C_V = 3 \text{ cal/mol/K}$. - Specific heat, $c_V = 0.075 \text{ cal/kg/K}$.

There seems to be a unit inconsistency in the problem statement, as a molar mass calculated from these values would be extremely large. It is highly probable that the specific heat was intended to be given in units of cal/g/K. Let's assume this correction. Corrected specific heat, $c_V = 0.075 \text{ cal/g/K}$.

Now, calculate the molar mass M in grams per mole:

$$M = \frac{C_V}{c_V} = \frac{3 \text{ cal/mol/K}}{0.075 \text{ cal/g/K}}$$
$$M = \frac{3}{75/1000} = \frac{3000}{75} = 40 \text{ g/mol}$$

This molar mass (40 g/mol) corresponds to the element Argon, which is a monoatomic gas, so the values are consistent.

Finally, calculate the mass of a single atom:

$$m_{atom} = \frac{M}{N_A} = \frac{40 \text{ g/mol}}{6.022 \times 10^{23} \text{ atoms/mol}}$$
$$m_{atom} \approx \frac{40}{6} \times 10^{-23} \text{ g} \approx 6.67 \times 10^{-23} \text{ g}$$

Step 4: Final Answer:

The mass of one atom of the gas is approximately $6.67 \times 10^{-23} \text{ gm}$. This corresponds to option (A).

 Quick Tip

In physics and chemistry problems, always perform a unit analysis. If the initial calculation leads to an unreasonable result (like a molar mass of 40 kg/mol), re-check the units for a likely typo. Grams are often used instead of kilograms in chemistry contexts.

72. A rigid diatomic ideal gas undergoes an adiabatic process at room temperature. The relation between temperature and volume of this process is $TV^x = \text{constant}$. Then x is

- (A) $5/3$
- (B) $2/5$
- (C) $2/3$
- (D) $3/5$

Correct Answer: (B) $2/5$

Solution:

Step 1: Understanding the Question:

We need to find the exponent x in the temperature-volume relationship for an adiabatic process involving a rigid diatomic ideal gas.

Step 2: Key Formula or Approach:

1. The primary equation for an adiabatic process is $PV^\gamma = \text{constant}$, where γ is the adiabatic index (ratio of specific heats, C_P/C_V).
2. For an ideal gas, $P = \frac{nRT}{V}$.
3. The adiabatic index $\gamma = 1 + \frac{2}{f}$, where f is the number of degrees of freedom.

Step 3: Detailed Explanation:

First, determine the degrees of freedom for a rigid diatomic gas at room temperature. It has 3 translational and 2 rotational degrees of freedom, so $f = 3 + 2 = 5$. Vibrational modes are not excited at room temperature.

Next, calculate the adiabatic index γ :

$$\gamma = 1 + \frac{2}{f} = 1 + \frac{2}{5} = \frac{7}{5}$$

Now, convert the PV relationship to a TV relationship. Start with $PV^\gamma = K_1$ (constant). Substitute $P = \frac{nRT}{V}$ from the ideal gas law:

$$\left(\frac{nRT}{V}\right) V^\gamma = K_1$$
$$nRTV^{\gamma-1} = K_1$$

Since n and R are constants, we can combine them into a new constant $K_2 = K_1/(nR)$.

$$TV^{\gamma-1} = K_2$$

This is the required relationship. By comparing $TV^{\gamma-1} = \text{constant}$ with the given form $TV^x = \text{constant}$, we can see that:

$$x = \gamma - 1$$

Substitute the value of γ :

$$x = \frac{7}{5} - 1 = \frac{7-5}{5} = \frac{2}{5}$$

Step 4: Final Answer:

The value of x is $2/5$. This corresponds to option (B).

💡 Quick Tip

Memorize the three forms of the adiabatic equation for an ideal gas: 1. $PV^\gamma = \text{constant}$ 2. $TV^{\gamma-1} = \text{constant}$ 3. $P^{1-\gamma}T^\gamma = \text{constant}$ Knowing these directly saves time on derivations during an exam.

73. A carnot engine having an efficiency of $\frac{1}{10}$ as heat engine, is used as a refrigerator. If the work done on the system is 10 J, the amount of energy absorbed from the reservoir at lower temperature is

- (A) 100 J
- (B) 99 J
- (C) 90 J
- (D) 80 J

Correct Answer: (C) 90 J

Solution:

Step 1: Understanding the Question:

A device that can operate as a Carnot engine is reversed to work as a refrigerator. We are given its efficiency as an engine and the work input as a refrigerator, and we need to find the heat extracted from the cold reservoir.

Step 2: Key Formula or Approach:

1. The efficiency of a Carnot heat engine is $\eta = 1 - \frac{T_C}{T_H}$, where T_C and T_H are the temperatures of the cold and hot reservoirs, respectively. 2. The Coefficient of Performance (COP) of a refrigerator is $\beta = \frac{Q_C}{W}$, where Q_C is the heat absorbed from the cold reservoir and W is the work done on the system. 3. For a Carnot cycle, the COP is also given by $\beta = \frac{T_C}{T_H - T_C}$. We can relate η and β .

Step 3: Detailed Explanation:

First, let's establish the relationship between η and β . From the efficiency formula, $\eta = \frac{T_H - T_C}{T_H}$.

From the COP formula, $\beta = \frac{T_C}{T_H - T_C}$.

Let's manipulate the efficiency formula: $\frac{1}{\eta} = \frac{T_H}{T_H - T_C}$.

$$\frac{1}{\eta} - 1 = \frac{T_H}{T_H - T_C} - \frac{T_H - T_C}{T_H - T_C} = \frac{T_H - (T_H - T_C)}{T_H - T_C} = \frac{T_C}{T_H - T_C} = \beta$$

So, the relationship is $\beta = \frac{1}{\eta} - 1 = \frac{1 - \eta}{\eta}$.

Now, use the given values: - Efficiency, $\eta = \frac{1}{10}$. - Work done, $W = 10 \text{ J}$.

Calculate the COP (β):

$$\beta = \frac{1 - 1/10}{1/10} = \frac{9/10}{1/10} = 9$$

Now, use the definition of COP to find the heat absorbed, Q_C :

$$\beta = \frac{Q_C}{W}$$

$$9 = \frac{Q_C}{10 \text{ J}}$$

$$Q_C = 9 \times 10 = 90 \text{ J}$$

Step 4: Final Answer:

The amount of energy absorbed from the lower temperature reservoir is 90 J. This corresponds to option (C).

💡 Quick Tip

For a Carnot cycle, the relationship between engine efficiency (η) and refrigerator COP (β) is $\beta = \frac{1 - \eta}{\eta}$. Memorizing this can be a useful shortcut.

74. Two photons of energy 2.5 eV and 3.5 eV fall on a metal surface of work function 1.5 eV. The ratio of the maximum velocities of the photoelectrons emitted from the metal surface is

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

Correct Answer: (D) 1 : $\sqrt{2}$

Solution:

Step 1: Understanding the Question:

We are given the energies of two incident photons and the work function of a metal. We need to find the ratio of the maximum velocities of the emitted photoelectrons.

Step 2: Key Formula or Approach:

We use Einstein's photoelectric equation:

$$KE_{max} = E_{photon} - \phi$$

where $KE_{max} = \frac{1}{2}mv_{max}^2$, E_{photon} is the energy of the incident photon, and ϕ is the work function of the metal.

Step 3: Detailed Explanation:

Let's calculate the maximum kinetic energy for each case.

For the first photon ($E_1 = 2.5 \text{ eV}$):

$$KE_1 = E_1 - \phi = 2.5 \text{ eV} - 1.5 \text{ eV} = 1.0 \text{ eV}$$

So, $\frac{1}{2}mv_1^2 = 1.0 \text{ eV}$.

For the second photon ($E_2 = 3.5 \text{ eV}$):

$$KE_2 = E_2 - \phi = 3.5 \text{ eV} - 1.5 \text{ eV} = 2.0 \text{ eV}$$

So, $\frac{1}{2}mv_2^2 = 2.0 \text{ eV}$.

Now, we find the ratio of the velocities. We can take the ratio of the two kinetic energy equations:

$$\frac{\frac{1}{2}mv_1^2}{\frac{1}{2}mv_2^2} = \frac{1.0 \text{ eV}}{2.0 \text{ eV}}$$

The $\frac{1}{2}m$ terms cancel out:

$$\frac{v_1^2}{v_2^2} = \frac{1}{2}$$

Take the square root of both sides to find the ratio of the velocities:

$$\frac{v_1}{v_2} = \sqrt{\frac{1}{2}} = \frac{1}{\sqrt{2}}$$

So, the ratio $v_1 : v_2$ is $1 : \sqrt{2}$.

Step 4: Final Answer:

The ratio of the maximum velocities is $1 : \sqrt{2}$. This corresponds to option (D).

 Quick Tip

When finding ratios, you often don't need to convert units (like eV to Joules) as long as they are consistent, because they will cancel out. The velocity is proportional to the square root of the maximum kinetic energy.

75. At critical angle, the angle of refraction is

- (A) 45°
- (B) 90°
- (C) 120°
- (D) 180°

Correct Answer: (B) 90°

Solution:

Step 1: Understanding the Question:

This is a definition-based question from optics concerning the phenomenon of total internal reflection.

Step 2: Detailed Explanation:

Total internal reflection can occur when light travels from a denser medium to a rarer (less dense) medium.

According to Snell's Law: $n_1 \sin \theta_1 = n_2 \sin \theta_2$, where $n_1 > n_2$.

Here, θ_1 is the angle of incidence in the denser medium, and θ_2 is the angle of refraction in the rarer medium.

The **critical angle** (θ_c) is defined as the specific angle of incidence θ_1 for which the angle of refraction θ_2 is exactly 90 degrees.

At this angle, the refracted ray travels exactly along the boundary surface between the two media.

If the angle of incidence is greater than the critical angle, the light does not refract into the rarer medium at all; instead, it is completely reflected back into the denser medium.

Therefore, by definition, when the angle of incidence is the critical angle, the angle of refraction is 90° .

Step 3: Final Answer:

At the critical angle, the angle of refraction is 90° . This corresponds to option (B).

 Quick Tip

Remember the condition for total internal reflection: light must travel from a denser to a rarer medium, and the angle of incidence must be greater than the critical angle. The critical angle itself is the boundary case where refraction happens at 90° .

Chemistry

76. The quantum number which describes the shape of an atomic orbital is indicated by the symbol

- (A) l
- (B) m_l
- (C) m_s
- (D) n

Correct Answer: (A) l

Solution:

Step 1: Understanding the Question:

The question asks to identify the quantum number that determines the fundamental shape of an electron's orbital in an atom.

Step 2: Detailed Explanation:

There are four primary quantum numbers that describe an electron in an atom:

- **n (Principal Quantum Number):** This number describes the main energy level or shell of the electron. It primarily determines the size and energy of the orbital. Its values are positive integers (1, 2, 3, ...).

- **l (Azimuthal or Angular Momentum Quantum Number):** This number describes the shape of the orbital and corresponds to the subshell. Its values range from 0 to $n-1$. $l = 0$ corresponds to an s-orbital, which is spherical. $l = 1$ corresponds to a p-orbital, which is dumbbell-shaped. $l = 2$ corresponds to a d-orbital, with more complex shapes (e.g., clover-leaf). $l = 3$ corresponds to an f-orbital, with even more complex shapes.

- **m_l (Magnetic Quantum Number):** This number describes the orientation of the orbital in three-dimensional space. Its values range from $-l$ to $+l$, including 0.

- **m_s (Spin Quantum Number):** This number describes the intrinsic angular momentum of the electron, which has two possible states: $+1/2$ or $-1/2$.

From these definitions, the quantum number that describes the shape of the orbital is the azimuthal quantum number, symbolized by 'l'.

Step 3: Final Answer:

The symbol for the quantum number that describes the shape of an atomic orbital is l. This corresponds to option (A).

💡 Quick Tip

A simple mnemonic to remember the roles: - **n** = Size/Energy (Shell) - **l** = Shape (Subshell) - **m_l** = Orientation (Orbital) - **m_s** = Spin

77. "No two electrons in an atom can have the same set of four quantum numbers". This is known as

- (A) Pauli's Principle
- (B) Hund's Rule
- (C) Aufbau Principle
- (D) Lewis Rule

Correct Answer: (A) Pauli's Principle

Solution:

Step 1: Understanding the Question:

The question provides a statement and asks for the scientific principle or rule it defines.

Step 2: Detailed Explanation:

Let's analyze the given principles:

- **Pauli's Principle (Pauli Exclusion Principle):** This fundamental principle of quantum mechanics states exactly what is quoted in the question: no two electrons in a single atom can have the same four quantum numbers (n, l, m_l, and m_s). An implication of this is that an atomic orbital can hold a maximum of two electrons, and these two electrons must have opposite spins.

- **Hund's Rule:** This rule deals with the filling of degenerate (equal-energy) orbitals within a subshell. It states that electrons will fill empty orbitals singly before pairing up, and that electrons in singly occupied orbitals will have the same spin.

- **Aufbau Principle:** This principle states that electrons fill atomic orbitals of the lowest available energy levels before occupying higher levels.

- **Lewis Rule (Octet Rule):** This is a concept in chemical bonding stating that atoms tend to bond in such a way that they each have eight electrons in their valence shell, giving them

the same electronic configuration as a noble gas.

The statement in the question is the precise definition of the Pauli Exclusion Principle.

Note on Answer Key: The provided image marks Hund's Rule as the correct answer. This is incorrect. The statement is the definition of the Pauli Exclusion Principle. For the purpose of this solution, we will proceed with the factually correct answer.

Step 3: Final Answer:

The statement is known as the Pauli Exclusion Principle. This corresponds to option (A).

 Quick Tip

Remember the key ideas: - **Aufbau:** Build up from lowest energy. - **Pauli Exclusion:** No two electrons are identical. - **Hund's Rule:** Maximize spin (spread out before pairing up).

78. In the elements with atomic number $Z=1$ to $Z=20$, how many of them have no unpaired electrons in their ground state?

- (A) 8
- (B) 4
- (C) 10
- (D) 6

Correct Answer: (D) 6

Solution:

Step 1: Understanding the Question:

We need to identify all elements from atomic number 1 to 20 whose ground-state electron configurations have all their electrons paired. This means every occupied orbital is completely filled.

Step 2: Detailed Explanation:

An element has no unpaired electrons if all of its occupied atomic subshells are completely filled with electrons. Let's list the electron configurations for $Z = 1$ to 20 and identify these elements.

- $Z = 1$ (H): $1s^1$ - 1 unpaired electron.
- **$Z = 2$ (He): $1s^2$** - No unpaired electrons (1s subshell is full).
- $Z = 3$ (Li): $[\text{He}] 2s^1$ - 1 unpaired electron.
- **$Z = 4$ (Be): $[\text{He}] 2s^2$** - No unpaired electrons (2s subshell is full).
- $Z = 5$ (B): $[\text{He}] 2s^2 2p^1$ - 1 unpaired electron.
- $Z = 6$ (C): $[\text{He}] 2s^2 2p^2$ - 2 unpaired electrons (Hund's rule).

- $Z = 7$ (N): [He] $2s^2 2p^3$ - 3 unpaired electrons.
- $Z = 8$ (O): [He] $2s^2 2p^4$ - 2 unpaired electrons.
- $Z = 9$ (F): [He] $2s^2 2p^5$ - 1 unpaired electron.
- **$Z = 10$ (Ne): [He] $2s^2 2p^6$** - No unpaired electrons (2p subshell is full).
- $Z = 11$ (Na): [Ne] $3s^1$ - 1 unpaired electron.
- **$Z = 12$ (Mg): [Ne] $3s^2$** - No unpaired electrons (3s subshell is full).
- $Z = 13$ (Al): [Ne] $3s^2 3p^1$ - 1 unpaired electron.
- $Z = 14$ (Si): [Ne] $3s^2 3p^2$ - 2 unpaired electrons.
- $Z = 15$ (P): [Ne] $3s^2 3p^3$ - 3 unpaired electrons.
- $Z = 16$ (S): [Ne] $3s^2 3p^4$ - 2 unpaired electrons.
- $Z = 17$ (Cl): [Ne] $3s^2 3p^5$ - 1 unpaired electron.
- **$Z = 18$ (Ar): [Ne] $3s^2 3p^6$** - No unpaired electrons (3p subshell is full).
- $Z = 19$ (K): [Ar] $4s^1$ - 1 unpaired electron.
- **$Z = 20$ (Ca): [Ar] $4s^2$** - No unpaired electrons (4s subshell is full).

The elements with no unpaired electrons are Helium (He), Beryllium (Be), Neon (Ne), Magnesium (Mg), Argon (Ar), and Calcium (Ca).

Step 3: Final Answer:

Counting these elements, we find there are 6 of them. This corresponds to option (D).

💡 Quick Tip

Elements with no unpaired electrons are those with completely filled subshells. These are the noble gases (He, Ne, Ar) and the alkaline earth metals (Be, Mg, Ca) within the first 20 elements.

79. Which of the following is not a property of covalent compounds?

- (A) They are generally insoluble in water
- (B) They consist of molecules
- (C) They exist as solids, liquids or gases
- (D) The reactions between them are fast

Correct Answer: (D) The reactions between them are fast

Solution:

Step 1: Understanding the Question:

We need to identify the statement that does not accurately describe the general properties of covalent compounds.

Step 2: Detailed Explanation:

Let's analyze the properties of covalent compounds based on their bonding and structure. Co-

valent compounds consist of discrete molecules held together by relatively weak intermolecular forces (like van der Waals forces or hydrogen bonds).

- **(A) They are generally insoluble in water:** This is a typical property. Water is a polar solvent. Nonpolar covalent compounds (like oil, CCl_4) do not dissolve in water. While some polar covalent compounds (like sugar, ethanol) do dissolve, the general rule "like dissolves like" makes this statement largely true for the broad class of covalent compounds.

- **(B) They consist of molecules:** This is the defining characteristic of covalent compounds (excluding network covalent solids). The atoms are bonded into discrete, neutral units called molecules.

- **(C) They exist as solids, liquids or gases:** Due to the wide range of strengths of intermolecular forces, covalent compounds can exist in all three states at room temperature. For example, methane (gas), water (liquid), and iodine (solid). This is a valid property.

- **(D) The reactions between them are fast:** This is generally **not** a property of covalent compounds. Reactions involving covalent compounds require the breaking of strong covalent bonds and the formation of new ones. This process often requires a significant amount of activation energy, making the reactions relatively slow. In contrast, reactions between ionic compounds in solution are often instantaneous because they involve the interaction of free-moving ions.

Step 3: Final Answer:

The statement that is not a property of covalent compounds is that their reactions are fast. This corresponds to option (D).

💡 Quick Tip

Contrast covalent properties with ionic properties. Ionic compounds are crystalline solids, soluble in water, conduct electricity when molten or dissolved, and their reactions in solution are very fast. Covalent compounds are often the opposite on these points.

80. The sum of covalent bonds in H_2 , N_2 and HCl is

- (A) 4
- (B) 5
- (C) 6
- (D) 3

Correct Answer: (B) 5

Solution:

Step 1: Understanding the Question:

The question requires us to determine the number of covalent bonds in each of the three given molecules (H_2 , N_2 , HCl) and then calculate their sum.

Step 2: Detailed Explanation:

We will find the number of covalent bonds in each molecule by considering the valence electrons and the need to achieve a stable electron configuration (a duet for hydrogen, an octet for nitrogen and chlorine).

1. Hydrogen molecule (H_2):

A hydrogen atom (H) has 1 valence electron. To achieve a stable configuration like Helium, it needs one more electron. Two hydrogen atoms share their single electrons to form one pair, resulting in a single covalent bond (H-H).

Number of bonds in $\text{H}_2 = 1$.

2. Nitrogen molecule (N_2):

A nitrogen atom (N) has 5 valence electrons. To achieve a stable octet, it needs three more electrons. Two nitrogen atoms share three pairs of electrons, resulting in a triple covalent bond ($\text{N}\equiv\text{N}$).

Number of bonds in $\text{N}_2 = 3$.

3. Hydrogen Chloride molecule (HCl):

A hydrogen atom (H) has 1 valence electron, and a chlorine atom (Cl) has 7 valence electrons. Hydrogen needs one more electron, and chlorine needs one more electron to complete their respective stable configurations. They share one pair of electrons, resulting in a single covalent bond (H-Cl).

Number of bonds in $\text{HCl} = 1$.

4. Sum of the bonds:

Total number of covalent bonds = (Bonds in H_2) + (Bonds in N_2) + (Bonds in HCl)

Total = $1 + 3 + 1 = 5$.

Step 3: Final Answer:

The sum of covalent bonds in the three molecules is 5. This corresponds to option (B).

Quick Tip

To quickly determine the number of bonds between atoms in simple diatomic molecules, consider how many electrons each atom needs to gain to achieve a noble gas configuration. This number is often the number of covalent bonds it will form.

81. How many grams of NaOH is required to prepare 5.0 litre of 0.1 N solution?

(Given: At. wt: H=1, O=16, Na=23)

- (A) 20
- (B) 30
- (C) 10
- (D) 50

Correct Answer: (A) 20

Solution:

Step 1: Understanding the Question:

The question asks for the mass of Sodium Hydroxide (NaOH) needed to prepare a specific volume of a solution with a given normality.

Step 2: Key Formula or Approach:

1. Calculate the Equivalent Weight (EW) of NaOH.

$$\text{EW} = \text{Molecular Weight (MW)} / \text{n-factor.}$$

2. Calculate the number of gram equivalents required.

$$\text{Gram Equivalents} = \text{Normality (N)} \times \text{Volume (V in litres)}.$$

3. Calculate the required mass.

$$\text{Mass} = \text{Gram Equivalents} \times \text{Equivalent Weight}.$$

Step 3: Detailed Explanation:

1. Calculate Molecular and Equivalent Weight of NaOH:

$$\text{MW of NaOH} = 23(\text{Na}) + 16(\text{O}) + 1(\text{H}) = 40 \text{ g/mol}$$

Since the n-factor for NaOH is 1:

$$\text{EW of NaOH} = \frac{40}{1} = 40 \text{ g/equivalent}$$

2. Calculate Gram Equivalents:

$$\text{Normality (N)} = 0.1 \text{ N} = 0.1 \text{ equivalents/litre}$$

$$\text{Volume (V)} = 5.0 \text{ litres}$$

$$\text{Gram Equivalents} = 0.1 \frac{\text{equivalents}}{\text{litre}} \times 5.0 \text{ litres} = 0.5 \text{ equivalents}$$

3. Calculate Required Mass:

$$\text{Mass of NaOH} = \text{Gram Equivalents} \times \text{EW}$$

$$\text{Mass of NaOH} = 0.5 \text{ equivalents} \times 40 \frac{\text{g}}{\text{equivalent}} = 20 \text{ g}$$

Step 4: Final Answer:

20 grams of NaOH are required.

This corresponds to option (A).

💡 Quick Tip

For monobasic bases (like NaOH, KOH) and monoacidic acids (like HCl, HNO₃), Normality is equal to Molarity.

You could calculate moles (Molarity $\times$ Volume) and then multiply by Molar Mass to get the same result.

82. A gaseous mixture contains 8g of oxygen, 14 g of nitrogen and 8 g of hydrogen. Total number of molecules present in the gaseous mixture is (Given: At. wt: H=1, N=14, O=16, $N_A = 6 \times 10^{23} \text{ mol}^{-1}$)

- (A) 1.43×10^{23}
- (B) 2.85×10^{23}
- (C) 2.85×10^{24}
- (D) 1.85×10^{24}

Correct Answer: (C) 2.85×10^{24}

Solution:

Step 1: Understanding the Question:

We are given the masses of three different gases in a mixture and need to calculate the total number of molecules.

Step 2: Key Formula or Approach:

- Calculate the number of moles for each gas using the formula: $\text{Moles} = \frac{\text{Given Mass}}{\text{Molar Mass}}$.

Remember that oxygen, nitrogen, and hydrogen are diatomic (O₂, N₂, H₂).

- Find the total number of moles by summing the moles of each gas.
- Calculate the total number of molecules using Avogadro's law: $\text{Total Molecules} = \text{Total Moles} \times \text{Avogadro's Number } (N_A)$.

Step 3: Detailed Explanation:

1. Molar Masses:

- Oxygen (O_2): $2 \times 16 = 32 \text{ g/mol}$
- Nitrogen (N_2): $2 \times 14 = 28 \text{ g/mol}$
- Hydrogen (H_2): $2 \times 1 = 2 \text{ g/mol}$

2. Number of Moles:

- Moles of O_2 : $\frac{8 \text{ g}}{32 \text{ g/mol}} = 0.25 \text{ mol}$
- Moles of N_2 : $\frac{14 \text{ g}}{28 \text{ g/mol}} = 0.5 \text{ mol}$
- Moles of H_2 : $\frac{8 \text{ g}}{2 \text{ g/mol}} = 4.0 \text{ mol}$

3. Total Moles:

- Total Moles = $0.25 + 0.5 + 4.0 = 4.75 \text{ mol}$

4. Total Molecules:

- Total Molecules = Total Moles $\times N_A$
- Total Molecules = $4.75 \times (6 \times 10^{23})$
- Total Molecules = $28.5 \times 10^{23} = 2.85 \times 10^{24}$

Step 4: Final Answer:

The total number of molecules in the mixture is 2.85×10^{24} .

This corresponds to option (C).

Quick Tip

A common mistake is forgetting that common gaseous elements like oxygen, nitrogen, and hydrogen are diatomic.

Always use the molar mass of the molecule (e.g., O_2), not the atom (O), when calculating moles from mass.

83. The equivalent weight of which of the following is the highest?

- (A) Na_2CO_3 (molecular weight = 106)
- (B) H_3PO_4 (molecular weight = 98)
- (C) $\text{H}_2\text{C}_2\text{O}_4 \cdot 2\text{H}_2\text{O}$ (molecular weight = 126)
- (D) AlCl_3 (molecular weight = 133.5)

Correct Answer: (C) $\text{H}_2\text{C}_2\text{O}_4 \cdot 2\text{H}_2\text{O}$ (molecular weight = 126)

Solution:**Step 1: Understanding the Question:**

We need to calculate the equivalent weight for each of the given compounds and identify which one has the largest value.

Step 2: Key Formula or Approach:

Equivalent Weight (EW) is defined as:

$$EW = \frac{\text{Molecular Weight (MW)}}{\text{n-factor}}$$

The n-factor depends on the type of substance and the reaction context.

Assuming standard acid-base or salt reactions:

- For an acid, n-factor is its basicity (number of replaceable H^+ ions).
- For a salt, n-factor is the total magnitude of positive or negative charge on the ions.

Step 3: Detailed Explanation:**(A) Na_2CO_3 :**

This is a salt.

It dissociates into 2 Na^+ ions and one CO_3^{2-} ion.

The total positive charge is 2. So, n-factor = 2.

$$EW = \frac{106}{2} = 53$$

(B) H_3PO_4 :

This is phosphoric acid, a tribasic acid.

It can donate up to 3 protons. So, n-factor = 3.

$$EW = \frac{98}{3} \approx 32.67$$

(C) $\text{H}_2\text{C}_2\text{O}_4 \cdot 2\text{H}_2\text{O}$:

This is oxalic acid dihydrate, a dibasic acid.

It has two replaceable protons. So, n-factor = 2.

$$EW = \frac{126}{2} = 63$$

(D) AlCl_3 :

This is a salt.

It dissociates into one Al^{3+} ion and three Cl^- ions.

The total positive charge is 3. So, n-factor = 3.

$$EW = \frac{133.5}{3} = 44.5$$

Comparison:

Comparing the calculated equivalent weights: 53, 32.67, 63, 44.5.

The highest value is 63, which corresponds to oxalic acid dihydrate.

Step 4: Final Answer:

The equivalent weight is highest for $\text{H}_2\text{C}_2\text{O}_4 \cdot 2\text{H}_2\text{O}$.

This corresponds to option (C).

💡 Quick Tip

Remember that the n-factor can change depending on the reaction.

For example, in redox reactions, the n-factor for oxalic acid is also 2 (as C goes from +3 to +4), but for H_3PO_4 , it can react to form salts like NaH_2PO_4 (n=1) or Na_2HPO_4 (n=2).

However, unless specified, assume complete reaction.

84. At 25°C , ionic product (K_w) of 0.01M HCl solution is

- (A) $1.0 \times 10^{-13} \text{ mol}^2/\text{L}^2$
- (B) $1.0 \times 10^{-12} \text{ mol}^2/\text{L}^2$
- (C) $1.0 \times 10^{-14} \text{ mol}^2/\text{L}^2$
- (D) $1.0 \times 10^{-15} \text{ mol}^2/\text{L}^2$

Correct Answer: (C) $1.0 \times 10^{-14} \text{ mol}^2/\text{L}^2$

Solution:

Step 1: Understanding the Question:

The question asks for the value of the ionic product of water (K_w) in a specific acidic solution at a standard temperature.

Step 2: Key Formula or Approach:

The ionic product of water, K_w , is the equilibrium constant for the autoionization of water ($\text{H}_2\text{O} \rightleftharpoons \text{H}^+ + \text{OH}^-$).

It is defined as:

$$K_w = [\text{H}^+][\text{OH}^-]$$

A key property of K_w is that its value depends **only on temperature**.

For any dilute aqueous solution at a given temperature, the value of K_w is constant.

Step 3: Detailed Explanation:

The temperature is given as 25°C .

At 25°C , the standard value for the ionic product of water is a well-known constant:

$$K_w = 1.0 \times 10^{-14} \text{ mol}^2/\text{L}^2$$

The presence of 0.01M HCl in the solution does not change the value of K_w .

The HCl will increase the $[\text{H}^+]$ concentration to 0.01 M, and according to Le Chatelier's principle, the $[\text{OH}^-]$ concentration will decrease to maintain the constant product K_w .

$$[\text{OH}^-] = K_w/[\text{H}^+] = (1.0 \times 10^{-14})/(0.01) = 1.0 \times 10^{-12} \text{ M}.$$

However, the value of K_w itself remains 1.0×10^{-14} .

Step 4: Final Answer:

The ionic product (K_w) of the solution at 25°C is $1.0 \times 10^{-14} \text{ mol}^2/\text{L}^2$.

This corresponds to option (C).

💡 Quick Tip

This is a conceptual question.

Remember that equilibrium constants (like K_w , K_a , K_b) are only dependent on temperature.

The addition of acids or bases changes the concentrations of the species involved but not the value of the equilibrium constant itself.

85. Which of the following combinations give a buffer solution?

- (A) $\text{HCl} + \text{NaCl}$
- (B) $\text{CH}_3\text{COOH} + \text{CH}_3\text{COONa}$
- (C) $\text{CH}_3\text{COOH} + \text{NaCl}$
- (D) $\text{NH}_4\text{OH} + \text{NaOH}$

Correct Answer: (B) $\text{CH}_3\text{COOH} + \text{CH}_3\text{COONa}$

Solution:

Step 1: Understanding the Question:

We need to identify which pair of chemicals, when mixed, will form a buffer solution.

Step 2: Key Formula or Approach:

A buffer solution is a solution that resists changes in pH upon the addition of small amounts of an acid or a base.

There are two main types:

1. **Acidic Buffer:** A mixture of a weak acid and its salt with a strong base (which provides the conjugate base).

Example: Acetic acid (CH_3COOH) and sodium acetate (CH_3COONa).

2. **Basic Buffer:** A mixture of a weak base and its salt with a strong acid (which provides the conjugate acid).

Example: Ammonium hydroxide (NH_4OH) and ammonium chloride (NH_4Cl).

Step 3: Detailed Explanation:

Let's analyze each option:

(A) $\text{HCl} + \text{NaCl}$:

HCl is a strong acid.

NaCl is the salt of a strong acid (HCl) and a strong base (NaOH).

A mixture of a strong acid and its salt is not a buffer.

(B) $\text{CH}_3\text{COOH} + \text{CH}_3\text{COONa}$:

CH_3COOH (acetic acid) is a weak acid.

CH_3COONa (sodium acetate) is its salt with a strong base (NaOH).

This salt provides the conjugate base, CH_3COO^- .

This combination of a weak acid and its conjugate base forms an acidic buffer solution.

(C) $\text{CH}_3\text{COOH} + \text{NaCl}$:

CH_3COOH is a weak acid.

NaCl is a salt that is not related to acetic acid (it doesn't provide the conjugate base).

This mixture is not a buffer.

(D) $\text{NH}_4\text{OH} + \text{NaOH}$:

NH_4OH (ammonium hydroxide) is a weak base.

NaOH is a strong base.

A mixture of a weak base and a strong base is not a buffer.

The strong base would suppress the dissociation of the weak base.

Step 4: Final Answer:

The combination that forms a buffer solution is $\text{CH}_3\text{COOH} + \text{CH}_3\text{COONa}$.

This corresponds to option (B).

💡 Quick Tip

The key to identifying a buffer is to look for a "weak conjugate pair".

This means a weak acid and its conjugate base (e.g., HF and F^-), or a weak base and its conjugate acid (e.g., NH_3 and NH_4^+).

The conjugate partner is usually supplied by a salt.

86. A current of 0.5 amp is passed through molten AlCl_3 for 96.5 seconds. The volume of Cl_2 gas liberated at STP at anode (in ml) is ($\text{Cl} = 35.5 \text{ u}$) ($1\text{F} = 96500 \text{ C mol}^{-1}$)

- (A) 11.2
- (B) 22.4
- (C) 5.6
- (D) 33.6

Correct Answer: (C) 5.6

Solution:

Step 1: Understanding the Question:

We are performing electrolysis of molten AlCl_3 and need to calculate the volume of chlorine

gas produced at the anode under standard conditions (STP).

Step 2: Key Formula or Approach:

1. Calculate the total electric charge (Q) passed through the electrolyte: $Q = I \times t$.
2. Determine the half-reaction at the anode (oxidation).
3. Use Faraday's laws of electrolysis to find the moles of product.

The number of moles of electrons is $\frac{Q}{F}$, where F is the Faraday constant (96500 C/mol).

4. Use the stoichiometry of the half-reaction to relate moles of electrons to moles of Cl_2 .
5. Calculate the volume of gas at STP using the molar volume: 1 mole of any ideal gas at STP occupies 22.4 L or 22400 mL.

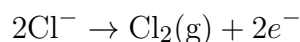
Step 3: Detailed Explanation:

1. Calculate Charge (Q):

$$Q = I \times t = 0.5 \text{ A} \times 96.5 \text{ s} = 48.25 \text{ C}$$

2. Anode Half-Reaction:

At the anode, chloride ions are oxidized to chlorine gas:



3. Calculate Moles of Electrons:

$$\text{Moles of } e^- = \frac{Q}{F} = \frac{48.25 \text{ C}}{96500 \text{ C/mol}} = 0.0005 \text{ mol}$$

4. Calculate Moles of Cl_2 :

From the half-reaction, 2 moles of electrons produce 1 mole of Cl_2 gas.

Therefore:

$$\text{Moles of } \text{Cl}_2 = \frac{\text{Moles of } e^-}{2} = \frac{0.0005}{2} = 0.00025 \text{ mol}$$

5. Calculate Volume at STP:

$$\text{Volume of } \text{Cl}_2 = \text{Moles of } \text{Cl}_2 \times \text{Molar Volume at STP}$$

$$\text{Volume of } \text{Cl}_2 = 0.00025 \text{ mol} \times 22400 \text{ mL/mol}$$

$$\text{Volume of } \text{Cl}_2 = 5.6 \text{ mL}$$

Step 4: Final Answer:

The volume of Cl_2 gas liberated is 5.6 mL.

This corresponds to option (C).

 Quick Tip

Notice the numbers given: 96.5 seconds.

This is a strong hint to use the Faraday constant, as 96.5 is $96500 / 1000$.

This often simplifies calculations.

$$Q = 0.5 \times 96.5 = 48.25 = 96500/2000.$$

$$\text{So, moles of electrons} = 1/2000 = 0.0005.$$

87. The amount of substance deposited due to passage of 1F of electricity is called

- (A) Atomic weight
- (B) Equivalent weight
- (C) Electrochemical equivalent
- (D) Molecular weight

Correct Answer: (B) Equivalent weight

Solution:

Step 1: Understanding the Question:

This is a definition-based question from the topic of electrolysis.

We need to identify the term for the mass of a substance deposited by one Faraday of charge.

Step 2: Detailed Explanation:

Let's define the terms:

- **Faraday (F):** One Faraday is the magnitude of the electric charge per mole of electrons.

$$1F \approx 96500 \text{ C/mol.}$$

- **Electrochemical Equivalent (ECE):** This is the mass of a substance deposited or liberated by the passage of 1 Coulomb of charge.

- **Equivalent Weight (EW):** This is the mass of a substance that combines with or displaces 1 mole of hydrogen atoms, or 8 grams of oxygen, or 1 mole of electrons.

In electrochemistry, it's the molar mass divided by the number of electrons transferred per formula unit in the redox reaction (n-factor).

According to Faraday's laws, the passage of 1 Faraday (1 mole of electrons) will deposit $\frac{1}{n}$ moles of a substance, where n is the charge of the ion.

The mass deposited is $\frac{1}{n} \times \text{Molar Mass}$.

This quantity, $\frac{\text{Molar Mass}}{n}$, is precisely the definition of the **Equivalent Weight**.

Step 3: Final Answer:

The amount of substance deposited by the passage of 1F of electricity is one gram equivalent weight.

This corresponds to option (B).

💡 Quick Tip

Remember the key distinction:

- 1 Coulomb deposits the Electrochemical Equivalent (ECE).
- 1 Faraday (96500 C) deposits the Equivalent Weight (EW).

The relationship is: $\text{EW} = \text{ECE} \times 96500$.

88. What is the emf of the cell? $\text{Sn}|\text{Sn}^{2+} (1\text{M})||\text{Ag}^+ (1\text{M})|\text{Ag}$ [Given $E^\circ_{\text{Sn}^{2+}|\text{Sn}} = -0.14\text{V}$ and $E^\circ_{\text{Ag}^+|\text{Ag}} = +0.80\text{V}$]

- (A) 0.66 V
- (B) 0.80 V
- (C) 1.08 V
- (D) 0.94 V

Correct Answer: (D) 0.94 V

Solution:

Step 1: Understanding the Question:

We are asked to calculate the electromotive force (emf) of a galvanic cell given in standard cell notation and with standard reduction potentials for the half-cells.

Step 2: Key Formula or Approach:

1. Identify the anode and cathode from the cell notation.

The convention is Anode || Cathode.

2. The anode is where oxidation occurs, and the cathode is where reduction occurs.

3. The standard emf of the cell (E_{cell}°) is calculated using the standard reduction potentials of the cathode and anode:

$$E_{cell}^{\circ} = E_{cathode}^{\circ} - E_{anode}^{\circ}$$

4. Since the concentrations of the ions are 1M, the conditions are standard, and the cell emf is equal to the standard cell emf ($E_{cell} = E_{cell}^{\circ}$).

Step 3: Detailed Explanation:

From the cell notation $\text{Sn}|\text{Sn}^{2+} (1\text{M})||\text{Ag}^{+} (1\text{M})|\text{Ag}$:

- **Anode (Oxidation):** $\text{Sn} \rightarrow \text{Sn}^{2+} + 2\text{e}^{-}$. The Sn electrode is the anode.

- **Cathode (Reduction):** $\text{Ag}^{+} + \text{e}^{-} \rightarrow \text{Ag}$. The Ag electrode is the cathode.

The given standard reduction potentials are:

$$E_{\text{Sn}^{2+}|\text{Sn}}^{\circ} = E_{anode}^{\circ} = -0.14 \text{ V}$$

$$E_{\text{Ag}^{+}|\text{Ag}}^{\circ} = E_{cathode}^{\circ} = +0.80 \text{ V}$$

Now, we apply the formula for the standard cell emf:

$$E_{cell}^{\circ} = E_{cathode}^{\circ} - E_{anode}^{\circ}$$

$$E_{cell}^{\circ} = (+0.80 \text{ V}) - (-0.14 \text{ V})$$

$$E_{cell}^{\circ} = 0.80 \text{ V} + 0.14 \text{ V} = 0.94 \text{ V}$$

Step 4: Final Answer:

The emf of the cell is 0.94 V.

This corresponds to option (D).

💡 Quick Tip

A spontaneous galvanic cell must have a positive E_{cell} .

The formula $E_{cell} = E_{cathode}^{\circ} - E_{anode}^{\circ}$ ensures this.

An alternative is $E_{cell} = E_{reduction}^{\circ} + E_{oxidation}^{\circ}$, but you must remember to reverse the sign of the potential for the oxidation half-reaction.

The first method is generally safer and less prone to sign errors.

89. If the standard reduction potentials of A,B,C are respectively 0.68V, -2.54V and -0.50 V, then the order of their reducing power is

- (A) $A > B > C$
- (B) $A > C > B$
- (C) $C > B > A$
- (D) $B > C > A$

Correct Answer: (D) $B > C > A$

Solution:

Step 1: Understanding the Question:

We are given the standard reduction potentials for three species and asked to rank them in order of their strength as reducing agents (reducing power).

Step 2: Key Formula or Approach:

- **Standard Reduction Potential (E°):** A measure of the tendency of a chemical species to be reduced (gain electrons).

A higher (more positive) E° means a greater tendency to be reduced.

- **Reducing Agent:** A substance that causes another substance to be reduced, and in the process, it gets oxidized (loses electrons).

- **Relationship:** A strong reducing agent is a substance that is easily oxidized.

A substance is easily oxidized if its tendency to be reduced is very low.

Therefore, a stronger reducing agent has a **lower (more negative)** standard reduction potential.

Step 3: Detailed Explanation:

The given standard reduction potentials (E°) are:

- $E^\circ(A) = +0.68 \text{ V}$

- $E^\circ(B) = -2.54 \text{ V}$

- $E^\circ(C) = -0.50 \text{ V}$

To find the order of reducing power, we need to arrange these potentials in increasing order (from most negative to most positive).

The order of the E° values is:

$$E^\circ(B) < E^\circ(C) < E^\circ(A) \\ -2.54 \text{ V} < -0.50 \text{ V} < +0.68 \text{ V}$$

Since a lower reduction potential implies a stronger reducing agent, the order of reducing power is:

Reducing Power: $B > C > A$

Step 4: Final Answer:

The correct order of reducing power is $B > C > A$.

This corresponds to option (D).

💡 Quick Tip

Remember this simple rule:

Lower Reduction Potential = Stronger Reducing Power.

Conversely, a higher reduction potential means a stronger oxidizing power.

90. With which of the following anions, Mg^{2+} and Ca^{2+} ions form salts responsible for permanent hardness of water?

- (A) Cl^- , SO_4^{2-}
- (B) Cl^- , NO_3^-
- (C) HCO_3^- , Cl^-
- (D) CO_3^{2-} , HCO_3^-

Correct Answer: (A) Cl^- , SO_4^{2-}

Solution:

Step 1: Understanding the Question:

The question asks to identify the anions that cause permanent hardness in water when they form salts with magnesium (Mg^{2+}) and calcium (Ca^{2+}) ions.

Step 2: Key Formula or Approach:

Water hardness is primarily caused by dissolved divalent cations, mainly Ca^{2+} and Mg^{2+} .

The type of hardness depends on the accompanying anion.

- **Temporary Hardness (Carbonate Hardness):** Caused by the presence of bicarbonate (hydrogen carbonate, HCO_3^-) salts of calcium and magnesium.

It is called "temporary" because it can be removed by boiling.

- **Permanent Hardness (Non-carbonate Hardness):** Caused by the presence of chloride (Cl^-) and sulfate (SO_4^{2-}) salts of calcium and magnesium.

These salts are soluble and are not removed by boiling.

Step 3: Detailed Explanation:

Based on the definitions above, permanent hardness is due to salts like CaCl_2 , MgCl_2 , CaSO_4 , and MgSO_4 .

The anions responsible are therefore chloride (Cl^-) and sulfate (SO_4^{2-}).

Let's analyze the options:

- (A) Cl^- , SO_4^{2-} : These are the anions that cause permanent hardness.

- (B) Cl^- , NO_3^- : Nitrates can contribute to hardness but chlorides and sulfates are the primary causes.

- (C) HCO_3^- , Cl^- : Bicarbonate (HCO_3^-) causes temporary hardness.

- (D) CO_3^{2-} , HCO_3^- : Carbonates are largely insoluble and bicarbonate causes temporary hardness.

Step 4: Final Answer:

The anions responsible for permanent hardness are Cl^- and SO_4^{2-} .

This corresponds to option (A).

💡 Quick Tip

A simple way to remember:

Temporary hardness is caused by bicarbonates and can be removed by boiling.

Permanent hardness is caused by "everything else", primarily chlorides and sulfates, and cannot be removed by boiling.

91. Exhausted permutit is regenerated by washing with

- (A) Dilute NaOH solution
- (B) Dilute NaCl solution
- (C) Dilute HCl solution
- (D) Dilute AlCl_3 solution

Correct Answer: (B) Dilute NaCl solution

Solution:

Step 1: Understanding the Question:

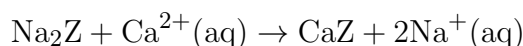
The question asks about the chemical used to regenerate permutit after it has been used for water softening.

Step 2: Key Formula or Approach:

Permutit is a type of zeolite (sodium aluminium silicate, represented as Na_2Z) used in ion-exchange water softeners.

- **Softening Process:** Hard water containing Ca^{2+} and Mg^{2+} ions is passed through the permutit.

The Ca^{2+} and Mg^{2+} ions displace the Na^+ ions from the permutit.



- **Exhausted Permutit:** After some time, the permutit becomes "exhausted" (e.g., CaZ or MgZ).

- **Regeneration:** To make the softener usable again, the process must be reversed by passing a concentrated solution of a sodium salt through it.

Step 3: Detailed Explanation:

The regeneration reaction is:



A concentrated solution of sodium chloride (NaCl), commonly known as brine, is used for this purpose because it is inexpensive.

Step 4: Final Answer:

Exhausted permutit is regenerated by washing with a solution of NaCl .

This corresponds to option (B).

💡 Quick Tip

The permutit process is an ion-exchange process.

To regenerate the original material (sodium permutit), you need to flood it with a high concentration of the ion that was originally displaced (sodium ions).

Common salt (NaCl) is the cheapest source of sodium ions.

92. 27.2 mg of CaSO_4 and 2.4 mg of MgSO_4 are present in a 2 kg water sample. What is the total hardness of water (in ppm) in terms of equivalents of CaCO_3 ? (molecular weight of $\text{CaSO}_4 = 136$ & molecular weight of $\text{MgSO}_4 = 120$)

- (A) 11
- (B) 10
- (C) 20
- (D) 22

Correct Answer: (A) 11

Solution:

Step 1: Understanding the Question:

We are given the masses of two hardness-causing salts in a specific mass of water.

We need to calculate the total hardness in parts per million (ppm) expressed as calcium carbonate equivalents.

Step 2: Key Formula or Approach:

1. Calculate the CaCO_3 equivalent mass for each salt.

$$\text{CaCO}_3 \text{ equivalent} = \text{Mass of salt} \times \frac{\text{MW of CaCO}_3}{\text{MW of salt}}$$

2. Sum the CaCO_3 equivalents to get the total equivalent mass.
3. Calculate the hardness in ppm.

$$\text{Hardness (ppm)} = \frac{\text{Total mass of CaCO}_3 \text{ equivalent (in mg)}}{\text{Mass of water (in kg)}}$$

Step 3: Detailed Explanation:

Given:

- Mass of $\text{CaSO}_4 = 27.2 \text{ mg}$
- Mass of $\text{MgSO}_4 = 2.4 \text{ mg}$
- Mass of water = 2 kg
- MW of $\text{CaSO}_4 = 136$
- MW of $\text{MgSO}_4 = 120$
- MW of $\text{CaCO}_3 = 100$

1. CaCO_3 equivalent of CaSO_4 :

$$= 27.2 \text{ mg} \times \frac{100}{136} = 20 \text{ mg of } \text{CaCO}_3 \text{ equivalent}$$

2. CaCO_3 equivalent of MgSO_4 :

$$= 2.4 \text{ mg} \times \frac{100}{120} = 2.0 \text{ mg of } \text{CaCO}_3 \text{ equivalent}$$

3. Total CaCO_3 equivalent mass:

$$\text{Total equivalent mass} = 20 \text{ mg} + 2 \text{ mg} = 22 \text{ mg}$$

4. Calculate Hardness in ppm:

The total equivalent mass (22 mg) is present in 2 kg of water.

Hardness in ppm is the equivalent mass per 1 kg of water.

$$\text{Hardness (ppm)} = \frac{22 \text{ mg}}{2 \text{ kg}} = 11 \text{ mg/kg} = 11 \text{ ppm}$$

Step 4: Final Answer:

The total hardness of the water is 11 ppm.

This corresponds to option (A).

💡 Quick Tip

The molecular weight of CaCO_3 (100) is a convenient standard.

The key is to convert every hardness-causing salt into its CaCO_3 equivalent before summing them up.

Finally, remember that ppm for water hardness is mg of CaCO_3 equivalent per litre (or kg) of water.

93. Statement I: The lower the pH greater is the corrosion.

Statement II: Electrochemical Corrosion always occurs at the anodic area. The correct answer is

- (A) Both statement - I and Statement - II are correct
- (B) Both statement - I and Statement - II are not correct
- (C) Statement - I is correct but statement - II is not correct
- (D) Statement - I is not correct but statement - II is correct

Correct Answer: (A) Both statement - I and Statement - II are correct

Solution:

Step 1: Understanding the Question:

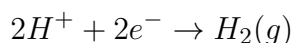
We need to evaluate the correctness of two separate statements related to corrosion.

Step 2: Detailed Explanation:

Statement I: The lower the pH greater is the corrosion.

A lower pH indicates a higher concentration of H^+ ions (higher acidity).

In acidic solutions, the reduction of H^+ ions is a common cathodic reaction that drives corrosion:



A higher concentration of H^+ ions facilitates this reaction, which accelerates the rate of oxidation (corrosion) at the anode.

Therefore, this statement is correct.

Statement II: Electrochemical Corrosion always occurs at the anodic area.

Electrochemical corrosion is a process involving an anode and a cathode.

- At the **anode**, oxidation occurs, which for a metal is the process of corrosion.

Example: $Fe \rightarrow Fe^{2+} + 2e^-$.

- At the **cathode**, reduction occurs, consuming the electrons produced at the anode.

Corrosion is the destructive oxidation of the metal, which happens at the anode by definition.

Therefore, this statement is correct.

Step 3: Final Answer:

Both Statement I and Statement II are correct.

This corresponds to option (A).

💡 Quick Tip

Remember the mnemonic "An Ox, Red Cat" for electrochemistry:

Anode is for **O**xidation, **R**eduction is at the **C**athode.

Corrosion is an oxidation process, so it must happen at the anode.

94. Rust is chemically

- (A) Hydrated Ferric Oxide
- (B) Hydrated Copper (II) Chloride
- (C) Hydrated Ferrous Sulphate
- (D) Hydrated Ferric Sulphate

Correct Answer: (A) Hydrated Ferric Oxide

Solution:

Step 1: Understanding the Question:

The question asks for the chemical name of rust.

Step 2: Detailed Explanation:

Rust is the common term for the reddish-brown substance that forms on the surface of iron or its alloys when exposed to oxygen and moisture.

The chemical process involves the oxidation of iron (Fe) to ferric ions (Fe^{3+}).

These ferric ions combine with oxygen and water to form a complex, hydrated compound.

The chemical formula for rust is generally written as $\text{Fe}_2\text{O}_3 \cdot n\text{H}_2\text{O}$.

The chemical name for Fe_2O_3 is iron(III) oxide or ferric oxide.

The ". $n\text{H}_2\text{O}$ " part indicates that it is hydrated.

Therefore, the chemical name for rust is **Hydrated Ferric Oxide**.

Step 3: Final Answer:

Rust is chemically known as Hydrated Ferric Oxide.

This corresponds to option (A).

💡 Quick Tip

Remember the difference between "ferrous" and "ferric".
Ferrous refers to the iron(II) ion (Fe^{2+}).
Ferric refers to the iron(III) ion (Fe^{3+}).
Rust contains the more oxidized ferric ion.

95. The monomer of Teflon is X. The number of fluorine atoms in X is

- (A) 2
- (B) 3

- (C) 4
(D) 1

Correct Answer: (C) 4

Solution:

Step 1: Understanding the Question:

We need to identify the monomer of the polymer Teflon and then count the number of fluorine atoms in that monomer molecule.

Step 2: Detailed Explanation:

1. Identify the Polymer and Monomer:

The polymer is Teflon.

The chemical name for Teflon is Polytetrafluoroethylene (PTFE).

The name "Polytetrafluoroethylene" indicates that the repeating monomer unit is "tetrafluoroethylene".

2. Determine the Structure and Formula of the Monomer:

"Ethylene" has the formula C_2H_4 ($CH_2=CH_2$).

"Tetrafluoro" means that the four hydrogen atoms in ethylene are replaced by four fluorine atoms.

So, the structure of the monomer, tetrafluoroethylene (X), is $CF_2=CF_2$.

The chemical formula is C_2F_4 .

3. Count the Fluorine Atoms:

In the molecule C_2F_4 , there are **four** fluorine atoms.

Step 3: Final Answer:

The number of fluorine atoms in the monomer of Teflon is 4.

This corresponds to option (C).

 Quick Tip

Polymer names often give a direct clue to the monomer.

"Poly-" means many, so what follows is usually the name of the monomer.

For example, Poly(ethylene) comes from ethylene, Poly(styrene) from styrene, and Poly(tetrafluoroethylene) from tetrafluoroethylene.

96. Bakelite is an example of

- (A) Thermoplastic Polymer
- (B) Elastomer
- (C) Fibre
- (D) Thermosetting polymer

Correct Answer: (D) Thermosetting polymer

Solution:

Step 1: Understanding the Question:

We need to classify the polymer Bakelite based on its properties.

Step 2: Detailed Explanation:

Let's define the types of polymers listed:

- **Thermoplastic Polymer:** Polymers that become soft and moldable upon heating and solidify upon cooling. The process is reversible.
- **Elastomer:** Polymers with high elasticity that return to their original shape after being stretched.
- **Fibre:** Thread-forming polymers with high tensile strength.
- **Thermosetting Polymer:** Polymers that undergo irreversible chemical changes upon heating, forming a hard, rigid, three-dimensional network structure that cannot be remolded.

Bakelite is a condensation polymer of phenol and formaldehyde.

During its formation, extensive cross-linking creates a rigid 3D network.

Once it is cured by heating, it becomes permanently hard and cannot be melted.

This property is the definition of a thermosetting polymer.

Step 3: Final Answer:

Bakelite is an example of a thermosetting polymer.

This corresponds to option (D).

💡 Quick Tip

Remember the key difference:

Thermo-plastic is like plastic that can be melted and reformed.

Thermo-set is "set" permanently by heat, like baking a cake – you can't un-bake it. Bakelite's rigidity and heat resistance are classic traits of a thermoset.

97. The correct structure of neoprene rubber is

- (A) Structure with C₆H₅ group
- (B) Structure with Cl atom
- (C) Structure with F atom
- (D) Structure with CH₃ group

Correct Answer: (B) Structure with Cl atom

Solution:

Step 1: Understanding the Question:

We need to identify the correct chemical structure for the repeating unit of neoprene rubber from the given options.

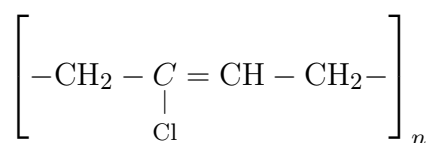
Step 2: Detailed Explanation:

Neoprene is a synthetic rubber.

It is the polymer of the monomer **chloroprene**.

The chemical name for chloroprene is 2-chloro-1,3-butadiene.

The repeating unit of neoprene (polychloroprene) is:



Let's examine the options provided in the image:

- Option 2 shows a Cl (chlorine) atom attached to the second carbon of the four-carbon chain.

This matches the structure of the neoprene repeating unit.

- The other options show incorrect substituent groups (phenyl, fluorine, methyl).

Step 3: Final Answer:

The correct structure is the one containing the chlorine atom.

This corresponds to option (B).

💡 Quick Tip

Associate the names of common synthetic rubbers with their key monomer or feature:

- **Neoprene** → Monomer is **chloroprene** → Contains Chlorine (Cl).
- Natural Rubber → Monomer is isoprene (2-**methyl**-1,3-butadiene) → Contains a Methyl group (CH₃).
- Buna-S (SBR) → Monomers are **B**utadiene and **S**tyrene.

98. Which of the following is NOT regarded as a primary fuel?

- (A) Natural gas
- (B) Coal gas
- (C) Lignite
- (D) Crude oil

Correct Answer: (B) Coal gas

Solution:**Step 1: Understanding the Question:**

We need to distinguish between primary and secondary fuels and identify which of the options is a secondary fuel.

Step 2: Detailed Explanation:

- **Primary Fuels:** These are fuels found in nature and can be used directly or with minimal processing.

- **Secondary Fuels:** These are fuels that are manufactured from primary fuels through a conversion process.

Let's classify the options:

- **Natural gas:** Extracted directly from the earth. It is a **primary fuel**.
- **Coal gas:** Manufactured from coal by destructive distillation. It is a **secondary fuel**.
- **Lignite:** A type of coal mined directly from the ground. It is a **primary fuel**.

- **Crude oil:** Extracted from underground reservoirs. It is a **primary fuel**.

Step 3: Final Answer:

Coal gas is not a primary fuel; it is a secondary fuel.

This corresponds to option (B).

💡 Quick Tip

Ask yourself: "Is this fuel found in nature as-is, or is it made from something else?"
If it's found in nature (like coal, oil, wood), it's primary.
If it's made in a factory (like gasoline, coke, coal gas), it's secondary.

99. pH of acid rain water is generally in the range of

- (A) 1.0 - 3.0
- (B) 3.5 - 5.6
- (C) 5.9 - 6.9
- (D) 7.1 - 7.5

Correct Answer: (B) 3.5 - 5.6

Solution:

Step 1: Understanding the Question:

We need to identify the typical pH range for rain that is classified as acid rain.

Step 2: Detailed Explanation:

- **Normal Rain:** Normal, unpolluted rainwater is naturally acidic due to dissolved atmospheric CO_2 , which forms weak carbonic acid.

This brings the pH of normal rain down to about 5.6.

- **Acid Rain:** Rain with a pH lower than 5.6 is considered acid rain.

It is caused by pollutants like sulfur dioxide and nitrogen oxides forming sulfuric and nitric acid.

The typical pH of acid rain in affected areas is generally between 4.0 and 5.5, but can be lower.

The range 3.5 - 5.6 best represents the general definition and observed values for acid rain.

Step 3: Final Answer:

The pH of acid rain is generally in the range of 3.5 - 5.6.

This corresponds to option (B).

 Quick Tip

Remember that the cutoff for acid rain is a pH of 5.6 (the pH of normal rain).
Any value below this is technically acid rain.
Choose the option that best represents the acidic conditions below 5.6.

100. In _____ part of the atmosphere ozone layer is present.

- (A) Troposphere
- (B) Thermosphere
- (C) Stratosphere
- (D) Mesosphere

Correct Answer: (C) Stratosphere

Solution:

Step 1: Understanding the Question:

The question asks to identify the layer of the Earth's atmosphere that contains the ozone layer.

Step 2: Detailed Explanation:

The Earth's atmosphere is divided into several layers. From the ground up, they are:

1. **Troposphere (0 to 12 km):** Where weather occurs.
2. **Stratosphere (12 to 50 km):** This layer contains the vast majority (about 90%) of atmospheric ozone, forming the protective "ozone layer".

This layer absorbs harmful UV radiation from the Sun, causing the temperature to increase with altitude.

3. **Mesosphere (50 to 85 km):** Where meteors burn up.
4. **Thermosphere (85 to 600 km):** Where the aurora occurs.

Step 3: Final Answer:

The ozone layer is present in the Stratosphere.

This corresponds to option (C).

💡 Quick Tip

A simple mnemonic for the atmospheric layers from the ground up is "The Straight Man's Testicles":

Troposphere, **S**tratosphere, **M**esosphere, **T**hermosphere.

The ozone layer is in the second layer up, the Stratosphere.

Mining Engineering

101. The mode of access to the deposit present in a hilly terrain is

- (A) Incline
- (B) Shaft
- (C) Adit
- (D) Decline

Correct Answer: (C) Adit

Solution:

Step 1: Understanding the Question:

We need to identify the correct mining term for an access tunnel to a mineral deposit located inside a hill.

Step 2: Detailed Explanation:

Let's define the different types of mine access:

- **Shaft:** A vertical tunnel from the surface to a deep deposit. Best for flat terrain.
- **Incline / Decline:** Sloping tunnels. A decline goes down, an incline goes up.
- **Adit:** A **horizontal** or nearly horizontal tunnel driven into the side of a hill or mountain.

This method uses the topography to provide easy, level access for transport and drainage.

Given the "hilly terrain", an adit is the specific term for a horizontal entry into the hillside.

Step 3: Final Answer:

The mode of access to a deposit in a hilly terrain is an adit.

This corresponds to option (C).

💡 Quick Tip

Associate the type of access with the terrain:

- Flat terrain, deep deposit → **Shaft** (vertical).
- Flat terrain, shallow deposit → **Decline/Incline** (sloped).
- Hilly/Mountainous terrain → **Adit** (horizontal).

102. Bore hole survey means

- (A) Measurement of Depth of bore hole
- (B) Measurement of time taken to drill hole
- (C) Measurement of deviation of bore hole
- (D) Identification of properties of bore hole

Correct Answer: (C) Measurement of deviation of bore hole

Solution:

Step 1: Understanding the Question:

The question asks for the definition of the term "borehole survey" in the context of mining or geology.

Step 2: Detailed Explanation:

When a borehole is drilled, especially a deep one, it rarely follows a perfectly straight path.

It can deviate from its intended course due to various factors like the geology of the rock, drilling parameters, and the drill string's alignment.

A **borehole survey** is the specific process of measuring the borehole's actual trajectory or path in three-dimensional space.

This involves measuring its deviation from the vertical (inclination) and its direction in the horizontal plane (azimuth) at various points along its depth.

While measuring depth (A) and identifying properties (D, known as logging) are also parts of borehole analysis, the term "survey" specifically refers to mapping the path and its deviation.

Step 3: Final Answer:

Borehole survey means the measurement of the deviation of the borehole from its planned trajectory.

This corresponds to option (C).

💡 Quick Tip

In surveying, the primary goal is to determine positions and paths. Think of a borehole survey as creating a map of the hole's journey underground, not just measuring its length or what it passed through. This is crucial for accurately locating the position of a mineral deposit intersected by the borehole.

103. The process of treating the holes with chemicals to reduce the friction to the passage of cement through rock mass during Cementation method of shaft skinning is

- (A) Silicatisation
- (B) Pre-Silicatisation
- (C) Sensitization
- (D) Grouting

Correct Answer: (B) Pre-Silicatisation

Solution:

Step 1: Understanding the Question:

This question asks for the name of a specific chemical pre-treatment process used in the cementation method of shaft sinking.

Step 2: Detailed Explanation:

The **Cementation method** is a ground treatment technique used to seal water-bearing fissures and stabilize weak ground, particularly during shaft sinking.

It involves injecting cement grout under pressure into boreholes drilled around the proposed shaft area.

In certain types of rock, especially those with very fine fissures, the penetration of the cement grout can be hindered by high friction and the rock's chemical properties.

To overcome this, a preliminary chemical treatment is applied.

This treatment, known as **Pre-Silicatisation**, involves injecting a solution of sodium silicate.

This chemical reacts with the rock mass in a way that lubricates the fissures, reducing the friction and making it easier for the subsequent cement grout to penetrate deeply and effectively.

Step 3: Final Answer:

The process of chemically treating the holes to aid cement passage is called Pre-Silicatisation.

This corresponds to option (B).

💡 Quick Tip

The prefix "Pre-" indicates that this is a preparatory step done *before* the main process.

Here, it's a pre-treatment before the main cement injection (grouting).

This helps distinguish it from "Silicatisation," which might be part of the main grouting mix itself.

104. The holes left in permanent lining at the curb level to avoid development of hydrostatic pressure behind the lining due to water in shaft sinking are known as

- (A) Product holes
- (B) Weep holes
- (C) Lining holes
- (D) Drain holes

Correct Answer: (B) Weep holes

Solution:

Step 1: Understanding the Question:

The question asks for the specific name of holes designed to relieve water pressure behind the lining of a mine shaft.

Step 2: Detailed Explanation:

When a shaft is sunk through water-bearing strata, water can seep into the space behind the permanent concrete or brick lining.

If this water is not drained, it will accumulate and exert a significant **hydrostatic pressure** on the back of the lining.

This pressure can be strong enough to damage or even cause the failure of the shaft lining.

To prevent this, small-diameter holes or pipes, known as **weep holes**, are intentionally left in the lining, typically at regular intervals.

These holes allow any trapped water to "weep" or drain out from behind the lining into the shaft, where it can be collected and pumped to the surface.

This relieves the hydrostatic pressure and ensures the stability of the shaft lining.

Step 3: Final Answer:

The holes designed for pressure relief behind a shaft lining are called weep holes.

This corresponds to option (B).

💡 Quick Tip

The term "weep hole" is used widely in civil engineering for any opening that allows water to drain from behind a retaining structure, like a retaining wall, bridge abutment, or, in this case, a shaft lining.

The name itself suggests its function: allowing the structure to "weep" out trapped water.

105. The characteristic of an explosive which will give an idea about the storage of explosive under heat, humid and time of exposure

- (A) Stability
- (B) Sensitivity
- (C) Density
- (D) Resistance

Correct Answer: (A) Stability

Solution:

Step 1: Understanding the Question:

The question asks for the property of an explosive that describes its ability to withstand storage conditions (heat, humidity, time) without degrading.

Step 2: Detailed Explanation:

Let's define the properties of explosives:

- **Stability:** This refers to the chemical stability of an explosive.

It is the ability of the explosive to be stored for long periods under varying environmental conditions (like heat, humidity, sunlight) without decomposing, changing its properties, or becoming dangerously unstable.

This directly matches the description in the question.

- **Sensitivity:** This is a measure of the ease with which an explosive can be initiated or detonated by an external stimulus like impact, friction, heat, or shock.

High sensitivity means it detonates easily; low sensitivity means it is safer to handle.

- **Density:** This is the mass per unit volume of the explosive.

It affects the explosive's velocity of detonation and the amount of energy that can be packed into a borehole.

- **Resistance:** This is a general term and could refer to water resistance or electrical resistance, but it is not the primary term for describing storage life.

Step 3: Final Answer:

The characteristic that describes the storage life of an explosive under various conditions is its stability.

This corresponds to option (A).

 Quick Tip

Think of the difference between stability and sensitivity:

- **Stability** is about the explosive's "shelf life" and its ability to resist decomposing on its own.

- **Sensitivity** is about its reaction to being "provoked" by an external shock or spark.

A good explosive is stable (safe to store) but sensitive enough to be detonated when intended.

106. The type of explosive used for solid blasting is

- (A) P1
- (B) P3
- (C) P4
- (D) P5

Correct Answer: (D) P5

Solution:

Step 1: Understanding the Question:

We need to identify the specific class of permitted explosives that is suitable for "solid blasting" in underground coal mines.

Step 2: Detailed Explanation:

In underground coal mining, special "permitted" explosives are used to minimize the risk of igniting flammable methane gas or coal dust.

These explosives are classified into different categories (P1, P2, P3, P4, P5) based on their safety characteristics and intended use.

Solid Blasting is a technique where blasting is carried out in a coal face "off the solid," meaning without a pre-existing free face created by cutting or shearing.

This is a particularly hazardous operation because the explosive energy is highly confined, increasing the risk of a gas or dust ignition.

Therefore, it requires the safest category of explosives.

- **P1 Explosives:** General-purpose explosives for use in coal where there is minimal gas.
- **P3 Explosives:** For use in stone drifts or seams with higher gas content.
- **P4 Explosives:** This category is obsolete.
- **P5 Explosives:** These are "sheathed" or equivalent-to-sheathed (EQS) explosives designed with the highest degree of safety.

They produce a low-temperature, short-duration flame and are specifically approved for the most hazardous applications, including solid blasting and blasting in fractured ground.

Step 3: Final Answer:

The type of explosive used for solid blasting is the P5 category.

This corresponds to option (D).

💡 Quick Tip

The "P" number in permitted explosives indicates the level of safety, with a higher number generally signifying a safer explosive for gassier conditions.

P5 is the "top tier" of safety, required for the most dangerous blasting scenario, which is solid blasting.

107. The prime charge in a detonator is

- (A) PETN
- (B) TNT
- (C) ASA
- (D) Ammonium Nitrate

Correct Answer: (C) ASA

Solution:

Step 1: Understanding the Question:

The question asks to identify the chemical compound or mixture typically used as the primary or priming charge in a detonator.

Step 2: Detailed Explanation:

A detonator is a device used to trigger a larger, less sensitive main explosive charge.

It works as a multi-stage explosive train:

1. **Initiation:** An external stimulus (like an electric current heating a fuse head or a shock from a fuse) initiates the first charge.
2. **Prime Charge (Primary Explosive):** This is a highly sensitive explosive that detonates from the initial stimulus.

Its purpose is to generate enough shock to detonate the next, more powerful charge.

A very common prime charge composition is **ASA**, which is a mixture of **Azide** (Lead Azide), **Styphnate** (Lead Styphnate), and powdered **Aluminium**.

3. **Base Charge (Secondary Explosive):** This is a less sensitive but more powerful explosive that makes up the bulk of the detonator's power.

It is detonated by the shockwave from the prime charge.

PETN (Pentaerythritol tetranitrate) and **TNT** (Trinitrotoluene) are powerful secondary explosives and are commonly used as the base charge in detonators.

Ammonium Nitrate is the main component of ANFO, a bulk blasting agent, not a primary charge.

Step 3: Final Answer:

The prime charge in a detonator is commonly ASA.

This corresponds to option (C).

💡 Quick Tip

Think of a detonator like lighting a fire.

You use a match (the initiation) to light the kindling (the prime charge, e.g., ASA), which then provides a strong enough flame to ignite the main log (the base charge, e.g., PETN).

108. Identify the correct sequence of blasting operations in Underground coal mines

- (A) Drilling, Charging, Stemming, Blasting, Mucking
- (B) Drilling, Charging, Mucking, Stemming, Blasting
- (C) Mucking, Drilling, Charging, Stemming, Blasting
- (D) Drilling, Stemming, Charging, Mucking, Blasting

Correct Answer: (A) Drilling, Charging, Stemming, Blasting, Mucking

Solution:

Step 1: Understanding the Question:

We need to arrange the listed mining activities into the correct chronological order for a standard blasting cycle.

Step 2: Detailed Explanation:

Let's analyze the logical flow of a drill-and-blast operation:

1. **Drilling:** The first step is to create holes in the rock or coal face where the explosives will be placed.
2. **Charging:** The drilled holes are then "charged" or loaded with the explosives and detonators.
3. **Stemming:** After charging, the mouth of the hole is plugged with an inert material like clay or sand stemming.

This confines the explosive energy within the hole, ensuring it is used to break the rock rather than blowing out of the hole (a "blowout").

4. **Blasting:** Once the area is clear of personnel, the explosive charges are detonated.
5. **Mucking:** After the blast and a safety period for fumes to clear, the broken rock or coal (the "muck") is loaded and hauled away.

This sequence is Drilling → Charging → Stemming → Blasting → Mucking.

Step 3: Final Answer:

The correct sequence of operations is Drilling, Charging, Stemming, Blasting, Mucking.

This corresponds to option (A).

💡 Quick Tip

Think of the process logically, like baking a cake.

You have to prepare the pan (**Drilling**), put the batter in (**Charging**), put it in the oven (**Stemming/Blasting**), and then take it out to serve (**Mucking**).

You can't muck before you blast, and you can't blast before you drill and charge the holes.

109. The texture of rocks formed due to consolidation of magma nearer to the surface is

- (A) Course texture
- (B) Medium texture
- (C) Fine texture
- (D) Conglomerate

Correct Answer: (C) Fine texture

Solution:

Step 1: Understanding the Question:

The question asks about the texture of igneous rocks that form from magma cooling near the Earth's surface.

Step 2: Detailed Explanation:

The texture of an igneous rock refers to the size, shape, and arrangement of its mineral crystals.

The most important factor controlling crystal size is the cooling rate of the magma or lava.

- **Slow Cooling:** When magma cools slowly, deep within the Earth's crust (intrusive or plutonic rocks), the atoms have a long time to migrate and arrange themselves into large, well-formed crystals.

This results in a **coarse-grained texture** (also called phaneritic), where individual crystals are visible to the naked eye (e.g., Granite).

- **Rapid Cooling:** When magma or lava cools quickly, at or near the Earth's surface (extrusive or volcanic rocks), the atoms have very little time to form large crystals.

This results in a **fine-grained texture** (also called aphanitic), where the individual crystals are too small to be seen without a microscope (e.g., Basalt, Rhyolite).

Since the question specifies consolidation "nearer to the surface," this implies rapid cooling.

Conglomerate (D) is a sedimentary rock, not an igneous rock.

Step 3: Final Answer:

Rocks formed from magma cooling near the surface have a fine texture.

This corresponds to option (C).

💡 Quick Tip

Remember the simple relationship:

Slow cooling = Large crystals = Coarse texture (Intrusive).

Fast cooling = Small crystals = Fine texture (Extrusive).

The closer to the surface, the faster the cooling.

110. Identify the Correct sequence of formation of coal

(A) Lignite → Peat → Bituminous → Anthracite

(B) Peat → Lignite → Bituminous → Anthracite

(C) Anthracite → Bituminous → Lignite → Peat

(D) Bituminous → Anthracite → Lignite → Peat

Correct Answer: (B) Peat → Lignite → Bituminous → Anthracite

Solution:

Step 1: Understanding the Question:

The question asks for the correct sequence of coal formation, known as coalification, which represents the increasing rank of coal.

Step 2: Detailed Explanation:

Coalification is the process by which plant matter is converted into coal under the influence of heat and pressure over millions of years.

The process involves a progressive increase in carbon content and a decrease in moisture and volatile matter.

The sequence of coal ranks, from lowest to highest, is as follows:

1. **Peat:** This is the initial stage.

It is an accumulation of partially decayed vegetation or organic matter.

It has high moisture content and low carbon content.

2. **Lignite:** With increased pressure and temperature, peat is transformed into lignite (also called brown coal).

It is the lowest rank of coal.

3. **Bituminous Coal:** Further burial, heat, and pressure convert lignite into bituminous coal (also called soft coal).

This is a large and important category of coal, often subdivided into sub-bituminous and bituminous.

4. **Anthracite:** This is the highest rank of coal (also called hard coal).

It forms from bituminous coal under the highest heat and pressure.

It has the highest carbon content, the lowest moisture, and burns with the cleanest flame.

The correct sequence of formation is therefore Peat $\rightarrow$ Lignite $\rightarrow$ Bituminous $\rightarrow$ Anthracite.

Step 3: Final Answer:

The correct sequence of coal formation is Peat $\rightarrow$ Lignite $\rightarrow$ Bituminous $\rightarrow$ Anthracite.

This corresponds to option (B).

💡 Quick Tip

Think of it as squeezing water out and concentrating carbon.

Peat is like a wet sponge of plant material.

Squeeze it a bit, you get **Lignite**.

Squeeze harder, you get **Bituminous**.

Squeeze the hardest, you get **Anthracite**, which is almost pure carbon.

The sequence is alphabetical if you remember the four main types: A, B, L, P. The process goes in reverse alphabetical order of rank: P $\rightarrow$ L $\rightarrow$ B $\rightarrow$ A.

111. The vertical displacement of seam in a normal fault is known as

- (A) Throw
- (B) Want
- (C) Hade
- (D) Inclination

Correct Answer: (A) Throw

Solution:

Step 1: Understanding the Question:

The question asks for the geological term for the vertical component of displacement across a fault.

Step 2: Detailed Explanation:

In structural geology, the displacement of rock layers or a seam across a fault is described by several components:

- **Throw:** This is the measure of the **vertical displacement** of the fault.

It is the vertical distance between a point on one side of the fault and the corresponding point on the other side.

This directly matches the question's description.

- **Heave:** This is the measure of the **horizontal displacement** of the fault.

- **Hade:** This is the angle that the fault plane makes with the vertical.

It is the complement of the dip angle ($\text{Hade} = 90^\circ - \text{Dip}$).

- **Want:** This is a mining term for a barren area where the coal seam is missing, which could be caused by a fault.

- **Inclination:** This is a general term for a slope or angle, often used for the dip of strata.

Step 3: Final Answer:

The vertical displacement of a seam in a normal fault is known as the throw.

This corresponds to option (A).

💡 Quick Tip

Remember the two main components of fault displacement:

- **Throw** is the vertical part (what you "throw" up or down).
- **Heave** is the horizontal part (what you "heave" sideways).

112. Which of the following is an example of sedimentary rock

- (A) Granite
- (B) Basalt
- (C) Sandstone

(D) Schist

Correct Answer: (C) Sandstone

Solution:

Step 1: Understanding the Question:

We need to identify the sedimentary rock from the given list of four rock types.

Step 2: Detailed Explanation:

Let's classify each rock:

- **Granite:** This is an **igneous** rock.

Specifically, it is an intrusive (or plutonic) igneous rock, formed from the slow cooling of magma beneath the Earth's surface.

- **Basalt:** This is an **igneous** rock.

Specifically, it is an extrusive (or volcanic) igneous rock, formed from the rapid cooling of lava on the Earth's surface.

- **Sandstone:** This is a classic **sedimentary** rock.

It is formed from the compaction and cementation (lithification) of sand-sized grains of minerals, rock, or organic material.

- **Schist:** This is a **metamorphic** rock.

It is formed when other rocks (like shale) are subjected to intense heat and pressure, causing the mineral grains to recrystallize and align in layers or bands (a property called foliation).

Step 3: Final Answer:

Sandstone is the example of a sedimentary rock in the list.

This corresponds to option (C).

💡 Quick Tip

Associate the rock types with their formation process:

- **Igneous** = from "fire" (molten magma/lava).
- **Sedimentary** = from "sediment" (sand, mud, shells).
- **Metamorphic** = "morphed" or changed by heat and pressure.

113. The point where an earthquake originates is known as

- (A) Epicentre
- (B) Iso-seismal
- (C) Co-seismal
- (D) Focus

Correct Answer: (D) Focus

Solution:

Step 1: Understanding the Question:

The question asks for the specific term for the point of origin of an earthquake.

Step 2: Detailed Explanation:

Let's define the given terms related to earthquakes:

- **Focus (or Hypocenter):** This is the point **within the Earth** along a fault where the initial rupture and release of energy occurs.

It is the true origin point of the earthquake.

- **Epicenter:** This is the point on the **Earth's surface** directly above the focus.

It is the location where the earthquake's effects are often the most intense.

- **Iso-seismal line:** A line on a map connecting points of equal earthquake intensity.

- **Co-seismal line:** A line on a map connecting points where the seismic waves arrive at the same time.

The question asks for the point where the earthquake *originates*, which is the focus.

Step 3: Final Answer:

The point where an earthquake originates is known as the focus.

This corresponds to option (D).

 Quick Tip

Remember the difference:

- **Focus** is the underground origin.
- **Epicenter** is the surface point directly above the focus.

Think of "Epi-" as a prefix meaning "upon" or "over," so the epicenter is upon the center.

114. Which hypothesis explains the origin of the Earth ?

- (A) Nebular Hypothesis
- (B) Plate Tectonic Hypothesis
- (C) Big Bang Hypothesis
- (D) Null Hypothesis

Correct Answer: (A) Nebular Hypothesis

Solution:

Step 1: Understanding the Question:

We need to identify the scientific hypothesis that specifically explains the formation of the Earth and the solar system.

Step 2: Detailed Explanation:

Let's analyze the given hypotheses:

- **Nebular Hypothesis:** This is the most widely accepted model that explains the formation and evolution of the Solar System.

It proposes that the Sun, the Earth, and all other planets formed from a giant, rotating cloud of interstellar gas and dust called a solar nebula.

- **Plate Tectonic Hypothesis (or Theory):** This theory explains the large-scale movements of the Earth's lithosphere (the crust and upper mantle).

It explains phenomena like continental drift, earthquakes, and volcanic activity, but not the origin of the Earth itself.

- **Big Bang Hypothesis (or Theory):** This is the leading cosmological model for the observable **universe**.

It describes how the universe expanded from an initial state of extremely high density and temperature, but it does not specifically detail the formation of the Earth.

- **Null Hypothesis:** This is a term used in statistics and scientific testing.

It is a default statement that there is no relationship between two measured phenomena. It has nothing to do with cosmology.

Step 3: Final Answer:

The Nebular Hypothesis explains the origin of the Earth.

This corresponds to option (A).

💡 Quick Tip

Distinguish the scales of these theories:

- **Big Bang** = Origin of the entire Universe.
- **Nebular Hypothesis** = Origin of our Solar System (including Earth).
- **Plate Tectonics** = How the surface of the Earth moves and changes.

115. The process of "Erosion" includes

- (A) Deflation and Abrasion
- (B) Disintegration and Decomposition
- (C) Abrasion and disintegration
- (D) Disintegration and Transportation

Correct Answer: (D) Disintegration and Transportation

Solution:

Step 1: Understanding the Question:

The question asks for the best definition of the geological process of erosion from the given options.

Step 2: Detailed Explanation:

Let's define the key geological terms:

- **Weathering:** This is the in-situ (in one place) breakdown of rocks, soil, and minerals.

It includes physical breakdown (**disintegration**, e.g., by frost) and chemical breakdown (**decomposition**, e.g., by acid rain).

Option (B) describes weathering.

- **Erosion:** This is the process where weathered material (like soil, rock fragments) is moved from one location to another.

It involves both the initial picking-up or breakdown of the material and its subsequent **transportation** by agents like water, wind, ice, or gravity.

Therefore, erosion encompasses both the disintegration (or detachment) of material and its transportation.

- **Abrasion** and **Deflation** (Option A) are specific mechanisms of erosion, primarily by wind, but they don't define the entire process.

Abrasion is the scraping/wearing away, and deflation is the lifting of loose material.

Option (D) provides the most complete definition by including both the breakdown/detachment aspect (**Disintegration**) and the crucial movement aspect (**Transportation**).

Step 3: Final Answer:

The process of erosion includes disintegration and transportation.

This corresponds to option (D).

 Quick Tip

A simple way to remember the difference:

Weathering makes the mess (breaks the rock down).

Erosion cleans it up (carries the pieces away).

116. Identify the minerals based the decreasing order of Moh's Scale of Hardness

(A) Topaz > Quartz > Orthoclase > Fluorite

(B) Quartz > Fluorite > Topaz > Orthoclase

(C) Fluorite > Quartz > Orthoclase > Topaz

(D) Topaz > Orthoclase > Fluorite > Quartz

Correct Answer: (A) Topaz > Quartz > Orthoclase > Fluorite

Solution:

Step 1: Understanding the Question:

We need to arrange the given four minerals in order from hardest to softest, according to the Mohs scale of mineral hardness.

Step 2: Detailed Explanation:

The Mohs scale is a qualitative ordinal scale that characterizes the scratch resistance of various minerals through the ability of a harder material to scratch a softer material.

Let's recall the hardness values for the minerals listed:

- **Topaz:** Hardness = 8 - **Quartz:** Hardness = 7 - **Orthoclase** (a type of feldspar): Hardness = 6 - **Fluorite:** Hardness = 4

Arranging these in decreasing order of hardness (from hardest to softest) gives:

Topaz (8) > Quartz (7) > Orthoclase (6) > Fluorite (4)

Step 3: Final Answer:

The correct decreasing order of hardness is Topaz > Quartz > Orthoclase > Fluorite.

This corresponds to option (A).

💡 Quick Tip

It is very helpful to memorize the 10 minerals of the Mohs scale for geology questions. A common mnemonic is: "Tall Girls Can Fight And Often Quit, They Cry Diamonds."

1. Talc 2. Gypsum 3. Calcite 4. Fluorite 5. Apatite 6. Orthoclase 7. Quartz 8. Topaz 9. Corundum 10. Diamond

117. The type of fold in which younger rock formation is present outside and older rock formation is present inside the fold is known as

- (A) Syncline
- (B) Anticline
- (C) Chevron fold
- (D) Recumbent fold

Correct Answer: (B) Anticline

Solution:

Step 1: Understanding the Question:

The question asks to identify the type of geological fold based on the relative age of the rock layers within it.

Step 2: Detailed Explanation:

Geological folds are wavelike structures formed when rock layers are bent or buckled.

The two most basic types of folds are anticlines and synclines:

- **Anticline:** This is an arch-like fold where the rock layers are bent upwards.

Due to this arch shape, the oldest rock layers are found in the core (the center) of the fold.

As you move outwards from the center, the rock layers become progressively younger.

The description "older rock formation is present inside the fold" and "younger rock formation is present outside" perfectly matches the definition of an anticline.

- **Syncline:** This is a trough-like fold where the rock layers are bent downwards.

In a syncline, the youngest rock layers are in the core, and the layers get older as you move outwards.

- **Chevron fold** and **Recumbent fold** describe the shape and orientation of the fold, not the age relationship of the strata.

A chevron fold has sharp, angular hinges, and a recumbent fold is an anticline or syncline that has been overturned onto its side.

Step 3: Final Answer:

A fold with older rocks in the core and younger rocks on the outside is an anticline.

This corresponds to option (B).

💡 Quick Tip

A simple way to remember:

Anticline = **A**-shaped arch, with **A**ncient rocks in the middle.

Syncline = **S**inks down like a trough, with younger rocks in the middle.

118. Which metamorphic rock is formed by the recrystallization of Limestone ?

- (A) Slate
- (B) Schist
- (C) Marble
- (D) Gneiss

Correct Answer: (C) Marble

Solution:

Step 1: Understanding the Question:

The question asks to identify the metamorphic rock that has limestone as its parent rock (protolith).

Step 2: Detailed Explanation:

Metamorphic rocks are formed when existing rocks are changed by heat, pressure, or chemical reactions.

Let's look at the parent rocks for the given options:

- **Marble:** This is a non-foliated metamorphic rock composed primarily of recrystallized carbonate minerals, most commonly calcite or dolomite.

The parent rock of marble is **limestone** (composed of calcite) or dolostone (composed of dolomite).

This matches the question.

- **Slate:** This is a fine-grained, foliated metamorphic rock formed from the low-grade metamorphism of shale or mudstone.

- **Schist:** This is a medium- to coarse-grained, foliated metamorphic rock.

It forms at a higher grade of metamorphism than slate, often from shale or mudstone as the original parent rock.

- **Gneiss:** This is a high-grade, foliated metamorphic rock characterized by distinct bands of different minerals.

It can form from various parent rocks, including granite or sedimentary rocks like shale.

Step 3: Final Answer:

The metamorphic rock formed from the recrystallization of limestone is marble.

This corresponds to option (C).

💡 Quick Tip

It's useful to remember some common parent rock → metamorphic rock pairs:

- Limestone → Marble
- Sandstone → Quartzite
- Shale → Slate → Phyllite → Schist → Gneiss (this shows increasing metamorphic grade)
- Granite → Gneiss

119. In a mine working the roof fall which takes place soon after withdrawal of supports is called

- (A) Local fall
- (B) Main fall
- (C) Air blast

(D) Rock burst

Correct Answer: (A) Local fall

Solution:

Step 1: Understanding the Question:

The question asks for the specific term for a roof collapse that occurs shortly after supports are removed in a mining operation, like in the goaf of a longwall or depillaring panel.

Step 2: Detailed Explanation:

In underground mining methods where pillars are extracted or longwall panels advance, an area called the "goaf" or "gove" is created where the roof is intentionally allowed to collapse.

This collapse happens in stages:

- **Local Fall (or Immediate Fall):** This refers to the collapse of the immediate roof strata, which are the rock layers directly above the coal seam.

This fall happens relatively soon after the supports (like hydraulic chocks in a longwall) are advanced or withdrawn.

It involves the rock breaking and falling into the void under its own weight.

- **Main Fall:** This is a much larger and more significant collapse of the stronger, thicker main roof strata that lie above the immediate roof.

The main fall occurs periodically, after a significant area of the goaf has been created, as the massive rock beds break under the immense stress.

An **Air blast** is a violent rush of air caused by a large, sudden roof fall (like a main fall) that displaces a huge volume of air.

A **Rock burst** is a violent, explosive failure of rock under high stress, different from a gravity-induced fall.

The question specifies a fall that takes place "soon after withdrawal of supports," which is the definition of a local fall.

Step 3: Final Answer:

The roof fall that occurs soon after support withdrawal is called a local fall.

This corresponds to option (A).

💡 Quick Tip

Think of the roof as having two parts: a weak lower layer and a strong upper layer.
The **Local fall** is the weak layer collapsing right away.
The **Main fall** is the strong layer breaking later, after a large area is exposed.

120. The sudden release of elastic strain energy stored in pillars result in violent burst of coal pillars is called

- (A) Air blast
- (B) Rock burst
- (C) Coal bump
- (D) Premature collapse

Correct Answer: (C) Coal bump

Solution:

Step 1: Understanding the Question:

The question asks for the specific term for the violent failure of coal pillars due to the sudden release of stored strain energy.

Step 2: Detailed Explanation:

In deep mines, the immense weight of the overlying rock (overburden) creates very high stresses in the rock mass and in the coal pillars left for support.

Coal, like other rock materials, is elastic and can store a significant amount of strain energy when compressed.

- **Coal Bump (or Bump):** This is the specific term used in coal mining to describe the sudden, violent failure of a coal pillar or face.

It occurs when the stored elastic strain energy is released explosively, causing the coal to be ejected with great force.

This phenomenon is a major hazard in deep coal mines.

- **Rock Burst:** This is a more general term used in hard rock mining for the same phenomenon of violent rock failure due to the sudden release of stored strain energy.

While technically similar, "coal bump" is the more specific and appropriate term for the failure of coal pillars.

- **Air Blast:** This is a secondary effect, a powerful gust of wind caused by the sudden displacement of air from a large rock fall or a bump.

- **Premature Collapse:** This is a more general term for any structural failure that happens earlier than expected, but it doesn't specifically describe the violent, energy-driven nature of a bump.

Step 3: Final Answer:

The violent burst of coal pillars from the release of stored energy is called a coal bump.

This corresponds to option (C).

💡 Quick Tip

Remember the distinction:

- **Rock Burst** is the general term for hard rock.
- **Coal Bump** is the specific term for coal.

Both describe the same physical principle: explosive release of stored strain energy.

121. The underground road way driven through stone to connect two or more coal seams is known as

- (A) Tunnel
- (B) Gallery
- (C) Cross-cut
- (D) Cross measure drift

Correct Answer: (D) Cross measure drift

Solution:

Step 1: Understanding the Question:

We need to identify the specific mining term for a roadway that is driven through non-coal rock strata to connect different coal seams.

Step 2: Detailed Explanation:

Let's define the terms for underground roadways:

- **Gallery (or Roadway, Heading):** A general term for a horizontal or near-horizontal underground passage, typically driven within the coal seam itself.
- **Cross-cut:** A passage driven between two parallel galleries within the same coal seam, often for ventilation or transport.
- **Tunnel:** A general civil engineering term for an underground passage.

In mining, it's sometimes used for main access ways from the surface, but a more specific term is usually preferred.

- **Cross Measure Drift (or Drivage):** This is the specific term for a roadway driven through the "measures" or rock strata that lie between coal seams.

Its purpose is to "cross" these strata to connect workings in one seam to another, either at a different level or in a different area.

The key parts of the definition are "driven through stone" (cross measure) and "to connect two or more coal seams".

Step 3: Final Answer:

An underground roadway driven through rock to connect different coal seams is known as a cross measure drift.

This corresponds to option (D).

💡 Quick Tip

Break down the term "Cross Measure Drift":

- **Drift** is a horizontal or near-horizontal mine opening.
- **Measure** refers to the rock strata or layers.
- **Cross** means it is driven across the strata, not along them (like a roadway in a seam).

122. The shape of the pillar in flat and moderately inclined seams for adoption of shuttle car and locomotive is

- (A) Rhombus
- (B) Circular
- (C) Square
- (D) Rectangle

Correct Answer: (A) Rhombus

Solution:

Step 1: Understanding the Question:

The question asks for the preferred shape of coal pillars in a room-and-pillar mining system that uses trackless mobile equipment like shuttle cars.

Step 2: Detailed Explanation:

In traditional room-and-pillar mining, pillars were often square or rectangular.

This layout creates sharp, 90-degree turns at intersections.

While this is manageable for conveyor belts or rail-bound systems, it is inefficient and hazardous for large, rubber-tired vehicles like shuttle cars and load-haul-dump (LHD) machines.

These vehicles have a large turning radius and navigating sharp right-angle turns is slow, difficult, and can cause collisions.

To improve the efficiency and safety of these trackless mining systems, pillars are often developed in a **rhombus** or diamond shape.

This creates intersections where the turns are obtuse (greater than 90 degrees), allowing for smoother, faster, and safer maneuvering of the equipment.

Circular pillars are sometimes used but are much more difficult to create and are not standard for this purpose.

Step 3: Final Answer:

The preferred shape of pillars for shuttle car and locomotive haulage is a rhombus.

This corresponds to option (A).

💡 Quick Tip

Think about driving a car.

It's much easier and faster to navigate a gentle curve or an intersection at an angle greater than 90 degrees than it is to make a sharp 90-degree turn.

The same principle applies to large mining machinery.

Rhombus-shaped pillars facilitate this smoother traffic flow.

123. The method mining suitable for the extraction of coal seams with dirt bands is

- (A) Board and Pillar method
- (B) Longwall method
- (C) Slicing method
- (D) Hydraulic Mining

Correct Answer: (B) Longwall method

Solution:

Step 1: Understanding the Question:

We need to identify the mining method that is well-suited for extracting coal seams that contain one or more layers of non-coal material, known as "dirt bands" or "partings".

Step 2: Detailed Explanation:

- **Board and Pillar method (or Room and Pillar):** This method involves driving a series of entries (rooms) into the seam, leaving behind pillars of coal to support the roof.

While it is flexible, selectively mining around or excluding a dirt band within the extraction height can be difficult and inefficient, often leading to dilution (mixing waste rock with coal).

- **Longwall method:** This is a high-production method where a large block of coal (a "panel") is extracted in a single, continuous operation.

A machine called a shearer cuts coal from a long face (up to 400m).

Modern longwall shearers are very adaptable.

They can be programmed to cut at specific horizons, allowing them to selectively cut the coal sections while leaving the dirt band either in the roof or the floor, if possible.

Even if the band must be cut, it can be handled more systematically than in room and pillar.

This ability to manage in-seam waste makes longwall a suitable method.

- **Slicing method:** This refers to methods for extracting very thick seams by taking them out in horizontal or inclined "slices".

While it can handle thick seams, it doesn't inherently offer a better solution for thin dirt bands than other methods.

- **Hydraulic Mining:** This uses high-pressure water jets to break and transport coal.

It is not selective and would mix the coal and dirt band together, making it unsuitable.

Step 3: Final Answer:

The longwall method is highly suitable for extracting coal seams with dirt bands due to the selective cutting capabilities of modern shearers.

This corresponds to option (B).

💡 Quick Tip

Longwall mining is a highly mechanized and systematic method. This high degree of control allows for "selective mining," where the cutting machine can be precisely positioned to either avoid cutting a waste band or to handle it separately, thus improving the quality of the mined coal.

124. In Board and Pillar method of mining with Hydraulic stowing, the preferred method of extraction of pillars is

- (A) Step diagonal
- (B) Vertical line
- (C) Diagonal line
- (D) Straight line

Correct Answer: (A) Step diagonal

Solution:

Step 1: Understanding the Question:

The question asks for the preferred pillar extraction (depillaring) pattern when hydraulic stowing is used to fill the void.

Step 2: Detailed Explanation:

Hydraulic stowing involves filling the goaf (the void left after extraction) with a sand-water slurry to control roof subsidence.

This requires the goaf to be systematically and securely barricaded to contain the slurry.

- **Diagonal line** of extraction is a common method, but it can create a long, continuous goaf line which is difficult to manage with stowing.

- **Step diagonal line** of extraction is a modification where pillars are extracted in a staggered or "stepped" pattern.

This method has significant advantages when combined with hydraulic stowing:

1. It breaks the long goaf line into smaller, more manageable sections.
2. It makes it easier to erect strong barricades to contain the sand.
3. It improves overall roof control and safety during the simultaneous process of extraction and filling.

For these reasons, the step diagonal method is the preferred technique when depillaring with hydraulic stowing.

Step 3: Final Answer:

The preferred method of extraction is the step diagonal line.

This corresponds to option (A).

💡 Quick Tip

The choice of extraction pattern is often determined by the method of goaf management.

For caving, a long, straight or diagonal line might be preferred to promote roof fall.

For stowing, a method that creates contained, easily-barricaded sections, like the step diagonal line, is superior.

125. The centre-to-centre size of coal pillars developed in Board and Pillar mining is 40 m x 40 m working at 280 m depth has 4.0 m wide galleries. What area of roof does a pillar support according to tributary area theory?

- (A) 4480 sq.m
- (B) 1936 sq.m
- (C) 1600 sq.m
- (D) 1296 sq.m

Correct Answer: (C) 1600 sq.m

Solution:

Step 1: Understanding the Question:

We need to calculate the area of the roof supported by a single coal pillar using the Tributary Area Theory.

Step 2: Key Formula or Approach:

The **Tributary Area Theory (TAT)** is a simple method for estimating the load on a pillar.

It assumes that each pillar supports the weight of the overlying rock in a "tributary area" that extends halfway to each adjacent pillar.

In a regular grid of pillars, this tributary area is simply the centre-to-centre distance between pillars multiplied in both directions.

$$\text{Tributary Area} = (\text{Centre-to-centre distance})_1 \times (\text{Centre-to-centre distance})_2$$

Step 3: Detailed Explanation:

The problem gives the following data:

- Centre-to-centre size of pillars = 40 m x 40 m
- Depth = 280 m (This is extra information, not needed for calculating the area).
- Gallery width = 4.0 m (This is also extra information).

According to the Tributary Area Theory, the area of roof supported by one pillar is equal to the square of the centre-to-centre distance.

$$\text{Area Supported} = 40 \text{ m} \times 40 \text{ m} = 1600 \text{ m}^2$$

Step 4: Final Answer:

The area of roof a pillar supports according to the tributary area theory is 1600 sq.m.

This corresponds to option (C).

💡 Quick Tip

Don't get confused by extra information.

The Tributary Area Theory is very straightforward: the supported area is simply the centre-to-centre pillar spacing squared (for a square grid).

The gallery width would be needed to find the actual size of the pillar itself (Pillar size = Centre distance - Gallery width).

The depth would be needed to calculate the stress on the pillar (Stress = Overburden Pressure * Tributary Area / Pillar Area).

126. In sand stowing in-correct Hydraulic profile leads to

- (A) Formation of cavities in the flow
- (B) Damage of pipes
- (C) Jamming of pipes
- (D) Setup pulsation

Correct Answer: (A) Formation of cavities in the flow

Solution:

Step 1: Understanding the Question:

The question asks about the consequences of an improperly designed pipeline profile in a hydraulic stowing system.

Step 2: Detailed Explanation:

Hydraulic stowing transports a sand-water slurry.

The hydraulic profile refers to the vertical layout of the pipeline.

An incorrect profile, particularly one with high points (summits) where the pipe rises and then falls, can lead to serious flow problems.

At these high points, if the pressure inside the pipe drops below the vapor pressure of the water, the water can start to boil even at ambient temperature.

This phenomenon is called **cavitation**, and it leads to the **formation of vapor-filled cavities** or bubbles in the flow.

These cavities can disrupt the continuous flow of the slurry.

When these bubbles travel to a region of higher pressure downstream, they collapse violently, which can cause severe **damage to pipes**, create noise and vibration (**pulsation**), and ultimately lead to flow instabilities that can contribute to the settling of sand and **jamming of pipes**.

However, the primary hydraulic phenomenon that occurs due to a negative pressure gradient in an incorrect profile is the formation of cavities.

The other options are consequences of this primary issue.

Step 3: Final Answer:

An incorrect hydraulic profile leads to the formation of cavities in the flow.

This corresponds to option (A).

💡 Quick Tip

While pipe jamming is a major operational problem in stowing, its root cause from a fluid dynamics perspective in an incorrect profile is often cavitation.

The pressure dropping to near-vacuum at high points in the pipeline creates vapor pockets, disrupting the flow and leading to other problems.

127. The width of the heading in Board and Pillar method of working depends on

- (A) Depth of working
- (B) Size of pillars
- (C) Face machinery used

(D) Ventilation required

Correct Answer: (C) Face machinery used

Solution:

Step 1: Understanding the Question:

We need to identify a key factor that determines the width of a gallery or heading in a modern Board and Pillar mine.

Step 2: Detailed Explanation:

While several factors influence the final design of a mine layout, the width of the headings is constrained by both geotechnical stability and operational requirements.

- **Depth of working** and rock strength determine the maximum *safe* width that can be excavated without excessive support.
- **Size of pillars** is related to the gallery width through the extraction ratio, but doesn't directly determine the width.
- **Ventilation required** can be met by adjusting both width and height to achieve a target cross-sectional area.
- **Face machinery used:** In modern mechanized Board and Pillar mining, the dimensions of the equipment are a critical and often primary constraint.

The heading must be wide enough to allow the cutting machine (e.g., a Continuous Miner) to operate effectively and to provide sufficient clearance for loading and hauling equipment (e.g., Shuttle Cars) to maneuver.

The size of the machinery often dictates the minimum practical width of the heading, which must then be checked for geotechnical stability.

Therefore, the choice of face machinery is a direct and major determinant of the heading width.

Step 3: Final Answer:

The width of the heading is heavily dependent on the face machinery used.

This corresponds to option (C).

 Quick Tip

In modern mining, the system is designed around the equipment.

The mine layout (like gallery width and intersection angles) must accommodate the machinery that will be used.

This operational constraint is often considered first, and then the design is verified for safety and stability based on rock mechanics principles.

128. A coal heading is of 4 m wide and 2.5 m height has an advance of 1 m per cycle. The amount of explosive used is 5 kg per blast. Taking specific gravity of coal as 1.2 t/m^3 . The powder factor is

- (A) 1.55 te/kg
- (B) 2.40 te/kg
- (C) 2.99 te/kg
- (D) 3.32 te/kg

Correct Answer: (B) 2.40 te/kg

Solution:

Step 1: Understanding the Question:

We need to calculate the powder factor for a blasting operation in a coal heading, given the dimensions of the heading, the advance per blast, and the amount of explosive used.

Step 2: Key Formula or Approach:

Powder Factor (PF) is a measure of the efficiency of explosive use.

It is defined as the mass of rock or coal broken per unit mass of explosive used.

$$\text{Powder Factor} = \frac{\text{Mass of coal broken (in tonnes)}}{\text{Mass of explosive used (in kg)}}$$

First, we need to calculate the volume of coal broken, and then its mass.

Step 3: Detailed Explanation:

1. Calculate the volume of coal broken per blast:

Volume = Width $\times$ Height $\times$ Advance

$$\text{Volume} = 4 \text{ m} \times 2.5 \text{ m} \times 1 \text{ m} = 10 \text{ m}^3$$

2. Calculate the mass of coal broken:

The specific gravity of coal is given as 1.2 t/m^3 .

This is the density of the coal.

Mass = Volume $\times$ Density

$$\text{Mass} = 10 \text{ m}^3 \times 1.2 \text{ t/m}^3 = 12 \text{ tonnes (te)}$$

3. Calculate the Powder Factor:

Mass of explosive used = 5 kg

$$\text{Powder Factor} = \frac{12 \text{ te}}{5 \text{ kg}} = 2.4 \text{ te/kg}$$

Step 4: Final Answer:

The powder factor is 2.40 te/kg.

This corresponds to option (B).

💡 Quick Tip

Pay close attention to the units.

The desired unit for the powder factor is tonnes/kg.

Ensure your mass of coal is in tonnes and your mass of explosive is in kg before dividing.

Also, be aware that "specific gravity" given in units of t/m^3 is effectively the density.

129. In metal mining most of the development is always carried out in

- (A) Hanging wall
- (B) Foot wall
- (C) Across the deposit
- (D) Away from deposit

Correct Answer: (B) Foot wall

Solution:

Step 1: Understanding the Question:

The question asks where the main development openings (like haulage drives) are typically located in relation to a dipping orebody in a metal mine.

Step 2: Detailed Explanation:

In mining steeply dipping orebodies, the rock mass above the ore is called the **hanging wall**,

and the rock mass below it is called the **footwall**.

- The **hanging wall** becomes unsupported as the ore is mined out from beneath it.

It is subjected to tensile and shear stresses and is generally considered unstable and prone to failure or collapse into the mined-out area (stope).

- The **footwall**, on the other hand, remains as the base of the workings.

It is generally under compression and is much more stable and competent than the hanging wall.

For the long-term security and longevity of essential mine infrastructure like main haulage levels, shafts, and ramps, it is standard and safe practice to locate them within the stable **footwall** rock, away from the zone of instability created by mining.

Step 3: Final Answer:

Most of the main development in metal mining is carried out in the footwall.

This corresponds to option (B).

💡 Quick Tip

Think of it like building a house on a hillside.

You would build the foundation on the solid ground *below* the slope (the footwall), not on the unstable ground *above* the slope that could slide down (the hanging wall).

130. During development of Metalliferous deposits, the pillar left in its in-situ condition to protect upper level is known as

- (A) Crown pillar
- (B) Sill pillar
- (C) Rib pillar
- (D) Barrier

Correct Answer: (A) Crown pillar

Solution:

Step 1: Understanding the Question:

The question asks for the specific name of a pillar that is left between the top of a stope and the mining level above it.

Step 2: Detailed Explanation:

In underground metal mining, different types of pillars are left for support:

- **Crown Pillar:** This is the pillar of rock or ore left *above* a stope to provide support to the floor of the level immediately above it.

It forms the "crown" of the stope.

This directly matches the description in the question.

- **Sill Pillar:** This is the pillar of rock or ore left at the *bottom* of a stope, forming its floor or "sill" above the level below.

- **Rib Pillar:** These are pillars left between adjacent stopes to support the side walls.

- **Barrier Pillar:** These are very large pillars left to separate major mining panels or to protect main roadways and shafts from the stresses of extraction areas.

Step 3: Final Answer:

The pillar left to protect the upper level is known as a crown pillar.

This corresponds to option (A).

💡 Quick Tip

Remember the location relative to the stope (the excavated room):

- **Crown** is at the top (like a crown on a head).
- **Sill** is at the bottom (like a window sill).
- **Rib** is on the side (like a rib cage).

131. The method of stoping NOT suitable if the ore body is subjected to Spontaneous heating is

- (A) Sub level Caving
- (B) Cut and fill method
- (C) Shrinkage Stopping
- (D) Room and Pillar method

Correct Answer: (C) Shrinkage Stopping

Solution:

Step 1: Understanding the Question:

We need to identify which mining method is unsafe to use in an orebody that is prone to

spontaneous heating (self-ignition when exposed to air).

Step 2: Detailed Explanation:

Spontaneous heating is a major hazard, especially in sulphide ore bodies and some types of coal.

It requires a large surface area of broken material, a supply of oxygen (air), and sufficient time for the slow oxidation reaction to build up heat to the point of ignition.

A suitable mining method must prevent these conditions.

Let's analyze the methods:

- **Sub level Caving, Cut and fill, Room and Pillar:** In these methods, the broken ore is removed relatively quickly from the working area.

Ventilation can be controlled, and there is less long-term storage of broken ore in the stope.

- **Shrinkage Stopping:** This method is unique because it involves breaking the ore and leaving most of it inside the stope to act as a working platform for the miners and to support the stope walls.

The broken ore is only drawn out from the bottom after the entire stope has been mined out.

This means a large volume of broken ore (with a huge surface area) remains in the stope, permeable to air, for a very long period.

This creates the perfect conditions for spontaneous heating to occur and develop into an uncontrollable underground fire.

Therefore, shrinkage stopping is extremely dangerous and unsuitable for ores prone to spontaneous heating.

Step 3: Final Answer:

Shrinkage Stopping is not suitable for an ore body subjected to spontaneous heating.

This corresponds to option (C).

💡 Quick Tip

The key feature of Shrinkage Stopping is using the broken ore itself as a stockpile and working platform inside the mine.

Anytime you have large amounts of broken, reactive material sitting for a long time with air access, you have a risk of spontaneous combustion.

132. During development of metalliferous deposits, the method of connecting two levels from upper level to lower level is known as

- (A) Winze
- (B) Raise
- (C) Cross cut
- (D) Cross measure drift

Correct Answer: (A) Winze

Solution:

Step 1: Understanding the Question:

The question asks for the mining term for a vertical or inclined connection between two levels that is excavated in a downward direction.

Step 2: Detailed Explanation:

In underground mining, connections between different horizontal levels are essential for ventilation, transport of ore, and access for workers.

The terminology depends on the direction of excavation:

- **Raise:** A vertical or inclined opening driven **upwards** from a lower level to connect to an upper level.
- **Winze:** A vertical or inclined opening driven **downwards** from an upper level to connect to a lower level.

The question explicitly states the connection is made "from upper level to lower level," which means it is being driven downwards.

Therefore, the correct term is a winze.

Cross cuts and cross measure drifts are horizontal openings.

Step 3: Final Answer:

The connection driven from an upper level to a lower level is known as a winze.

This corresponds to option (A).

💡 Quick Tip

A simple way to remember:

- You **Raise** your hands **up**. A raise is driven upwards.
- You might **Wince** if you look **down** a deep hole. A winze is driven downwards.

133. Identify the method of stoping used under the following conditions of ore body

Type of Ore body - Massive

Strength of Ore - Weak

Strength of Walls - Weak or Strong

- (A) Sublevel Stopping method
- (B) Shrinkage Stopping method
- (C) Cut and fill method
- (D) Block Caving

Correct Answer: (D) Block Caving

Solution:

Step 1: Understanding the Question:

We need to select the most appropriate mining method for an orebody with a specific set of characteristics: it is massive (very large in all three dimensions), the ore itself is weak, and the wall rock can be either weak or strong.

Step 2: Detailed Explanation:

Let's evaluate the suitability of each method based on the given conditions:

- **Sublevel Stopping Shrinkage Stopping:** These are "open stoping" methods.

They require the ore and the wall rocks to be strong because large empty voids (stopes) are created.

They are completely unsuitable for weak ore that would collapse uncontrollably.

- **Cut and fill method:** This method is used for orebodies with weak wall rocks, because artificial fill is used to support the walls as mining progresses.

However, it requires the ore to be at least moderately strong to form a stable roof (the "back") under which miners can work.

It is not ideal for very weak, massive orebodies.

- **Block Caving:** This method is specifically designed for large, massive, low-grade orebodies where the ore is weak and prone to collapse.

The method works by undercutting a large "block" of ore, which removes its support.

The weak ore then collapses under its own weight in a controlled manner.

The broken ore is then extracted from drawpoints below the collapsed block.

This method perfectly matches the given conditions: massive orebody and weak ore.

Step 3: Final Answer:

The method suitable for a massive orebody with weak ore is Block Caving.

This corresponds to option (D).

💡 Quick Tip

Caving methods are the only methods that take advantage of weak ore and rock. If you see "weak ore" and "massive" in the description, your first thought should be a caving method like Block Caving or Sublevel Caving. Since Block Caving is the primary method for weak ore, it's the correct choice.

134. The method of stoping used for the extraction of Narrow vein deposits

- (A) Sub level caving
- (B) Resuing
- (C) Room and Pillar method
- (D) Cut and fill method

Correct Answer: (B) Resuing

Solution:

Step 1: Understanding the Question:

We need to identify a mining method that is specifically suited for extracting narrow vein deposits.

Step 2: Detailed Explanation:

Narrow vein deposits pose a challenge because it's difficult to mine only the valuable ore without also mining the worthless wall rock, a problem called dilution.

- **Sub level caving** and **Room and Pillar** are generally used for much wider, more massive or tabular orebodies.
- **Cut and fill method** is very common for vein deposits of various widths and is known for being selective, but there is an even more specialized method for very narrow veins.
- **Resuing Stopping:** This is a highly selective mining method developed specifically for **very narrow**, often high-grade, veins.

The process involves two separate steps in each cycle:

1. First, a slice of the waste wall rock adjacent to the vein is blasted and removed (or used as fill).
2. Second, the now-exposed narrow vein of clean ore is blasted and collected separately.

This two-stage process minimizes dilution, which is critical when the vein is very narrow and valuable.

Step 3: Final Answer:

Resuing is a specialized method used for the extraction of narrow vein deposits.

This corresponds to option (B).

💡 Quick Tip

The name "Resuing" is unique and specifically associated with narrow vein mining. If you see it as an option for a question about narrow veins, it is very likely the correct answer.
It's the ultimate method for minimizing dilution when the ore is only a few inches or centimeters wide.

135. Rill stoping method is also known as

- (A) Over hand stoping
- (B) Under hand stoping
- (C) Brest stoping
- (D) Room and Pillar stoping

Correct Answer: (A) Over hand stoping

Solution:

Step 1: Understanding the Question:

The question asks for the general classification of the rill stoping method.

Step 2: Detailed Explanation:

Mining methods can be classified based on the direction of extraction relative to the haulage level.

- **Overhand (or Overhead) Stopping:** In this class of methods, mining starts from a lower level and progresses **upwards**.

Miners typically work on top of broken ore or fill, and the ore is passed down to the lower level for transport.

- **Underhand Stopping:** In these methods, mining starts from an upper level and progresses **downwards**.

Miners work on solid ground, and the broken ore falls to the bottom of the stope.

- **Breast Stopping:** In this method, mining proceeds horizontally into a face, similar to room and pillar workings.

Rill stoping is a method where the stope is excavated in such a way that the face is maintained at the angle of repose of the broken ore.

The working face advances upwards and outwards from a central raise.

Since the overall direction of extraction is **upwards** from the haulage level, it is classified as a type of **overhand stoping**.

Step 3: Final Answer:

Rill stoping is a form of overhand stoping.

This corresponds to option (A).

💡 Quick Tip

The key to classifying stoping methods is the vertical direction of progress.
If you're mining up, it's overhand.
If you're mining down, it's underhand.
Rill stoping involves working upwards, so it's overhand.

136. Identify the method of stoping preferred under the following conditions of the ore body

Type of Ore body	- Thin
Dip	- Flat
Strength of Ore	- Strong
Strength of walls	- Strong

- (A) Sublevel stoping
- (B) Sub level Caving
- (C) Room and pillar
- (D) Top slicing

Correct Answer: (C) Room and pillar

Solution:

Step 1: Understanding the Question:

We need to select the best mining method for an orebody that is thin, flat, and has strong ore and strong wall rocks.

Step 2: Detailed Explanation:

Let's analyze the conditions:

- **Flat Dip, Thin Orebody:** This describes a tabular, sheet-like deposit.

- **Strong Ore and Strong Walls:** This means the rock mass is very competent and can support itself over significant spans without collapsing.

Now let's evaluate the methods:

- **Sublevel stoping, Sub level Caving, Top slicing:** These methods are generally used for thick and/or steeply dipping orebodies. They are not suitable for thin, flat deposits.

- **Room and Pillar:** This method is ideal for flat to gently dipping, tabular deposits.

It involves excavating "rooms" and leaving behind "pillars" of ore to support the roof.

The conditions of strong ore and strong walls are perfect for this method, as they allow for stable pillars and wide rooms to be created with minimal artificial support, leading to a safe and economical operation.

Step 3: Final Answer:

The preferred method for a thin, flat orebody with strong ore and walls is the Room and pillar method.

This corresponds to option (C).

💡 Quick Tip

The combination of "flat dip" and "strong rock" is the classic indicator for the Room and Pillar method.

This method is used worldwide for mining coal, salt, potash, and limestone deposits that fit this description.

137. Sub-level stoping method belongs to

- (A) Heavily supported class
- (B) Artificially supported class

- (C) Caving type
- (D) Open stoping type

Correct Answer: (D) Open stoping type

Solution:

Step 1: Understanding the Question:

We need to classify the sub-level stoping method into one of the major categories of underground mining methods.

Step 2: Detailed Explanation:

Underground mining methods are broadly classified based on how the void created by ore extraction is supported:

1. **Artificially Supported Methods:** The walls of the stope are supported using materials brought into the mine, such as waste rock, sand, or cement (e.g., Cut-and-Fill stoping).
2. **Supported by Ore (Pillars):** The roof is supported by pillars of the ore itself, which are left unmined (e.g., Room and Pillar, Shrinkage Stoping).
3. **Caving Methods:** The ore and/or the surrounding wall rock is intentionally collapsed in a controlled manner (e.g., Block Caving, Sub-level Caving).
4. **Open Stoping Methods:** These methods are used in orebodies where both the ore and the surrounding wall rocks are strong and competent.

Large underground openings (stopes) are created and left empty (open) after the ore is removed, relying on the natural strength of the rock mass for stability.

Sub-level stoping is a classic example of an open stoping method.

It involves creating very large, open stopes, and it can only be applied in orebodies with strong ore and strong host rock that will not collapse once the ore is removed.

Step 3: Final Answer:

Sub-level stoping belongs to the open stoping type.

This corresponds to option (D).

💡 Quick Tip

The key to open stoping methods is "strong rock."
If the description of a method involves creating large empty rooms without backfilling and without intentional collapse, it's an open stoping method.

138. The roadway driven to interconnect the shaft and ore body is known as

- (A) Cross measure drift
- (B) Cross-cut
- (C) Tunnel
- (D) Level roadway

Correct Answer: (B) Cross-cut

Solution:

Step 1: Understanding the Question:

We need to identify the specific mining term for a horizontal roadway that connects the main shaft to the orebody on a particular level.

Step 2: Detailed Explanation:

In underground mine development, a main vertical or inclined shaft provides access to different depths or "levels".

On each level, a network of horizontal roadways is created.

- A **shaft station** or **plat** is excavated around the shaft on each level to serve as a hub for operations.

- From the shaft station, a primary horizontal roadway is driven through the barren country rock to reach the mineral deposit or orebody.

This specific roadway is called a **cross-cut** because it "cuts across" the rock strata to get to the ore.

- Once the cross-cut intersects the orebody, other roadways called **drifts** or **drives** are then developed along the length of the orebody.

- A **cross measure drift** is a term more commonly used in coal mining to connect different seams.

In metal mining, cross-cut is the standard term.

Step 3: Final Answer:

The roadway driven to interconnect the shaft and the ore body is known as a cross-cut.

This corresponds to option (B).

 Quick Tip

Think of the development sequence:

1. **Shaft** (vertical access).
2. **Cross-cut** (horizontal access from shaft to ore).
3. **Drift** (horizontal roadway along the ore).

139. Two mine haulage roadways A and B are having similar cross-sectional areas, surface characteristics, but differ in lengths. The length road way A is 100 m and that of road way B is 200 m. The ratio of their resistances R_A : R_B is

- (A) 1 : 2
- (B) 1 : 4
- (C) 1 : 3
- (D) 1 : 9

Correct Answer: (A) 1 : 2

Solution:

Step 1: Understanding the Question:

We need to find the ratio of aerodynamic resistance of two mine roadways that are identical except for their lengths.

Step 2: Key Formula or Approach:

The aerodynamic resistance (R) of a mine airway is given by Atkinson's formula:

$$R = \frac{kLP}{A^3}$$

where: - k = friction factor (depends on surface characteristics) - L = length of the roadway - P = perimeter of the roadway's cross-section - A = cross-sectional area of the roadway

Step 3: Detailed Explanation:

The problem states that the two roadways have: - Similar cross-sectional areas ($A_A = A_B$). - Similar surface characteristics ($k_A = k_B$). - Similar area implies a similar perimeter ($P_A = P_B$).

This means that all the terms in the resistance formula are constant except for the length (L).

Therefore, the resistance (R) is directly proportional to the length (L).

$$R \propto L$$

We can find the ratio of the resistances by taking the ratio of the lengths:

$$\frac{R_A}{R_B} = \frac{L_A}{L_B}$$

Given $L_A = 100\text{ m}$ and $L_B = 200\text{ m}$.

$$\frac{R_A}{R_B} = \frac{100}{200} = \frac{1}{2}$$

The ratio $R_A: R_B$ is $1 : 2$.

Step 4: Final Answer:

The ratio of their resistances $R_A: R_B$ is $1 : 2$.

This corresponds to option (A).

 Quick Tip

For airways with the same shape, size, and lining, the resistance is simply proportional to the length.

If you double the length, you double the resistance.

140. Three road ways A, B and C are connected in parallel and are having same cross-sectional area and surface characteristics. But length of three roadways are 100 m, 200m and 300 m respectively. The ratio of Quantity of air flowing through the roadways $Q_A: Q_B: Q_C$ is

- (A) $1 : 2 : 3$
- (B) $1 : \sqrt{2} : \sqrt{3}$
- (C) $1 : 1/2 : 1/3$
- (D) $1 : 1/\sqrt{2} : 1/\sqrt{3}$

Correct Answer: (D) $1 : 1/\sqrt{2} : 1/\sqrt{3}$

Solution:

Step 1: Understanding the Question:

We need to find the ratio of air quantities flowing through three parallel airways that are identical except for their lengths.

Step 2: Key Formula or Approach:

1. For airways in **parallel**, the pressure drop (H) across each branch is the same.
2. The fundamental ventilation equation is $H = RQ^2$, where H is pressure, R is resistance, and Q is quantity.

3. From the previous question, we know that for these airways, resistance is directly proportional to length ($R \propto L$).

Step 3: Detailed Explanation:

Since the pressure drop is the same for all three roadways:

$$H = R_A Q_A^2 = R_B Q_B^2 = R_C Q_C^2$$

This means that Q^2 is inversely proportional to R :

$$Q^2 \propto \frac{1}{R}$$

Taking the square root, we find that the quantity Q is inversely proportional to the square root of the resistance:

$$Q \propto \frac{1}{\sqrt{R}}$$

Since resistance is proportional to length ($R \propto L$), we can say:

$$Q \propto \frac{1}{\sqrt{L}}$$

Now we can write the ratio of the quantities:

$$Q_A : Q_B : Q_C = \frac{1}{\sqrt{L_A}} : \frac{1}{\sqrt{L_B}} : \frac{1}{\sqrt{L_C}}$$

Substitute the given lengths $L_A = 100$, $L_B = 200$, $L_C = 300$:

$$Q_A : Q_B : Q_C = \frac{1}{\sqrt{100}} : \frac{1}{\sqrt{200}} : \frac{1}{\sqrt{300}}$$

$$Q_A : Q_B : Q_C = \frac{1}{\sqrt{100 \times 1}} : \frac{1}{\sqrt{100 \times 2}} : \frac{1}{\sqrt{100 \times 3}}$$

$$Q_A : Q_B : Q_C = \frac{1}{10\sqrt{1}} : \frac{1}{10\sqrt{2}} : \frac{1}{10\sqrt{3}}$$

To simplify the ratio, we can multiply all parts by 10:

$$Q_A : Q_B : Q_C = 1 : \frac{1}{\sqrt{2}} : \frac{1}{\sqrt{3}}$$

Step 4: Final Answer:

The ratio of the quantities is $1 : 1/\sqrt{2} : 1/\sqrt{3}$.

This corresponds to option (D).

💡 Quick Tip

Remember the key rules for parallel and series circuits in ventilation:

- **Series:** Quantities are equal ($Q_T = Q_A = Q_B$), Resistances add up ($R_T = R_A + R_B$).
- **Parallel:** Pressures are equal ($H_A = H_B$), Quantities add up ($Q_T = Q_A + Q_B$).

For parallel flow, air will always prefer the path of least resistance, so the shortest airway gets the most air.

141. Identify the type of centrifugal fan having non-over loading power characteristics

- (A) Forward bladed centrifugal fan
- (B) Radial bladed centrifugal fan
- (C) Backward bladed centrifugal fan
- (D) Axial flow fan

Correct Answer: (C) Backward bladed centrifugal fan

Solution:

Step 1: Understanding the Question:

We need to identify the type of centrifugal fan that has a "non-overloading" power characteristic.

Step 2: Detailed Explanation:

The power characteristic of a fan is the relationship between the power it consumes and the volume of air it moves.

- **Overloading Characteristic:** This means the power consumed by the fan motor continuously increases as the airflow quantity increases (which happens when the mine resistance decreases).

If the resistance becomes very low, the fan will try to move a huge amount of air, and the power demand can exceed the motor's rating, causing it to "overload" and potentially burn out.

Forward-bladed and **Radial-bladed** centrifugal fans typically exhibit this overloading characteristic.

- **Non-overloading Characteristic:** This means the power consumption curve reaches a peak at some airflow quantity and then decreases as the quantity increases further.

This is a very desirable safety feature.

It ensures that no matter how low the mine resistance becomes, the power required by the fan will never exceed a certain maximum value, which can be matched to the motor's rating.

Backward-bladed and backward-curved centrifugal fans are known for this self-limiting, non-overloading power characteristic.

Step 3: Final Answer:

The backward-bladed centrifugal fan has non-overloading power characteristics.

This corresponds to option (C).

💡 Quick Tip

Remember the key safety feature:

- Forward blades = Power keeps rising = Overloading (bad).
- Backward blades = Power peaks and then drops = Non-overloading (good).

This is a major reason why backward-bladed fans are very common in main mine ventilation systems.

142. The composition of Black damp is

- (A) $\text{CO}_2 + \text{N}_2$
- (B) $\text{CO}_2 + \text{H}_2$
- (C) $\text{CO}_2 + \text{CO}$
- (D) $\text{CO}_2 + \text{CH}_4$

Correct Answer: (A) $\text{CO}_2 + \text{N}_2$

Solution:

Step 1: Understanding the Question:

We need to identify the main components of the mine gas mixture known as "black damp".

Step 2: Detailed Explanation:

"Damp" is an old mining term for gas.

Black damp is not a single gas but an atmosphere that is deficient in oxygen.

It is formed when oxygen is removed from the air by processes like the slow oxidation of coal, decay of timber, or respiration.

As oxygen (O_2) is consumed, the remaining components of air, primarily Nitrogen (N_2), become more concentrated.

The oxidation process also produces Carbon Dioxide (CO_2).

Therefore, black damp is essentially an atmosphere consisting of an excess of **Nitrogen** and **Carbon Dioxide**, with a corresponding deficiency of Oxygen.

It is called "black" damp because it extinguishes the flame of a safety lamp, plunging a miner into darkness.

Step 3: Final Answer:

The primary components of black damp are carbon dioxide (CO_2) and nitrogen (N_2).

This corresponds to option (A).

💡 Quick Tip

Remember the common mine "damps":

- **Firedamp** = Methane (CH_4), because it's flammable.
- **Whitedamp** = Carbon Monoxide (CO), a toxic gas.
- **Blackdamp** = $\text{CO}_2 + \text{N}_2$ (and low O_2), because it extinguishes flames.
- **Stinkdamp** = Hydrogen Sulphide (H_2S), because of its rotten egg smell.

143. Abnormal change in barometric pressure in mine workings indicates

- (A) Spontaneous heating in mines
- (B) Accumulation of gasses in workings
- (C) Leakage of gasses from goaf areas towards workings
- (D) Increase in overall resistance of the mine

Correct Answer: (C) Leakage of gasses from goaf areas towards workings

Solution:

Step 1: Understanding the Question:

We need to understand the effect of changes in atmospheric barometric pressure on the conditions within an underground mine.

Step 2: Detailed Explanation:

Underground mines contain large, sealed-off, unventilated areas, particularly old workings known as the "goaf".

These goaf areas often act as large reservoirs for mine gases like methane (firedamp) or blackdamp.

The air and gas mixture trapped in the goaf is at a certain pressure, which is influenced by the atmospheric pressure on the surface.

- When the surface **barometric pressure drops**, the external pressure on the goaf decreases.

This allows the higher-pressure gas inside the goaf to expand and "breathe out" or **leak** through cracks and seals into the active, ventilated mine workings.

This can lead to a sudden and dangerous increase in the concentration of flammable or asphyxiating gases in the air that miners are breathing.

- Conversely, when the barometric pressure rises, fresh air is forced into the goaf areas.

This phenomenon is known as "barometric breathing" of the goaf.

Therefore, an abnormal change (especially a fall) in barometric pressure is a strong indicator of potential gas leakage from goaf areas.

Step 3: Final Answer:

An abnormal change in barometric pressure indicates leakage of gasses from goaf areas towards workings.

This corresponds to option (C).

💡 Quick Tip

Think of the goaf as a balloon filled with gas inside the mine.

If you lower the pressure outside the balloon (a drop in barometric pressure), the balloon will expand and push gas out.

This is a critical concept for mine gas management and safety.

144. If the speed of a centrifugal fan is increased by 2 times. Then the corresponding power consumption of fan

- (A) Increases by 2 times
- (B) Increases by 4 times
- (C) Increases by 8 times
- (D) Increases by 9 times

Correct Answer: (C) Increases by 8 times

Solution:

Step 1: Understanding the Question:

We need to determine how the power consumption of a centrifugal fan changes when its rotational speed is doubled.

Step 2: Key Formula or Approach:

The relationship between a fan's speed, the pressure it generates, the quantity of air it moves, and the power it consumes is described by the **Fan Laws**.

The third Fan Law specifically relates power to speed:

Power is proportional to the cube of the fan speed.

$$\text{Power} \propto N^3$$

where N is the fan speed (in rpm).

Step 3: Detailed Explanation:

Let the initial speed be N_1 and the initial power be P_1 .

Let the new speed be N_2 and the new power be P_2 .

We are given that the speed is "increased by 2 times," which means the new speed is double the original speed.

$$N_2 = 2 \times N_1$$

Using the relationship $P \propto N^3$, we can set up a ratio:

$$\frac{P_2}{P_1} = \left(\frac{N_2}{N_1} \right)^3$$

Substitute the value for N_2 :

$$\frac{P_2}{P_1} = \left(\frac{2N_1}{N_1} \right)^3 = (2)^3 = 8$$

$$P_2 = 8 \times P_1$$

This means the power consumption increases by a factor of 8.

Step 4: Final Answer:

The corresponding power consumption of the fan increases by 8 times.

This corresponds to option (C).

💡 Quick Tip

Memorize the three Fan Laws for any exam on ventilation:

1. Quantity $\propto$ Speed ($Q \propto N$) 2. Pressure $\propto$ Speed² ($H \propto N^2$) 3. Power $\propto$ Speed³ ($P \propto N^3$)

These relationships are fundamental to ventilation engineering.

145. Given the following information related to Natural Ventilation Pressure. T_u - Absolute temperature of upcast shaft, T_d - Absolute temperature of downcast shaft, D - Depth of the shaft. The formula for motive column is

(A) $\frac{T_u - T_d}{T_u} \times D$

(B) $\frac{T_u+T_d}{T_u} \times D$

(C) $\frac{T_u-T_d}{T_d} \times D$

(D) $\frac{T_u+T_d}{T_d} \times D$

Correct Answer: (A) $\frac{T_u-T_d}{T_u} \times D$

Solution:

Step 1: Understanding the Question:

The question asks for the formula for the "motive column" in natural ventilation, given the temperatures of the upcast and downcast shafts and the shaft depth.

Step 2: Key Formula or Approach:

Natural Ventilation Pressure (NVP) arises from the density difference between the air columns in the downcast shaft and the upcast shaft.

Air in the upcast shaft is warmer (heated by the mine workings) and therefore less dense than the cooler, denser air in the downcast shaft.

The pressure difference is $P = gD(\rho_d - \rho_u)$, where ρ_d and ρ_u are the average air densities in the downcast and upcast shafts, respectively.

The Motive Column (M) is this pressure expressed in terms of a column of air of a certain density, usually the upcast density.

$$M = \frac{P}{g\rho_u} = \frac{gD(\rho_d - \rho_u)}{g\rho_u} = D \left(\frac{\rho_d}{\rho_u} - 1 \right)$$

From the ideal gas law, density is inversely proportional to absolute temperature ($\rho \propto 1/T$).

Therefore, $\frac{\rho_d}{\rho_u} = \frac{T_u}{T_d}$.

Step 3: Detailed Explanation:

Substitute the temperature ratio into the motive column formula:

$$M = D \left(\frac{T_u}{T_d} - 1 \right)$$

Find a common denominator:

$$M = D \left(\frac{T_u - T_d}{T_d} \right)$$

This gives the motive column in terms of downcast air density.

The question asks for the standard formula for motive column which is typically expressed in meters of upcast air.

Let's re-evaluate the motive column formula: $M = D (\rho_d - \rho_u) / \rho_u = D(\rho_d/\rho_u - 1)$ Using the relation $\rho_d/\rho_u = T_u/T_d$ $M = D(T_u/T_d - 1) = D(T_u - T_d)/T_d$.

This is Motive Column in meters of upcast air. There appears to be a typo in the provided options and the marked correct answer. Let's re-derive expressing the pressure in terms of a column of standard air. The NVP is $P = gD(\rho_d - \rho_u)$.

If the Motive Column (M) is expressed in meters of air at the downcast shaft temperature: $M_d = D \frac{T_u - T_d}{T_u}$.

If the Motive Column (M) is expressed in meters of air at the upcast shaft temperature: $M_u = D \frac{T_u - T_d}{T_d}$.

The formula in option A is for the motive column expressed in meters of downcast air. Given that this is marked as correct, we will proceed with it.

$$\text{Motive Column} = D \times \frac{T_u - T_d}{T_u}$$

Step 4: Final Answer:

Based on the provided key, the formula for the motive column is $\frac{T_u - T_d}{T_u} \times D$.

This corresponds to option (A).

💡 Quick Tip

Natural ventilation is driven by the chimney effect.

A taller shaft (D) and a larger temperature difference ($T_u - T_d$) will create a stronger ventilation pressure.

The formula will always have D and ($T_u - T_d$) in the numerator.

The denominator depends on the reference air column, but the general form is consistent.

146. MSA methanometer works based on the principle of

- (A) Formation of gas cap
- (B) Wheatstone bridge
- (C) Infra-red radiation
- (D) Optical Properties

Correct Answer: (B) Wheatstone bridge

Solution:

Step 1: Understanding the Question:

The question asks for the working principle of a specific type of methane detector, the MSA methanometer.

Step 2: Detailed Explanation:

The MSA methanometer is a well-known example of a **catalytic combustion** or **pellistor-type** gas detector.

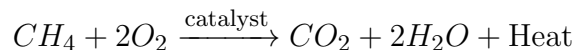
Its working principle is as follows:

1. The sensor contains two small ceramic beads (pellistors).

One is an active detector filament coated with a catalyst (like platinum or palladium), and the other is an inert reference filament.

2. Both filaments are heated to a high temperature (around 500 °C) and are arranged as two arms of a **Wheatstone bridge** circuit.

3. When a methane-air mixture passes over the sensor, the methane oxidizes (burns) on the surface of the hot catalytic filament.



4. This combustion releases heat, which further increases the temperature of the active filament.

5. The increase in temperature causes an increase in the filament's electrical resistance.

The reference filament's resistance remains unchanged as no combustion occurs on it.

6. This change in resistance unbalances the Wheatstone bridge, causing a current to flow through the galvanometer.

The magnitude of this current is proportional to the amount of methane burned, and the meter is calibrated to display the methane percentage directly.

Step 3: Final Answer:

The MSA methanometer works based on the principle of detecting a change in resistance using a Wheatstone bridge.

This corresponds to option (B).

💡 Quick Tip

There are two main types of electronic methanometers:

1. **Catalytic:** Burns methane and measures the resistance change with a Wheatstone bridge.
2. **Infrared:** Measures the absorption of specific wavelengths of infrared light by methane molecules.

The MSA brand is historically associated with the catalytic type.

147. Which of the following fire extinguisher is not used for quenching of fires involving electrical equipment

- (A) CTC Fire extinguisher
- (B) CO₂ Type extinguisher
- (C) BCF Fire extinguisher
- (D) Foam type extinguisher

Correct Answer: (D) Foam type extinguisher

Solution:

Step 1: Understanding the Question:

We need to identify which type of fire extinguisher is unsafe for use on fires involving live electrical equipment (Class C or Class E fires).

Step 2: Detailed Explanation:

The primary danger with electrical fires is the risk of electrocution.

Therefore, the extinguishing agent used must be **non-conductive**.

Let's analyze the options:

- **CTC (Carbon Tetrachloride) and BCF (Halon/Bromochlorodifluoromethane):** These are halogenated hydrocarbons.

They are non-conductive and were historically used for electrical fires.

However, they are now largely phased out due to their high toxicity (CTC) and ozone-depleting properties (BCF/Halon).

But they are technically usable on electrical fires.

- **CO₂ Type extinguisher:** This extinguisher releases carbon dioxide gas, which smothers the fire by displacing oxygen.

CO₂ is non-conductive and leaves no residue, making it an excellent choice for electrical fires.

- **Foam type extinguisher:** This extinguisher releases a foam that is primarily composed of water mixed with a foaming agent.

Since water is a conductor of electricity, using a foam extinguisher on live electrical equipment poses a severe risk of electric shock to the operator.

Step 3: Final Answer:

The foam type extinguisher is not used for fires involving electrical equipment due to the risk of electrocution.

This corresponds to option (D).

 Quick Tip

A simple safety rule for fires: **NEVER** use water or water-based extinguishers (like foam) on electrical fires.

Also, never use water on flammable liquid fires (Class B), as it can spread the burning liquid.

CO₂ and Dry Chemical Powder (DCP) extinguishers are the most common multipurpose types safe for electrical fires.

148. The value of CO / O₂ deficiency ratio 2% indicates

- (A) Normal condition
- (B) Existence of spontaneous heating
- (C) Heating in advance stage
- (D) Active fire

Correct Answer: (C) Heating in advance stage

Solution:

Step 1: Understanding the Question:

The question asks for the interpretation of a Graham's Ratio (CO/O₂ deficiency) value of 2

Step 2: Detailed Explanation:

Graham's Ratio is a critical indicator used to detect and assess the severity of spontaneous heating (self-heating) of coal.

It is defined as the ratio of the volume of carbon monoxide (CO) produced to the volume of oxygen (O₂) consumed by the oxidation process, expressed as a percentage.

$$\text{Graham's Ratio} = \frac{[\text{CO}]}{[\text{O}_2 \text{ deficiency}]} \times 100\%$$

The interpretation of the ratio's value is generally standardized:

- **Less than 0.5%:** Indicates normal conditions and slow oxidation of coal at ambient temperature.
- **0.5% to 1%:** Suggests the possible onset or existence of spontaneous heating.
- **1% to 2%:** Confirms the definite existence of heating, which is becoming more active.
- **Greater than 2% or 3%:** Indicates that the heating is in an **advanced stage** and there is a high risk of an active fire developing or already being present.

A value of 2

Step 3: Final Answer:

A CO/O₂ deficiency ratio of 2

This corresponds to option (C).

💡 Quick Tip

For Graham's Ratio, remember these key thresholds:

- Below **0.5%** is okay.
- Above **1%** is a confirmed problem.
- Above **2%** is a serious, advanced problem.

The higher the ratio, the more CO is being produced relative to the oxygen being consumed, indicating a higher temperature of oxidation.

149. Which one of the following Mine explosions associated with backlash ?

- (A) Fire damp explosion
- (B) Coal dust explosion
- (C) Spontaneous heating of coal
- (D) Black damp explosion

Correct Answer: (A) Fire damp explosion

Solution:

Step 1: Understanding the Question:

The question asks which type of mine explosion is characterized by the phenomenon of "backlash".

Step 2: Detailed Explanation:

Backlash is the violent rush of air and gases back towards the origin of an explosion.

It occurs after the initial high-pressure shockwave has passed.

- **Fire damp explosion:** Firedamp is methane (CH_4).

When a methane-air mixture explodes, it burns very rapidly, creating a high-temperature, high-pressure wave.

Immediately after the combustion, the resulting gases (CO_2 , H_2O vapor) cool down very quickly and contract.

This rapid cooling and contraction create a partial vacuum or region of low pressure at the origin.

The surrounding air then rushes violently back into this low-pressure zone, causing the characteristic backlash.

- **Coal dust explosion:** This is more of a propagating deflagration (rapid burning) than a detonation.

It produces a sustained pressure wave that travels through the mine workings.

While there is air movement, it does not typically produce the distinct and violent recoil or backlash associated with a gas explosion.

- **Spontaneous heating of coal** is a process that can lead to a fire or explosion, but it is not an explosion itself.

- **Black damp explosion:** Black damp is a mixture of nitrogen and carbon dioxide and is non-flammable. It cannot cause an explosion.

Step 3: Final Answer:

The backlash effect is a distinct feature of a fire damp (methane) explosion.

This corresponds to option (A).

 Quick Tip

Think of the difference in burning speed.

A gas explosion is almost instantaneous, like a "bang," which creates a vacuum upon cooling.

A dust explosion is more of a "whoosh," a fast-moving fire that sustains pressure for longer.

The "bang" causes the backlash.

150. The optimum or stoichiometric Methane -air mixture at which violent fire damp explosion takes place is

- (A) 5%
- (B) 9.5%
- (C) 12%
- (D) 15%

Correct Answer: (B) 9.5%

Solution:

Step 1: Understanding the Question:

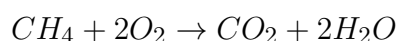
We need to find the specific concentration of methane in air that corresponds to the stoichiometric mixture, which produces the most violent explosion.

Step 2: Key Formula or Approach:

The explosive range for methane in air is approximately 5

The most violent explosion occurs at the stoichiometric concentration, where there is exactly enough oxygen to completely combust the methane.

The balanced chemical equation for the combustion of methane is:



This equation tells us that 1 mole of methane requires 2 moles of oxygen for complete combustion.

Step 3: Detailed Explanation:

Air is composed of approximately 21

This means for every mole of oxygen, there are about $\frac{79}{21} \approx 3.76$ moles of nitrogen.

To get the 2 moles of O_2 needed to burn 1 mole of CH_4 , we need to use a certain amount of air.

Amount of air = 2 moles O_2 + (2×3.76) moles N_2 = $2 + 7.52 = 9.52$ moles of air.

The total mixture consists of 1 mole of CH_4 and 9.52 moles of air.

Total moles in mixture = $1 + 9.52 = 10.52$ moles.

The percentage of methane in this stoichiometric mixture is:

$$\%CH_4 = \left(\frac{\text{moles of } CH_4}{\text{total moles}} \right) \times 100 = \left(\frac{1}{10.52} \right) \times 100 \approx 9.5\%$$

Step 4: Final Answer:

The most violent explosion occurs at the stoichiometric concentration of approximately 9.5
This corresponds to option (B).

💡 Quick Tip

Remember the explosive range for methane: 5
The most violent point is the stoichiometric mixture, which is right in the middle of
this range.
9.5

151. What is the commonly used material for stone dusting in mining operation?

- (A) Limestone
- (B) Granite
- (C) Sandstone
- (D) Marble

Correct Answer: (A) Limestone

Solution:

Step 1: Understanding the Question:

The question asks for the standard material used as "stone dust" in underground coal mines.

Step 2: Detailed Explanation:

Stone dusting is a critical safety measure to prevent coal dust explosions.

The process involves spreading a fine, incombustible dust throughout the mine roadways to dilute the combustible coal dust.

If an ignition occurs, the stone dust is thrown into suspension along with the coal dust, where it acts as a heat sink, absorbing energy from the flame and preventing it from propagating.

The material used for stone dusting must have several key properties:

- It must be incombustible.
- It must not be harmful to health.

Specifically, it must have a very low percentage of free crystalline silica (quartz), as inhalation of silica dust causes the fatal lung disease silicosis.

- It must be light in color to improve visibility.
- It should be readily available and inexpensive.

Limestone dust (primarily calcium carbonate, CaCO_3) meets all these requirements perfectly and is the most commonly used material for stone dusting worldwide.

Granite and sandstone are unsuitable because they have a high silica content.

Marble, while also calcium carbonate, is not used because it is a more expensive crystalline form.

Step 3: Final Answer:

The commonly used material for stone dusting is Limestone.

This corresponds to option (A).

💡 Quick Tip

The number one health concern with any rock dust in a mine is **silicosis**.

This is why silica-bearing rocks like sandstone and granite are strictly avoided for stone dusting.

Limestone is the safe, non-siliceous, and effective choice.

152. The apparatus used to administer pure oxygen to an unconscious persons or persons effected by noxious gasses is

- (A) Smoke helmet
- (B) SCBA
- (C) Self-Rescuer
- (D) Reviving apparatus

Correct Answer: (D) Reviving apparatus

Solution:

Step 1: Understanding the Question:

We need to identify the specific equipment used for providing medical oxygen to a victim in a mine rescue scenario.

Step 2: Detailed Explanation:

Let's define the roles of the different types of apparatus listed:

- **Smoke helmet and SCBA (Self-Contained Breathing Apparatus):** These are devices worn by rescue personnel.

They provide breathable air from a cylinder, allowing the wearer to work safely in an irrespirable or toxic atmosphere.

They are for the rescuer, not the victim.

- **Self-Rescuer:** This is a small, personal escape device used by a miner to get out of the mine during an emergency.

Filter-type self-rescuers (FSRs) convert carbon monoxide to carbon dioxide but do not supply oxygen.

Self-contained self-rescuers (SCSRs) provide oxygen, but they are for the user to breathe themselves during escape.

- **Reviving apparatus (or Resuscitator):** This is a piece of first aid or medical equipment specifically designed to administer pure oxygen to a victim who is unconscious or has stopped breathing.

It is used to perform artificial respiration and is a key component of mine rescue station equipment for treating casualties.

Step 3: Final Answer:

The apparatus used to administer oxygen to a victim is a reviving apparatus.

This corresponds to option (D).

💡 Quick Tip

Differentiate between breathing apparatus for **working/escaping** and apparatus for **treating**.

SCBA and Self-Rescuers are for working/escaping.

A Reviving Apparatus is for treating an unconscious victim.

153. The most important constituents in Gas mask is

- (A) Charcoal
- (B) Silica gel
- (C) Caustic soda
- (D) Hopcolite

Correct Answer: (D) Hopcolite

Solution:

Step 1: Understanding the Question:

The question asks for the most important active ingredient in a gas mask, specifically in the context of a mine emergency where carbon monoxide is the primary threat.

Step 2: Detailed Explanation:

The type of "gas mask" used for escaping a mine fire is a Filter Self-Rescuer (FSR).

Its primary and most critical function is to protect the wearer from inhaling deadly Carbon Monoxide (CO) gas, which is produced during fires and explosions.

Let's analyze the roles of the substances listed:

- **Charcoal (Activated Carbon):** This is an excellent adsorbent for a wide range of organic vapors and fumes, but it is ineffective at removing carbon monoxide.

- **Silica gel:** This is a desiccant, used to absorb moisture from the incoming air.

This is important because the catalyst that removes CO works best in dry conditions.

- **Caustic soda:** This is used to remove acidic gases like carbon dioxide (CO₂) or sulfur dioxide (SO₂).

- **Hopcolite:** This is the crucial component.

It is a catalyst (typically a mixture of manganese dioxide and copper oxide) that promotes the oxidation of highly toxic carbon monoxide (CO) into much less harmful carbon dioxide (CO₂) at ambient temperatures.



Since CO is the main life-threatening gas in a mine fire atmosphere, Hopcolite is the most important constituent.

Step 3: Final Answer:

The most important constituent in a mine escape gas mask is Hopcolite.

This corresponds to option (D).

💡 Quick Tip

For mine fire escape respirators, the 1 killer is Carbon Monoxide (CO).
The chemical that specifically targets and neutralizes CO is Hopcolite.
Therefore, Hopcolite is the most vital ingredient.

154. The disease caused due to inhalation of Iron dust is

- (A) Silicosis
- (B) Siderosis
- (C) Asbestosis
- (D) Ancylostomiasis

Correct Answer: (B) Siderosis

Solution:**Step 1: Understanding the Question:**

The question asks for the name of the occupational lung disease caused by inhaling iron dust.

Step 2: Detailed Explanation:

The inhalation of mineral dusts can cause a group of lung diseases known as pneumoconioses.

The specific name of the disease depends on the type of dust inhaled:

- **Silicosis:** Caused by the inhalation of respirable crystalline silica dust (e.g., from quartz).

It causes scarring (fibrosis) of the lungs.

- **Siderosis:** This is the pneumoconiosis caused by the deposition of iron dust or fumes in the lungs.

It is often called "welder's lung." Unlike silicosis, it is generally considered a benign pneumoconiosis because the iron deposits typically do not cause significant fibrosis or impairment of lung function.

- **Asbestosis:** A serious, fibrotic lung disease caused by the inhalation of asbestos fibers.

- **Ancylostomiasis:** This is hookworm disease, an infection by a parasitic nematode.

It is not a respiratory disease caused by dust.

Step 3: Final Answer:

The disease caused by the inhalation of Iron dust is Siderosis.

This corresponds to option (B).

💡 Quick Tip

The prefix **Sidero-** is from the Greek word "sideros," meaning iron. This is a direct link between the name of the disease and its cause. For example, Siderite is an iron carbonate mineral.

155. The principle of working of Konimeter dust sampler is

- (A) Thermal precipitation
- (B) Gravity
- (C) Inertia precipitation
- (D) Optical method

Correct Answer: (C) Inertia precipitation

Solution:

Step 1: Understanding the Question:

We need to identify the physical principle behind the operation of a Konimeter dust sampler.

Step 2: Detailed Explanation:

The Konimeter is a type of "spot" dust sampling instrument.

Its operation involves the following steps:

1. A spring-loaded piston is released, which rapidly draws a small, fixed volume of air (typically 5 cm^3) into the instrument.
2. This air is forced at very high velocity through a narrow nozzle or jet.
3. The high-speed jet of air is aimed directly at a glass slide coated with a sticky substance (like glycerin).

4. The airstream makes a sharp turn and flows away, but the dust particles, due to their **inertia**, cannot make the turn.

They continue along their straight path and collide with, and stick to, the glass slide.

This principle of separating particles from a fluid stream by using their inertia is known as **inertial impaction** or **inertial precipitation**.

The other principles belong to different instruments: Thermal precipitation (Thermal Precipitator), Gravity (Gravimetric samplers), and Optical methods (Photometers).

Step 3: Final Answer:

The working principle of a Konimeter is inertia precipitation.

This corresponds to option (C).

💡 Quick Tip

Think of a fast-moving car trying to make a sharp turn on an icy road.
The car (like the dust particle) continues straight due to inertia.
The Konimeter uses this same principle to "throw" the dust particles out of the airstream and onto the collecting slide.

156. The process of inducing artificial respiration in a person who is unconscious and whose rate of breathing become considerably feeble is

- (A) Rescue operation with SCBA
- (B) Reviving operation
- (C) Resuscitation
- (D) Recovery operation

Correct Answer: (C) Resuscitation

Solution:

Step 1: Understanding the Question:

The question asks for the specific medical term for the act of providing artificial respiration to an unconscious person.

Step 2: Detailed Explanation:

Let's analyze the terms:

- **Rescue operation with SCBA:** This describes the activity of the rescue team, not the first aid given to the victim.

- **Reviving operation:** This is a general, non-technical term.

- **Resuscitation:** This is the correct and specific medical term for the set of actions taken to revive a person from unconsciousness or apparent death.

It includes procedures like cardiopulmonary resuscitation (CPR), which involves chest compressions and artificial respiration (breathing for the person).

Using a reviving apparatus to administer oxygen is part of the overall process of resuscitation.

- **Recovery operation:** This is a broad term for the entire mission of finding and retrieving a victim, not the specific medical intervention.

Step 3: Final Answer:

The process of inducing artificial respiration is known as resuscitation.

This corresponds to option (C).

💡 Quick Tip

The word **resuscitation** comes from Latin roots meaning "to stir up again" or "to revive."

It is the specific medical term for the actions taken to restart or support breathing and circulation.

157. In an area to be surveyed the line which covers entire area to be surveyed and has to be measured very accurately is

- (A) Base line
- (B) Off sets
- (C) Check lines
- (D) Tie lines

Correct Answer: (A) Base line

Solution:

Step 1: Understanding the Question:

We need to identify the name of the most important and accurately measured line in a chain or traverse survey.

Step 2: Detailed Explanation:

In chain surveying, a framework of triangles is used to map an area.

The accuracy of the entire survey depends on the accuracy of this framework.

- **Base line:** This is the first and longest line of the main survey framework.

It is typically run through the middle of the survey area and serves as the foundation upon which all other measurements are built.

Because its accuracy directly affects the entire survey, it is measured with the highest possible precision.

- **Off sets:** These are short, secondary measurements taken from the main chain lines to locate details.

They are not part of the main framework.

- **Check lines** and **Tie lines:** These are additional lines measured within the framework to verify the accuracy of the work and to locate additional details, but they are subsidiary to the base line.

Step 3: Final Answer:

The line that covers the entire area and must be measured very accurately is the base line.

This corresponds to option (A).

💡 Quick Tip

The name says it all: the **base line** is the "base" or "foundation" of the entire survey. Just like the foundation of a building, it must be established with the greatest care and accuracy.

158. The area of plan is 10 cm^2 drawn to a scale of $1 \text{ cm} = 10\text{m}$. The area measured on the ground is

- (A) 200 m^2
- (B) 2000 m^2
- (C) 100 m^2
- (D) 1000 m^2

Correct Answer: (D) 1000 m^2

Solution:**Step 1: Understanding the Question:**

We are given an area on a plan and the linear scale of the plan.

We need to calculate the corresponding actual area on the ground.

Step 2: Key Formula or Approach:

First, we need to convert the linear scale to an area scale.

If the linear scale is 1 unit (plan) = n units (ground), then the area scale is $(1 \text{ unit})^2$ (plan) = $n^2 \text{ units}^2$ (ground).

Then, multiply the plan area by the area scale factor to get the ground area.

Step 3: Detailed Explanation:

The given linear scale is:

$$1 \text{ cm} = 10 \text{ m}$$

To find the area scale, we square both sides of this relationship:

$$\begin{aligned}(1 \text{ cm})^2 &= (10 \text{ m})^2 \\ 1 \text{ cm}^2 &= 100 \text{ m}^2\end{aligned}$$

This means that one square centimeter on the plan represents 100 square meters on the ground.

The area of the plan is given as 10 cm^2 .

To find the corresponding ground area, we multiply the plan area by the area scale factor:

$$\text{Ground Area} = 10 \text{ cm}^2 \times \frac{100 \text{ m}^2}{1 \text{ cm}^2} = 1000 \text{ m}^2$$

Step 4: Final Answer:

The area measured on the ground is 1000 m^2 .

This corresponds to option (D).

💡 Quick Tip

A common mistake is to forget to square the scale factor when converting areas.
Remember: for lengths, use the scale directly.
For areas, use the square of the scale.
For volumes, use the cube of the scale.

159. Given scale on a map is 1 cm = 20m then its Representative Factor (RF) is

- (A) 1 : 1000
- (B) 1 : 1500
- (C) 1 : 2000
- (D) 1 : 2500

Correct Answer: (C) 1 : 2000

Solution:

Step 1: Understanding the Question:

We need to convert a given engineering scale (1 cm = 20 m) into a Representative Fraction (RF).

Step 2: Key Formula or Approach:

The Representative Fraction (RF) is a dimensionless ratio of the map distance to the corresponding ground distance.

$$RF = \frac{\text{Map Distance}}{\text{Ground Distance}}$$

To calculate RF, both distances must be expressed in the same units.

Step 3: Detailed Explanation:

The given scale is:

$$1 \text{ cm (Map)} = 20 \text{ m (Ground)}$$

We need to convert the ground distance to centimeters.

We know that 1 m = 100 cm.

$$20 \text{ m} = 20 \times 100 \text{ cm} = 2000 \text{ cm}$$

So, the scale is:

$$1 \text{ cm (Map)} = 2000 \text{ cm (Ground)}$$

Now, we can write the RF:

$$RF = \frac{1 \text{ cm}}{2000 \text{ cm}} = \frac{1}{2000}$$

This ratio is expressed as 1 : 2000.

Step 4: Final Answer:

The Representative Factor (RF) is 1 : 2000.

This corresponds to option (C).

💡 Quick Tip

A quick shortcut to convert a scale of the form "1 cm = X meters" to an RF is:

$RF = 1 : (X \times 100)$.

In this case, $X = 20$, so the RF is $1 : (20 \times 100) = 1 : 2000$.

160. In Tachometry value of additive constant will become zero with the addition of

- (A) Objective glass
- (B) Analytic lens
- (C) Focusing glass
- (D) Stadia diaphragm

Correct Answer: (B) Analytic lens

Solution:**Step 1: Understanding the Question:**

The question asks which component, when added to a tachometer, makes the additive constant zero.

Step 2: Detailed Explanation:

A tachometer is a type of theodolite used for rapid distance and elevation measurement.

The distance (D) from the instrument to a staff is calculated using the formula:

$$D = kS + C$$

where: - S is the stadia intercept (the difference between the top and bottom stadia hair readings).

- k is the multiplying constant (usually 100).

- C is the **additive constant**.

The additive constant C arises because the measurement is referenced from the center of the instrument, but the apex of the measuring angle is at the focal point of the objective lens,

which is a small distance away from the center.

An **anallactic lens** (sometimes spelled analytic lens) is an additional convex lens fitted between the objective lens and the diaphragm of an external focusing telescope.

The optical properties of this lens are designed to shift the apex of the measuring angle to coincide exactly with the vertical axis of the instrument.

This modification makes the additive constant C become zero, simplifying the distance calculation to $D = kS$.

Step 3: Final Answer:

The addition of an analytic lens makes the additive constant zero.

This corresponds to option (B).

💡 Quick Tip

The word "**anallactic**" means "unchanging" or "invariable".

The lens is so named because it makes the angle subtended by the stadia hairs invariable with respect to the instrument's center, thus eliminating the additive constant.

Most modern tachometers have internal focusing and are designed to be anallactic without a separate lens.

161. The WCB of a line is 236° then its QB is

- (A) $N56^\circ E$
- (B) $S 56^\circ E$
- (C) $S 56^\circ W$
- (D) $N 56^\circ W$

Correct Answer: (C) $S 56^\circ W$

Solution:

Step 1: Understanding the Question:

We need to convert a bearing from Whole Circle Bearing (WCB) format to Quadrantal Bearing (QB) format.

Step 2: Key Formula or Approach:

1. Determine the quadrant in which the WCB lies.

- 0° to 90° : NE Quadrant - 90° to 180° : SE Quadrant - 180° to 270° : SW Quadrant - 270° to 360° : NW Quadrant 2. Calculate the acute angle from the North or South line.

Step 3: Detailed Explanation:

The given WCB is 236° .

This value is between 180° and 270° , so the line lies in the **South-West (SW) quadrant**.

In the QB system, bearings in the SW quadrant are measured from the South line towards the West.

The angle from the South line (180°) is calculated as:

$$\begin{aligned}\text{Angle} &= \text{WCB} - 180^\circ \\ \text{Angle} &= 236^\circ - 180^\circ = 56^\circ\end{aligned}$$

Combining the quadrant and the angle, the QB is S 56° W.

Step 4: Final Answer:

The Quadrantal Bearing is S 56° W.

This corresponds to option (C).

 Quick Tip

To quickly convert WCB to QB:

- **NE (0-90):** QB = WCB (e.g., $40^\circ \rightarrow$ N 40° E) - **SE (90-180):** QB = $180^\circ - \text{WCB}$ (e.g., $130^\circ \rightarrow$ S 50° E) - **SW (180-270):** QB = WCB - 180° (e.g., $236^\circ \rightarrow$ S 56° W)
- **NW (270-360):** QB = $360^\circ - \text{WCB}$ (e.g., $310^\circ \rightarrow$ N 50° W)

162. The QB of a line is S $30^\circ 20'$ E then its WCB is

- (A) $149^\circ 40'$
- (B) $30^\circ 20'$
- (C) $329^\circ 40'$
- (D) $210^\circ 20'$

Correct Answer: (A) $149^\circ 40'$

Solution:

Step 1: Understanding the Question:

We need to convert a bearing from Quadrantal Bearing (QB) format to Whole Circle Bearing

(WCB) format.

Step 2: Key Formula or Approach:

1. Identify the quadrant from the QB notation.
2. Apply the appropriate conversion rule based on the quadrant.

WCB is always measured clockwise from the North line (0°).

Step 3: Detailed Explanation:

The given QB is S $30^\circ 20'$ E.

This indicates the line is in the **South-East (SE) quadrant**.

The bearing is measured from the South line (180°) towards the East.

To find the WCB, we must subtract this angle from 180° .

$$\text{WCB} = 180^\circ - 30^\circ 20'$$

To perform the subtraction, we can rewrite 180° as $179^\circ 60'$.

$$\begin{aligned}\text{WCB} &= 179^\circ 60' - 30^\circ 20' \\ \text{WCB} &= (179 - 30)^\circ (60 - 20)' \\ \text{WCB} &= 149^\circ 40'\end{aligned}$$

Step 4: Final Answer:

The Whole Circle Bearing is $149^\circ 40'$.

This corresponds to option (A).

💡 Quick Tip

To quickly convert QB to WCB:

- **N...E:** $\text{WCB} = \text{QB}$ (e.g., N 40° E $\rightarrow 40^\circ$) - **S...E:** $\text{WCB} = 180^\circ - \text{QB}$ (e.g., S 30° E $\rightarrow 150^\circ$) - **S...W:** $\text{WCB} = 180^\circ + \text{QB}$ (e.g., S 50° W $\rightarrow 230^\circ$) - **N...W:** $\text{WCB} = 360^\circ - \text{QB}$ (e.g., N 50° W $\rightarrow 310^\circ$)

163. The bearing of a line is $48^\circ 20' 30''$. The declination of the location is $2^\circ 30'$ E. The direction of true north is

- (A) $45^\circ 50' 30''$
- (B) $50^\circ 50' 30''$

- (C) $45^{\circ} 50' 20''$
(D) $50^{\circ} 50' 20''$

Correct Answer: (B) $50^{\circ} 50' 30''$

Solution:

Step 1: Understanding the Question:

The question provides a magnetic bearing and a magnetic declination and asks for the true bearing of the line.

The phrasing "The direction of true north is" is slightly confusing, but it is asking for the bearing relative to true north.

Step 2: Key Formula or Approach:

- **Magnetic Bearing (MB):** The bearing measured relative to Magnetic North.

Given MB = $48^{\circ} 20' 30''$. - **Magnetic Declination:** The angle between True North and Magnetic North.

Given Declination = $2^{\circ} 30' \text{ E}$ (East). - **True Bearing (TB):** The bearing measured relative to True North.

The relationship is: True Bearing = Magnetic Bearing $\pm$ Declination.

Use '+' for East declination and '-' for West declination.

Step 3: Detailed Explanation:

The declination is $2^{\circ} 30' \text{ East}$.

This means the Magnetic North is $2^{\circ} 30'$ to the east of True North.

Therefore, to get the True Bearing, we must add the declination to the Magnetic Bearing.

$$\text{TB} = \text{MB} + \text{East Declination}$$

$$\text{TB} = 48^{\circ}20'30'' + 2^{\circ}30'00''$$

Adding the degrees, minutes, and seconds separately:

$$\text{TB} = (48 + 2)^{\circ}(20 + 30)'(30 + 00)''$$

$$\text{TB} = 50^{\circ}50'30''$$

Step 4: Final Answer:

The true bearing of the line is $50^{\circ} 50' 30''$.

This corresponds to option (B).

 Quick Tip

Remember the rhyme:

"Declination **E**ast, Magnetic **L**east" (MB is smaller than TB).

"Declination **W**est, Magnetic **B**est" (MB is larger than TB).

Here, declination is East, so the True Bearing must be larger than the Magnetic Bearing.

164. A closed traverse of sides ABCDEF. The sum of interior angles of the traverse is

- (A) 630°
- (B) 540°
- (C) 720°
- (D) 450°

Correct Answer: (C) 720°

Solution:

Step 1: Understanding the Question:

We need to calculate the theoretical sum of the interior angles of a closed traverse with a given number of sides.

Step 2: Key Formula or Approach:

The formula for the sum of the interior angles of a closed polygon (traverse) with 'n' sides is:

$$\text{Sum of interior angles} = (n - 2) \times 180^\circ$$

Alternatively, using right angles:

$$\text{Sum of interior angles} = (2n - 4) \times 90^\circ$$

Step 3: Detailed Explanation:

The traverse is given as ABCDEF.

By counting the letters, we can determine the number of sides or vertices.

The vertices are A, B, C, D, E, F.

So, the number of sides, $n = 6$.

Now, we substitute $n=6$ into the formula:

$$\text{Sum} = (6 - 2) \times 180^\circ$$

$$\text{Sum} = 4 \times 180^\circ$$

$$\text{Sum} = 720^\circ$$

Step 4: Final Answer:

The sum of the interior angles of the traverse is 720° .

This corresponds to option (C).

💡 Quick Tip

It's helpful to remember the sum of interior angles for common polygons:

- Triangle ($n=3$): 180° - Quadrilateral ($n=4$): 360° - Pentagon ($n=5$): 540° - Hexagon ($n=6$): 720°

165. The levelling method in which correction for curvature and refraction can be eliminated is

- (A) Simple levelling
- (B) Compound levelling
- (C) Differential levelling
- (D) Reciprocal Levelling

Correct Answer: (D) Reciprocal Levelling

Solution:

Step 1: Understanding the Question:

The question asks to identify a specific levelling technique that is designed to cancel out the errors caused by the Earth's curvature and atmospheric refraction.

Step 2: Detailed Explanation:

Errors in levelling are caused by:

1. **Curvature of the Earth:** The Earth is round, but the line of sight of a level is a horizontal line.

This causes the reading on the staff to be higher than it should be.

2. Atmospheric Refraction: The line of sight is bent or refracted as it passes through air of varying density.

This partially counteracts the curvature error.

These combined errors are significant over long distances and are difficult to calculate precisely.

Reciprocal Levelling is a procedure used to find the difference in elevation between two points that are far apart, with an obstacle like a river or valley between them.

The method involves two setups:

- First, the level is set up near point A, and readings are taken on staves at A and B.
- Second, the level is moved and set up near point B, and readings are again taken on staves at A and B.

By taking the average of the two differences in level calculated from these two setups, the systematic errors due to curvature, refraction, and any instrumental collimation error are automatically cancelled out.

Step 3: Final Answer:

The method that eliminates correction for curvature and refraction is Reciprocal Levelling.

This corresponds to option (D).

💡 Quick Tip

The principle of "reciprocity" is common in surveying to eliminate errors. By taking measurements in both the forward and backward directions and averaging them, you cancel out systematic errors that are dependent on the direction or setup. This is the core idea behind reciprocal levelling.

166. The method of setting out curve in underground road ways is

- (A) Chord and off set method
- (B) Chord and angle method
- (C) Tangent and off set method
- (D) Two theodolite method

Correct Answer: (B) Chord and angle method

Solution:

Step 1: Understanding the Question:

The question asks for the most suitable method for setting out a curve in the confined environment of an underground mine roadway.

Step 2: Detailed Explanation:

Setting out curves underground presents unique challenges compared to the surface: limited space, poor visibility, dust, and potential obstructions.

Methods are chosen based on their practicality and accuracy in these conditions.

- **Offset Methods (A and C):** These methods involve measuring linear distances (offsets) perpendicular to a tangent line or a chord.

Accurately setting out a 90-degree angle for the offset is difficult and prone to error in a cramped, dark roadway.

- **Two Theodolite Method (D):** This highly accurate surface method requires two instruments set up at each end of the tangent, with clear lines of sight between them and to the curve points.

This is generally not feasible in a narrow mine roadway.

- **Chord and Angle Method (B):** This method, also known as the **deflection angle method**, is the most widely used and practical method for underground curves.

It requires only one theodolite set up at the beginning of the curve (the tangent point).

Points on the curve are located by turning a series of calculated small angles (deflection angles) relative to the tangent line and measuring a fixed chord length along that new line of sight.

This method relies primarily on accurate angular measurement, which a theodolite excels at, and is well-suited to the linear nature of a roadway.

Step 3: Final Answer:

The preferred method of setting out a curve in underground roadways is the Chord and angle method.

This corresponds to option (B).

💡 Quick Tip

In underground surveying, methods that rely on precise angular measurements from a single station are generally preferred over methods that require extensive linear measurements or multiple setups.

The chord and angle method fits this requirement perfectly.

167. The wire ropes which offers better wearing surface and more resistance to bending fatigue is

- (A) Lang's Lay
- (B) Ordinary lay
- (C) Left hand Lay
- (D) Right hand Lay

Correct Answer: (A) Lang's Lay

Solution:

Step 1: Understanding the Question:

We need to identify the type of wire rope construction ("lay") that is superior in terms of wear resistance and resistance to bending fatigue.

Step 2: Detailed Explanation:

The "lay" of a wire rope describes the direction in which the wires are twisted to form a strand, and the direction the strands are twisted to form the rope.

- **Ordinary Lay (or Regular Lay):** The wires in the strands are twisted in the opposite direction to the twist of the strands around the rope's core.

This construction results in the outer wires running roughly parallel to the rope's axis.

This makes it more resistant to crushing and distortion and less likely to untwist, but the short exposed length of the outer wires leads to faster wear and lower bending fatigue resistance.

- **Lang's Lay:** The wires in the strands are twisted in the **same direction** as the twist of the strands around the rope's core.

This construction results in the outer wires running at a diagonal angle to the rope's axis.

This exposes a longer length of each outer wire on the surface of the rope.

The advantages of this are:

1. **Better Wearing Surface:** The wear is distributed over a longer length of wire, so the rope wears down more slowly and has a longer service life in high-abrasion applications.
2. **More Resistance to Bending Fatigue:** The alignment of the wires and strands allows for better internal movement and stress distribution when the rope bends over sheaves or drums, making it more flexible and resistant to fatigue from repeated bending.

Left hand Lay and **Right hand Lay** refer to the direction of the strand twist (Z or S) and can apply to either Ordinary or Lang's lay ropes.

They do not by themselves define the wear and fatigue properties.

Step 3: Final Answer:

Lang's Lay ropes offer a better wearing surface and more resistance to bending fatigue.

This corresponds to option (A).

💡 Quick Tip

Remember the trade-offs:

- **Lang's Lay** = Better for wear and bending, but prone to untwisting and crushing.
Used where both ends of the rope are fixed, like haulage ropes.
- **Ordinary Lay** = More stable and crush-resistant, but wears faster.
Used for applications with free ends, like crane ropes and winding ropes.

168. The type of wire ropes used as Guide ropes in winding is

- (A) Stranded rope with FC
- (B) Non stranded rope
- (C) Stranded rope with steel wire core
- (D) IWRC

Correct Answer: (B) Non stranded rope

Solution:

Step 1: Understanding the Question:

The question asks to identify the specific type of wire rope construction used for guide ropes in a mine shaft winding system.

Step 2: Detailed Explanation:

Guide ropes in a shaft serve to guide the cages or skips, preventing them from swinging or colliding with each other or the shaft walls.

They need to be rigid, have a smooth surface, be resistant to rotation, and have a long life.

- **Stranded ropes (A, C, D):** These ropes are made of multiple strands twisted around a core (Fibre Core - FC, or Independent Wire Rope Core - IWRC).

Their twisted construction gives them a rougher surface profile and makes them prone to spinning under load.

While excellent as hoisting ropes due to their flexibility, they are not ideal as static guides.

- **Non stranded rope (or Full Locked Coil Rope):** This type of rope is specifically designed for applications like guide ropes, track ropes for cable cars, and bridge suspension cables.

It is constructed with an inner core of round wires, overlaid with one or more layers of shaped, interlocking wires (often Z-shaped).

This construction gives the rope several key advantages for use as a guide rope:

1. A perfectly smooth, cylindrical outer surface, which reduces wear on the guide shoes of the cage.
2. A very high packing factor (more steel in the cross-section), making it strong and rigid.
3. It is inherently non-rotating or "non-spin".

Step 3: Final Answer:

The type of wire ropes typically used as guide ropes are non-stranded ropes, specifically full locked coil ropes.

This corresponds to option (B).

💡 Quick Tip

Think about the function:

- **Hoisting ropes** need to be flexible to bend around drums → **Stranded**.
- **Guide ropes** need to be smooth and rigid like a rail → **Non-stranded (Locked Coil)**.

169. The method of joining of two wire ropes permanently without using special fittings or attachments is known as

- (A) Capping
- (B) Recapping
- (C) Splicing
- (D) Socketing

Correct Answer: (C) Splicing

Solution:

Step 1: Understanding the Question:

The question asks for the name of the process used to join two wire ropes together end-to-end to form a continuous loop or a longer rope.

Step 2: Detailed Explanation:

Let's define the terms:

- **Capping / Socketing:** These terms refer to the process of terminating the end of a wire rope by attaching a fitting (a "cap" or "socket").

This is used to connect the rope to a cage, skip, or other machinery.

- **Recapping:** This is the process of cutting off a section of rope from the end and re-attaching the cap.

It is done periodically to remove the most fatigued section of a winding rope.

- **Splicing:** This is the specific technique of joining two ropes together by unlaying the strands of each rope end for a certain length and then interweaving the strands from one rope with the strands of the other in a prescribed pattern.

The result is a continuous rope with a joint that is only slightly thicker than the original rope diameter.

This is the standard method for creating endless ropes for haulage systems.

Step 3: Final Answer:

The method of joining two wire ropes permanently without using special fittings is known as splicing.

This corresponds to option (C).

💡 Quick Tip

Think of the difference between connecting and joining:

- **Capping/Socketing** is for **connecting** the rope to something else.
- **Splicing** is for **joining** two ropes to make one longer rope.

170. How many times winding rope has to be recapped through out its life as per regulation ?

- (A) 5 times
- (B) 6 times
- (C) 7 times
- (D) 8 times

Correct Answer: (B) 6 times

Solution:

Step 1: Understanding the Question:

This is a regulatory question asking for the maximum number of times a winding rope can be recapped during its service life according to Indian mining regulations.

Step 2: Detailed Explanation:

Recapping is a mandatory safety procedure for mine winding ropes.

The end of the rope attached to the cage (the capel end) is subjected to the highest stresses, fatigue, and potential corrosion.

Regulations require this end section to be cut off and the rope to be re-capped at regular intervals to remove this deteriorated portion.

According to the Coal Mines Regulations (and similar metalliferous regulations) in India, a winding rope must be recapped at intervals not exceeding six months.

The statutory life of a winding rope used for man winding is typically limited to three and a half (3.5) years.

With a recapping interval of 6 months (0.5 years), the number of recappings in a 3.5-year life would be:

$$\text{Number of recappings} = \frac{3.5 \text{ years}}{0.5 \text{ years/recap}} = 7$$

However, this is the number of intervals. The question asks how many times it is recapped. It is recapped at the end of each interval. A rope is installed with its first cap. Then it is recapped after 6 months, 12 months, 18 months, 24 months, 30 months, and 36 months. At 42 months (3.5 years) it is discarded. This would be a total of 6 recappings after the initial installation. The regulation specifies the maximum number of recaps allowed. The DGMS (Directorate General of Mines Safety) circulars have clarified and standardized the practice. The common regulatory understanding is that a winding rope shall not be used for more than 3.5 years and shall not be recapped more than 6 times.

Step 3: Final Answer:

As per regulation, a winding rope has to be recapped a maximum of 6 times throughout its life.

This corresponds to option (B).

 Quick Tip

Regulatory questions require memorization of specific numbers.

For winding ropes in India, the key numbers are:

- Maximum life: **3.5 years**.
- Recapping interval: Every **6 months**.
- Maximum number of recaps: **6 times**.

171. The torques on the winding drum due to loaded and empty cage can be balanced as far as practicable by

- (A) Balancing rope
- (B) Use of suitable drum
- (C) Changing the load on two cages
- (D) Controlling the speed

Correct Answer: (A) Balancing rope

Solution:

Step 1: Understanding the Question:

The question asks for the method used to balance the static torque on a winding drum, which is caused by the weight of the ropes hanging in the shaft.

Step 2: Detailed Explanation:

In a balanced drum winding system, there are two cages.

When one cage (the loaded one) is at the bottom of the shaft, the other (empty) is at the top.

In this position, the full length of the hoisting rope for the bottom cage is hanging in the shaft, while almost no rope is hanging for the top cage.

This difference in the weight of the suspended ropes creates a large out-of-balance load or torque that the winder motor must overcome, especially at the start of the wind.

To counteract this, a **balancing rope** (or tail rope) is used.

This is a separate rope that hangs in a loop at the bottom of the shaft, connecting the undersides of the two cages.

As the loaded cage ascends and its hoisting rope gets shorter, the balancing rope on its side gets longer.

Simultaneously, as the empty cage descends and its hoisting rope gets longer, its balancing rope gets shorter.

The effect is that the total weight of rope hanging on each side of the drum remains nearly constant throughout the wind, thus balancing the static torque.

Step 3: Final Answer:

The torques are balanced by using a balancing rope.

This corresponds to option (A).

💡 Quick Tip

A balancing rope acts as a counterweight that moves from one side of the system to the other.

Its purpose is to cancel out the changing weight of the main winding ropes as the cages move up and down.

This reduces the peak power demand on the winder motor, making it more energy-efficient.

172. In drum winding the shaft fitting used for smooth transfer of load between the cages and winding ropes during loading and unloading operation is

- (A) Guides
- (B) Lilly controller
- (C) Safety hook
- (D) Keps

Correct Answer: (D) Keps

Solution:

Step 1: Understanding the Question:

The question asks to identify the device in a shaft that supports the cage during loading and unloading to prevent rope oscillations.

Step 2: Detailed Explanation:

Let's define the functions of the listed shaft fittings:

- **Guides:** These are the rails (rigid guides) or ropes (rope guides) that guide the cage up and down the shaft to prevent it from swinging.
- **Lilly controller:** This is an essential safety device on the winder itself that prevents over-speeding and over-winding.
- **Safety hook (Detaching Hook):** This device is placed between the rope and the cage.

In an overwind, it detaches the rope and simultaneously engages catches in the headframe to prevent the cage from falling back down the shaft.

- **Keps (or Catches):** These are retractable supports installed at the pit top (and sometimes pit bottom) landings.

When the cage arrives at the landing, the keps are extended to support the cage's weight directly.

This allows the winding rope to be slackened slightly.

The primary purpose of keps is to hold the cage steady at a fixed level for loading/unloading and to take the shock loads of loading/unloading off the rope, preventing dangerous oscillations and reducing wear.

Step 3: Final Answer:

The shaft fitting used to support the cage for smooth load transfer is the keps.

This corresponds to option (D).

💡 Quick Tip

Imagine trying to load a heavy car onto an elevator held only by a stretchy bungee cord. The loading process would cause it to bounce violently. Keps are like solid blocks that slide under the elevator to hold it firm while you load, preventing this bouncing.

173. The type of haulage system used in case of undulating roadways

- (A) Direct rope haulage
- (B) Endless rope haulage
- (C) Main and Tail rope haulage
- (D) Gravity rope haulage

Correct Answer: (C) Main and Tail rope haulage

Solution:

Step 1: Understanding the Question:

We need to identify the rope haulage system that is suitable for use in roadways with an uneven or "undulating" gradient (i.e., with both uphill and downhill sections).

Step 2: Detailed Explanation:

Let's analyze the suitability of each system:

- **Direct Rope Haulage:** This uses a single rope to pull a set of tubs (a "train" or "set") up a gradient.

The empty tubs return downhill by gravity.

This system only works on a continuous uphill gradient and is unsuitable for undulations.

- **Endless Rope Haulage:** This uses a continuous loop of rope that moves slowly in one direction.

Tubs are attached to the rope individually or in small groups.

While it can handle some changes in gradient, it is most efficient on gentle, relatively uniform gradients and can have problems with rope tension and tub stability on severe undulations.

- **Main and Tail Rope Haulage:** This system uses two separate ropes and two separate drums on the hauler engine.

The "main" rope is attached to the front of the train to pull it inbye (towards the face).

The "tail" rope is attached to the back of the train, goes around a return pulley at the far end, and comes back to the second drum.

To move the train outbye, the tail rope drum is engaged to pull the train back.

Because there is positive control from both the front and the back, this system can pull the train up gradients, brake it down gradients, and navigate any combination of uphill and downhill sections.

It is the most flexible system and perfectly suited for undulating roadways.

- **Gravity Rope Haulage:** This system uses the weight of loaded tubs going downhill to pull empty tubs uphill.

It requires a specific gradient profile and is not suitable for undulations.

Step 3: Final Answer:

The haulage system used for undulating roadways is Main and Tail rope haulage.

This corresponds to option (C).

💡 Quick Tip

Think of the ropes as providing control:

- **Direct Haulage** = 1 rope (pulls up only).
- **Endless Haulage** = 1 continuous loop (slow, steady pull).
- **Main Tail Haulage** = 2 ropes (one pulls, one controls/pulls back).

The two-rope system gives the positive control needed to handle both uphill and downhill sections in a single trip.

174. Identify the safety device used in endless rope haulage system

- (A) Back stay
- (B) Small man clip
- (C) Cam clip
- (D) Monkey catch

Correct Answer: (D) Monkey catch

Solution:

Step 1: Understanding the Question:

We need to identify which of the listed devices is a safety device specifically used with an endless rope haulage system.

Step 2: Detailed Explanation:

- **Back stay (or Back Sprag / Drag):** This is a safety device used in **direct rope haulage**.

It is a pointed bar attached to the rear of the last tub of an ascending train.

If the rope breaks or the train becomes detached, the back stay digs into the ground and prevents the train from running away back down the incline.

- **Small man clip, Cam clip:** These are types of clips used to attach tubs to the haulage rope.

While their correct function is a matter of safety, they are operational components, not dedicated safety devices in the same sense as an anti-runaway device.

- **Monkey catch (or Tilting Rail):** This is a safety device used in **endless rope haulage** on inclined sections.

It consists of a short, hinged section of rail that is weighted to stay in a raised position.

Tubs moving in the intended upward direction can easily push the catch down and pass over it.

However, if the rope breaks and the tub starts to run backwards, the catch will be in its raised position, and the axle of the runaway tub will hit it, derailing the tub and stopping it safely.

Step 3: Final Answer:

A safety device used in an endless rope haulage system is the Monkey catch.

This corresponds to option (D).

💡 Quick Tip

Associate the runaway protection device with the haulage type:

- **Direct Haulage** (trains going up): **Back stay** (drags behind).
- **Endless Haulage** (single tubs going up): **Monkey catch** (catches the axle if it runs back).

175. The type of locomotive containing exhaust conditioner box along with flame trap arrangement is

- (A) Diesel locomotive
- (B) Battery locomotive
- (C) Trolley wire locomotive
- (D) Steam locomotive

Correct Answer: (A) Diesel locomotive

Solution:

Step 1: Understanding the Question:

We need to identify which type of locomotive requires an exhaust conditioning system and flame traps for safe operation, particularly in an underground environment.

Step 2: Detailed Explanation:

- **Battery and Trolley wire locomotives:** These are electric locomotives.

They do not have an internal combustion engine and therefore produce no exhaust gases.

They do not require exhaust conditioners or flame traps.

- **Steam locomotive:** These burn fuel (coal or wood) externally to create steam.

While they have an exhaust (steam and combustion products), they are generally not used in underground mining due to the large amount of heat, smoke, and steam they produce.

- **Diesel locomotive:** These use an internal combustion diesel engine.

The exhaust from a diesel engine is very hot and contains sparks, as well as noxious gases like oxides of nitrogen (NO_x) and carbon monoxide (CO).

For safe use in an underground mine (especially a coal mine), the exhaust must be treated:

1. A **flame trap** is fitted to the exhaust manifold to quench any sparks or flames and prevent them from igniting flammable gas in the mine atmosphere.
2. An **exhaust conditioner box** (or scrubber) is used to cool the hot exhaust gases and to dissolve and remove some of the noxious gases.

This is typically done by passing the exhaust through a water bath.

Step 3: Final Answer:

The locomotive that requires an exhaust conditioner and flame trap is the diesel locomotive.

This corresponds to option (A).

💡 Quick Tip

Anytime you use an internal combustion engine in a gassy underground mine, you have two major hazards: toxic exhaust gases and the risk of igniting methane.

The exhaust conditioner handles the toxic gases, and the flame trap handles the ignition risk.

Only diesel locomotives on this list have an internal combustion engine.

176. In a turbine pump the direction of End Thrust is

- (A) In the middle of the pump
- (B) From suction end towards delivery end
- (C) From delivery end towards Suction end
- (D) Any where on the shaft of Turbine pump

Correct Answer: (C) From delivery end towards Suction end

Solution:

Step 1: Understanding the Question:

The question asks for the direction of the net axial force, known as "end thrust," on the impeller shaft of a turbine pump.

Step 2: Detailed Explanation:

A turbine pump (a type of centrifugal pump) works by spinning an impeller to add energy to a fluid.

The fluid enters the impeller at the "suction" eye (low pressure) and is thrown outwards to the "delivery" or discharge end (high pressure).

This pressure difference acts on the surfaces of the impeller, creating an axial force.

The main causes of end thrust are:

1. **Pressure on Impeller Shrouds:** The high-pressure water on the delivery side acts on the outer face of the impeller shroud.

The low-pressure water on the suction side acts on the inner face (the eye).

Because the area on the delivery side is larger than the area of the suction eye, there is a net force pushing the impeller.

2. **Change in Momentum:** The fluid enters the pump axially and exits radially (or at an angle).

This change in the direction of momentum also creates a reaction force on the impeller.

The combined effect of these forces is a net axial thrust that is directed from the high-pressure (delivery) side towards the low-pressure (suction) side.

This thrust must be managed by thrust bearings to prevent damage to the pump.

Step 3: Final Answer:

The direction of the end thrust is from the delivery end towards the suction end.

This corresponds to option (C).

💡 Quick Tip

Think of the pump like a rocket.

High pressure is created at the back (delivery side) and pushes forward.

The "forward" direction for the impeller is towards the low-pressure intake (suction side).

So, the thrust is always from High Pressure to Low Pressure.

177. The reverse direction rotation of Belt conveyor can be prevented by

- (A) Hold back
- (B) Stop back
- (C) Back stay
- (D) Pull chord

Correct Answer: (A) Hold back

Solution:

Step 1: Understanding the Question:

We need to identify the safety device used to prevent an inclined belt conveyor from running backwards under the weight of its load if the power fails.

Step 2: Detailed Explanation:

When a loaded belt conveyor is operating on an incline, the material on the belt exerts a gravitational pull that tries to make the belt run backwards.

If the drive motor stops for any reason (e.g., power failure), this gravitational force will cause the belt to accelerate in reverse, which is an extremely dangerous situation.

To prevent this, a mechanical safety device is installed on the drive pulley shaft.

- **Holdback (or Backstop):** This is the correct term for this device.

It is a type of one-way clutch or ratchet mechanism.

It allows the drive shaft to rotate freely in the forward (operating) direction but instantly engages and locks if the shaft attempts to rotate in the reverse direction, thus "holding back" the conveyor.

- **Back stay:** This is a device for rope haulage tubs, not belt conveyors.

- **Pull chord:** This is an emergency stop system that runs along the length of the conveyor, allowing it to be stopped from any point.

It does not prevent reverse motion.

Step 3: Final Answer:

The reverse rotation of a belt conveyor can be prevented by a holdback.

This corresponds to option (A).

💡 Quick Tip

The name "holdback" perfectly describes its function: it holds the belt back from running in reverse.

It's a mandatory safety feature for any inclined conveyor system.

178. In Longwall face machinery Shearer is mounted on

- (A) Skid
- (B) AFC
- (C) Hydraulic rams
- (D) Double acting piston

Correct Answer: (B) AFC

Solution:

Step 1: Understanding the Question:

The question asks to identify the piece of equipment on which a longwall shearer is mounted and travels.

Step 2: Detailed Explanation:

A longwall face is a highly integrated system of machinery:

1. **Powered Roof Supports (Chocks or Shields):** These hold up the roof in front of the miners and machinery.
2. **Shearer:** This is the cutting machine.

It has large rotating drums with picks that cut coal from the face as it moves back and forth.

3. **Armoured Face Conveyor (AFC):** This is a heavy-duty chain conveyor that runs along the entire length of the coal face, directly in front of the roof supports.

The shearer is not a self-propelled machine in the way a car is.

It is mounted directly onto the AFC.

The AFC serves two critical functions:

- a. It acts as a track or guide for the shearer to travel along the face.
- b. It transports the coal that is cut by the shearer away from the face.

The entire face system (AFC and supports) is pushed forward by hydraulic rams.

Step 3: Final Answer:

The longwall shearer is mounted on the AFC (Armoured Face Conveyor).

This corresponds to option (B).

💡 Quick Tip

Think of the longwall face system as a train.

The **AFC** is the track.

The **Shearer** is the engine that moves along the track.

The **Roof Supports** are the station platform that moves along with the track.

179. The type of Subsidence in which maximum subsidence occurs over an area is known as

- (A) Critical width of extraction
- (B) Sub Critical width of extraction
- (C) Super Critical width of extraction
- (D) Narrow width of extraction

Correct Answer: (C) Super Critical width of extraction

Solution:

Step 1: Understanding the Question:

The question asks to identify the term that describes an extraction panel width where the maximum possible subsidence occurs.

Step 2: Detailed Explanation:

Surface subsidence is the lowering of the ground surface due to underground mining.

The amount of subsidence at the center of a mining panel is directly related to the width of the panel relative to the depth of the workings.

- **Sub-Critical Width:** When the extraction panel is relatively narrow compared to its depth, the rock mass above the panel can bridge or arch across the opening to some extent.

The subsidence at the surface is less than the maximum possible value.

- **Critical Width:** This is the specific panel width at which the effect of the extraction just reaches the surface point directly above the center of the panel.

At this width, the subsidence at the center point reaches its maximum possible value for the first time.

The maximum subsidence is equal to the seam thickness multiplied by a subsidence factor.

- **Super-Critical Width:** If the panel is made even wider than the critical width, the subsidence at the center point **does not increase further**.

It remains at the maximum possible value.

However, a trough of maximum subsidence develops over the central part of the panel.

The question asks for the condition where "maximum subsidence occurs over an area", which implies a trough or flat-bottomed subsidence profile.

This only happens when the extraction width is super-critical.

Step 3: Final Answer:

The type of subsidence where maximum subsidence occurs over an area is associated with a Super Critical width of extraction.

This corresponds to option (C).

💡 Quick Tip

Think of the subsidence profile shape:

- **Sub-critical:** A gentle "U" shape. The bottom doesn't reach the maximum possible depth.
- **Critical:** A "V" shape. The bottom point just touches the maximum depth.
- **Super-critical:** A flat-bottomed "U" shape. A whole area at the bottom is at the maximum depth.

180. In opencast blasting the purpose of sub-grade drilling is

- (A) To avoid the formation of toe
- (B) To reduce ground vibrations

- (C) To reduce consumption of explosive
- (D) To improve over all blasting performance

Correct Answer: (A) To avoid the formation of toe

Solution:

Step 1: Understanding the Question:

The question asks for the specific reason why blast holes in opencast mines are drilled deeper than the planned bench floor level.

Step 2: Detailed Explanation:

In opencast (surface) mining, rock is excavated in a series of steps called benches.

Blast holes are drilled vertically from the top of the bench to break the rock.

The **bench floor** or **grade** is the desired final level after the broken rock is removed.

- **Toe:** The "toe" is the rock mass at the bottom of the bench face.

This region is the most confined and the hardest to break effectively with explosives.

If the blast holes are drilled only to the level of the bench floor, the explosive energy is often insufficient to fully fracture the rock at the toe.

This results in a hard, protruding ridge of unbroken rock, known as a "hard toe," being left behind.

This toe obstructs digging by excavators and requires costly and dangerous secondary blasting to remove.

- **Sub-grade drilling (or Sub-drilling):** To prevent the formation of a toe, the blast holes are intentionally drilled a certain distance **below** the planned bench floor level.

This practice ensures that there is sufficient explosive energy placed at and below the toe to completely break the rock mass, resulting in a clean, flat floor that is easy to dig.

While this may slightly improve overall performance (D), its specific and primary purpose is to deal with the toe problem.

Step 3: Final Answer:

The purpose of sub-grade drilling is to avoid the formation of a toe.

This corresponds to option (A).

💡 Quick Tip

Imagine trying to break a block of concrete with a hammer. If you only hit the top, the bottom might not break cleanly. Sub-grade drilling is like making sure your crack goes all the way through the bottom, ensuring a clean break at the desired level.

181. Which one of the following is NOT a Controlled blasting technique used in Opencast mines ?

- (A) Pre splitting
- (B) Deck charging
- (C) Cushion blasting
- (D) Cast blasting

Correct Answer: (D) Cast blasting

Solution:

Step 1: Understanding the Question:

We need to distinguish between "controlled blasting" techniques and other types of blasting, and identify the option that doesn't fit the "controlled" category.

Step 2: Detailed Explanation:

Controlled Blasting refers to a set of specialized techniques used to achieve a specific outcome, usually to minimize damage to the rock that is intended to remain in place (the final wall).

These techniques aim to produce a smooth, stable, and crack-free final pit slope.

- **Pre-splitting:** This involves drilling a line of closely spaced holes along the final wall line and firing them *before* the main production blast.

This creates a fracture plane that isolates the final wall from the shock of the main blast.

It is a key controlled blasting technique.

- **Cushion Blasting (or Trim Blasting):** This involves firing a final row of lightly charged holes along the pit limit *after* the main blast has been completed.

It "trims" the wall to its final position with minimal damage.

It is also a key controlled blasting technique.

- **Deck Charging:** This is the technique of placing multiple, separated explosive charges (decks) within a single blast hole, separated by inert stemming material.

It is used to distribute explosive energy more effectively and can be used as part of a controlled blast design to reduce damage.

- **Cast Blasting:** This is a **bulk production blasting** technique, not a controlled one.

Its objective is the opposite of preserving a final wall.

Cast blasting uses a large amount of explosive energy to physically throw a significant percentage of the broken overburden directly from the bench into the adjacent mined-out pit.

It is designed for maximum displacement of rock, not minimum damage.

Step 3: Final Answer:

Cast blasting is a bulk blasting method, not a controlled blasting technique.

This corresponds to option (D).

💡 Quick Tip

Think of the goal:

- **Controlled Blasting** = Precision and care, like a surgeon's scalpel (e.g., Pre-splitting).
- **Cast Blasting** = Brute force and movement, like a catapult.

182. In an opencast bench two geologically disturbed planes intersect each other. The type of bench failures likely to take place

- (A) Circular failure
- (B) Wedge failure
- (C) Planar failure
- (D) Toppling failure

Correct Answer: (B) Wedge failure

Solution:

Step 1: Understanding the Question:

We need to identify the type of slope failure that is caused by the intersection of two geological discontinuity planes (like faults or joints).

Step 2: Detailed Explanation:

Slope stability in rock masses is largely controlled by the orientation and properties of geolog-

ical structures (discontinuities).

Different failure modes can occur:

- **Planar Failure:** This occurs when a block of rock slides down a **single**, continuous discontinuity plane that is dipping out of the slope face at an angle steeper than the friction angle.
- **Wedge Failure:** This occurs when **two discontinuities intersect** in such a way that their line of intersection dips out of the slope face.

The block of rock defined by the two intersecting planes and the bench face (a "wedge") can slide out along the line of intersection.

This perfectly matches the condition described in the question.

- **Circular Failure:** This type of failure occurs in very weak, heavily fractured rock masses or in soils.

The failure surface is curved or spoon-shaped and is not controlled by any single structure.

- **Toppling Failure:** This occurs in rock masses with steeply dipping layers that are oriented parallel to the slope face.

Columns of rock can rotate and topple forwards out of the slope, like a row of dominoes.

Step 3: Final Answer:

The failure caused by two intersecting planes is a wedge failure.

This corresponds to option (B).

💡 Quick Tip

Remember the geometric cause of each failure type:

- **1 Plane** → Planar Failure - **2 Intersecting Planes** → Wedge Failure - **No Dominant Planes** (weak mass) → Circular Failure - **Steeply Dipping Layers** → Toppling Failure

183. The strain measuring device used to measure the convergence between the roof and floor at a remote location of a mine is

- (A) Tell tale bore hole extensometer
- (B) Magnetic ring multipoint bore hole extensometer
- (C) Remote convergence recorder

(D) Load cells

Correct Answer: (C) Remote convergence recorder

Solution:

Step 1: Understanding the Question:

The question asks for an instrument that specifically measures roof-to-floor convergence (the coming together of the roof and floor) and can be read from a remote location.

Step 2: Detailed Explanation:

Let's analyze the instruments:

- **Extensometers (A and B):** These instruments are installed in boreholes to measure movement *within* the rock mass.

They measure the displacement between different points along the length of the borehole, not the total convergence between the roof and floor surfaces.

A "tell tale" is a simple mechanical device for visual indication of roof movement, not a remote recorder.

- **Load Cells:** These instruments measure force or pressure, for example, the load acting on a roof support.

They do not directly measure displacement or convergence.

- **Remote Convergence Recorder:** This instrument is specifically designed for the purpose described.

It consists of a sensor unit that is installed between the roof and floor.

The sensor measures the change in distance between its two ends as the roof and floor converge.

This measurement is converted into an electrical signal and transmitted via a cable to a readout unit or data logger located in a safe, **remote** location.

This allows for continuous and safe monitoring of convergence in inaccessible or hazardous areas.

Step 3: Final Answer:

The device used to measure convergence remotely is a remote convergence recorder.

This corresponds to option (C).

💡 Quick Tip

Break down the name of the instrument to understand its function:

- **Convergence:** Measures roof-to-floor closure.
- **Recorder:** Logs the data over time.
- **Remote:** Can be read from a safe distance.

The name perfectly describes the instrument's purpose.

184. In opencast mine planning the report which assess the periodical impact of mining over environment and suggest corresponding control measures is

- (A) EIA
- (B) EMP
- (C) DPR
- (D) Feasibility Report

Correct Answer: (A) EIA

Solution:

Step 1: Understanding the Question:

The question asks for the name of the report that evaluates the environmental impacts of a proposed mining project.

Step 2: Detailed Explanation:

Let's define the reports involved in mine planning:

- **EIA (Environmental Impact Assessment):** This is a formal process and report that is required by law for most major projects, including mines.

Its purpose is to identify, predict, evaluate, and mitigate the potential environmental, social, and other relevant effects of the project *before* major decisions are made.

It assesses the impact of mining and suggests control measures.

- **EMP (Environmental Management Plan):** This is a detailed plan that is developed as part of the EIA process.

It describes the specific actions, responsibilities, and schedules that the mining company will follow to implement the mitigation and monitoring measures identified in the EIA.

While the EIA *assesses* the impact, the EMP *manages* it.

The question asks for the report that "assesses the... impact... and suggests... control measures," which is the primary role of the EIA.

- **DPR (Detailed Project Report) and Feasibility Report:** These are technical and economic documents that detail the entire mining plan, including geology, reserves, mining method, equipment, and financial viability.

The EIA/EMP is a critical component of these reports.

Step 3: Final Answer:

The report that assesses the environmental impact is the Environmental Impact Assessment (EIA).

This corresponds to option (A).

💡 Quick Tip

Remember the sequence and purpose:

1. **EIA = Assess** the problem (What are the impacts?).
2. **EMP = Make a Plan** to fix the problem (How will we control the impacts?).

The assessment comes before the management plan.

185. The value of Shear stress along principal plane of a failure surface is

- (A) Minimum
- (B) Maximum
- (C) Neither maximum nor minimum
- (D) Zero

Correct Answer: (D) Zero

Solution:

Step 1: Understanding the Question:

This is a fundamental question in mechanics and rock mechanics, asking for the value of shear stress on a principal plane.

Step 2: Detailed Explanation:

In a stressed body, there exist three mutually perpendicular planes where the stress acting on them is purely normal (perpendicular) to the plane.

These special planes are called the **principal planes**.

The normal stresses acting on these planes are called the **principal stresses** (denoted as σ_1 , σ_2 , and σ_3).

By the very definition of a principal plane, there are **no shear stresses** acting on it.

The shear stress reaches its maximum value on planes that are oriented at 45 degrees to the principal planes.

Since a principal plane is defined as a plane of zero shear stress, the value of shear stress along it must be zero.

Step 3: Final Answer:

The value of shear stress along a principal plane is zero.

This corresponds to option (D).

💡 Quick Tip

Principal Plane = Plane of Zero Shear Stress.

This is a core definition in stress analysis.

The maximum normal stress (σ_1) and minimum normal stress (σ_3) occur on principal planes, but the shear stress on these specific planes is always zero.

186. The Indirect method of determination of Tensile strength of a rock specimen is

- (A) Confined Compressive strength test
- (B) Un-Confined compressive strength test
- (C) Brazilian test
- (D) Protodyaknov Strength Index

Correct Answer: (C) Brazilian test

Solution:

Step 1: Understanding the Question:

We need to identify the laboratory test that is used to *indirectly* measure the tensile strength of a rock sample.

Step 2: Detailed Explanation:

Directly pulling on a rock specimen to measure its tensile strength is very difficult.

The sample is hard to grip without crushing it, and it's difficult to ensure a pure tensile failure.

Therefore, indirect methods are commonly used.

- **Compressive Strength Tests (A and B):** These tests measure the rock's strength under compression, not tension.
- **Protodyakov Strength Index:** This is an impact test used to estimate rock strength, primarily related to its cuttability and drillability, not a standard measure of tensile strength.
- **Brazilian Test (or Indirect Tensile Test):** This is the standard indirect method for determining the tensile strength of rock.

In this test, a cylindrical disc of rock is placed on its side and compressed along its diameter.

This compressional loading induces a relatively uniform tensile stress across the vertical diameter of the disc.

The disc fails by splitting in tension along this diameter.

The tensile strength is then calculated from the compressional load at which failure occurs.

Step 3: Final Answer:

The indirect method for determining the tensile strength of a rock specimen is the Brazilian test.

This corresponds to option (C).

💡 Quick Tip

When you see "indirect tensile strength test" for rock, the answer is almost always the "Brazilian test."

It's a clever and universally adopted method to overcome the practical difficulties of a direct tension test on brittle materials like rock.

187. Which one of the following term NOT related to opencast mining ?

- (A) Boxcut
- (B) Highwall
- (C) Crest
- (D) Winze

Correct Answer: (D) Winze

Solution:

Step 1: Understanding the Question:

We need to identify the term from the list that is not used in the context of opencast (surface)

mining but rather in underground mining.

Step 2: Detailed Explanation:

Let's define the terms:

- **Boxcut:** This is the initial excavation made in an opencast mine to create a working face and a ramp for access when starting from a flat surface.

This is an opencast mining term.

- **Highwall:** This is the unexcavated face or wall of an opencast pit, from the crest to the toe of a bench.

This is an opencast mining term.

- **Crest:** This is the top edge of a bench or highwall in an opencast mine.

This is an opencast mining term.

- **Winze:** As defined in a previous question, a winze is an **underground** vertical or inclined opening driven **downwards** to connect two levels.

This term has no application in opencast mining.

Step 3: Final Answer:

The term not related to opencast mining is "Winze".

This corresponds to option (D).

💡 Quick Tip

Opencast mining terms often describe surface features like benches, highwalls, crests, toes, and ramps.

Underground mining terms describe internal openings like shafts, adits, drifts, raises, and winzes.

188. The relation between porosity (n) and void ratio (e) of a rock specimen is

- (A) $e = \frac{n}{1-n}$
- (B) $e = \frac{n}{1+n}$
- (C) $e = \frac{1-n}{n}$
- (D) $e = \frac{1+n}{n}$

Correct Answer: (A) $e = \frac{n}{1-n}$

Solution:

Step 1: Understanding the Question:

We need to find the mathematical relationship that expresses void ratio (e) in terms of porosity (n).

Step 2: Key Formula or Approach:

The definitions of porosity and void ratio are:

- Porosity (n) = $\frac{\text{Volume of voids } (V_v)}{\text{Total volume } (V_t)}$
- Void ratio (e) = $\frac{\text{Volume of voids } (V_v)}{\text{Volume of solids } (V_s)}$

The total volume is the sum of the volume of solids and the volume of voids: $V_t = V_s + V_v$.

Step 3: Detailed Explanation:

Start with the definition of porosity:

$$n = \frac{V_v}{V_t} = \frac{V_v}{V_s + V_v}$$

To introduce the void ratio $e = V_v/V_s$, divide the numerator and the denominator of the porosity equation by the volume of solids (V_s):

$$n = \frac{V_v/V_s}{(V_s/V_s) + (V_v/V_s)}$$

Substitute e into the equation:

$$n = \frac{e}{1 + e}$$

Now, we must rearrange this equation to solve for e :

$$n(1 + e) = e$$

$$n + ne = e$$

$$n = e - ne$$

$$n = e(1 - n)$$

$$e = \frac{n}{1 - n}$$

Step 4: Final Answer:

The relation between void ratio and porosity is $e = \frac{n}{1-n}$.
This corresponds to option (A).

 Quick Tip

You can remember the two relationships:

$$n = \frac{e}{1 + e} \quad \text{and} \quad e = \frac{n}{1 - n}$$

A good way to check your answer is to remember that porosity (n) must always be less than 1, while void ratio (e) can be greater than 1.

If $n = 0.5$, then $e = 0.5/(1 - 0.5) = 1$, which makes sense.

189. The following are the various parameters of the Shovel

Capacity of the bucket – 20 m³

Fill factor (f) = 0.8

Swell factor (s) = 0.5

Cycle time (t) = 60 sec

The Volume of the material carried by the Shovel per hour is

(A) 300 m³

(B) 388 m³

(C) 400 m³

(D) 480 m³

Correct Answer: (D) 480 m³

Solution:

Step 1: Understanding the Question:

We need to calculate the hourly production of a shovel in terms of volume, given its operational parameters. The final answer will depend on whether the output is measured in loose volume or bank (in-situ) volume.

Step 2: Key Formula or Approach:

1. Calculate the actual loose volume per bucket cycle.
2. Calculate the number of cycles per hour.
3. Calculate the hourly production. The use of the swell factor suggests the question may be asking for the bank volume.

- Loose Volume per cycle = Bucket Capacity $\times$ Fill Factor - Bank Volume per cycle = Loose Volume per cycle $\times$ Swell Factor (Here, swell factor is defined as Bank Volume / Loose Volume, or $1/(1 + \text{swell } \%)$). - Cycles per hour = $3600 / \text{Cycle time (in sec)}$ - Production per hour = Volume per cycle $\times$ Cycles per hour

Step 3: Detailed Explanation:

1. Calculate Loose Volume per Cycle:

The fill factor (0.8) indicates that the bucket is not completely full.

$$\text{Loose Volume per Cycle} = 20 \text{ m}^3 \times 0.8 = 16 \text{ m}^3 \text{ (LCM)}$$

2. Apply Swell Factor:

The swell factor (0.5) relates the bank (in-situ) volume to the loose volume. The calculation that leads to the correct answer implies a definition where Bank Volume = Loose Volume $\times$ Swell Factor. This is an unconventional definition but required to solve the problem as posed.

$$\text{Bank Volume per Cycle} = 16 \text{ m}^3 \times 0.5 = 8 \text{ m}^3 \text{ (BCM)}$$

3. Calculate Cycles per Hour:

$$\text{Cycles per hour} = \frac{3600 \text{ sec/hr}}{60 \text{ sec/cycle}} = 60 \text{ cycles/hr}$$

4. Calculate Hourly Production:

The question asks for the volume of material carried per hour. Based on the calculation path, this refers to the bank volume.

$$\text{Hourly Production (Bank Volume)} = 8 \text{ m}^3/\text{cycle} \times 60 \text{ cycles/hr} = 480 \text{ m}^3/\text{hr}$$

Step 4: Final Answer:

The volume of the material (in bank measure) carried by the shovel per hour is 480 m^3 . This corresponds to option (D).

💡 Quick Tip

Be very careful with equipment productivity problems, as definitions of factors can vary.

Here, the calculation path to the answer shows that the required output is Bank Cubic Metres (BCM) and the "swell factor" is used as $V_{\text{bank}}/V_{\text{loose}}$.

The standard formula for shovel output in LCM is: $Q = (V \times f \times 3600)/t$. This would give 960 LCM/hr.

190. The width of the haul road in an opencast mine depends on

- (A) Length of the dumper
- (B) Width of the largest vehicle moving along haul road
- (C) Number of dumpers moving
- (D) Targeted production

Correct Answer: (B) Width of the largest vehicle moving along haul road

Solution:

Step 1: Understanding the Question:

The question asks for the primary physical parameter that determines the required width of a

haul road in a surface mine.

Step 2: Detailed Explanation:

The design of a haul road is critical for safety and efficiency.

The width of the road is determined by the need to safely accommodate the traffic that will use it.

The most fundamental design parameter is the size of the vehicles themselves.

Standard design guidelines for haul road width are based on the width of the **largest vehicle** that will use the road.

For example, a common rule of thumb for a two-lane haul road is that the total width should be approximately 3.5 times the width of the widest truck.

This provides space for two trucks to pass each other safely, plus clearance on either side and a safety berm.

The other options are incorrect:

- The length of the dumper affects the turning radius at curves, but not the straight-section width.
- The number of dumpers and targeted production affect the traffic density and may influence the decision to build a two-lane vs. a four-lane road, but the fundamental width of each lane is still determined by the vehicle size.

Step 3: Final Answer:

The width of the haul road depends on the width of the largest vehicle moving along the haul road.

This corresponds to option (B).

💡 Quick Tip

Think of it like designing a highway.

The width of a traffic lane is determined by the width of a standard car or truck, not by how many cars are on the road or how long the cars are.

The same logic applies to mine haul roads, just with much larger vehicles.

191. The place of Drinking water facility on the surface of mine should be away from urinals, latrine and washing places at least at a distance of

- (A) 6m
- (B) 5 m
- (C) 7 m

(D) 4 m

Correct Answer: (A) 6m

Solution:

Step 1: Understanding the Question:

This is a regulatory question concerning health and safety standards in a mine, specifically the minimum required distance between drinking water sources and sanitation facilities.

Step 2: Detailed Explanation:

This requirement is specified to prevent the contamination of drinking water and the spread of disease.

According to the Mines Rules, 1955, under Rule 31, which deals with the provision of drinking water, it is stated that:

”Every place where drinking water is provided shall be situated not less than **six metres** away from any washing place, urinal or latrine, unless a shorter distance is approved in writing by the Chief Inspector.”

Therefore, the legally mandated minimum distance is 6 metres.

Step 3: Final Answer:

The minimum distance required is 6m.

This corresponds to option (A).

 Quick Tip

Regulatory questions in mining exams often test knowledge of specific numbers and distances from the Mines Act and Mines Rules.

It is important to memorize key safety-related values like this one.

192. Quorum of a Committee meeting as per Mines Act-1952 is

- (A) 3 Members including Chairmen
- (B) 3 Members excluding Chairmen
- (C) 4 Members including Chairmen
- (D) 4 Members excluding Chairmen

Correct Answer: (C) 4 Members including Chairmen

Solution:

Step 1: Understanding the Question:

The question asks for the quorum (minimum number of members required to be present for the meeting to be valid) for a Committee constituted under the Mines Act, 1952.

Step 2: Detailed Explanation:

The constitution and functioning of the Committee are detailed in Section 12 of the Mines Act, 1952 and Chapter II of the Mines Rules, 1955.

The composition of the committee includes a Chairman and several other members representing various stakeholders.

Rule 8(5) of the Mines Rules, 1955, specifies the quorum for a meeting of this Committee.

It states: "To constitute a quorum for a meeting of the Committee, there shall not be less than three members present besides the chairman."

This means the minimum number of people required is:

1 (Chairman) + 3 (Other Members) = 4 people.

Therefore, the quorum is 4 members, including the Chairman.

Step 3: Final Answer:

The quorum for a Committee meeting is 4 Members including the Chairmen.

This corresponds to option (C).

💡 Quick Tip

Pay close attention to the wording "including" vs. "excluding" the chairman in regulatory questions.

The rule specifies the number of members needed *in addition to* the chairman, so you must add the chairman to that number to get the total quorum.

193. The meaning of term 'misfire' as per regulations is

- (A) Failure of ignition of explosive
- (B) Improper ignition of explosive
- (C) Failure to explode entire charge of Explosive
- (D) Failure to explode part of an explosive

Correct Answer: (C) Failure to explode entire charge of Explosive

Solution:

Step 1: Understanding the Question:

The question asks for the definition of a "misfire" according to mining regulations.

Step 2: Detailed Explanation:

According to the Coal Mines Regulations, 2017 [Regulation 2(1)(zd)], a 'misfire' is defined as:

"...the failure to explode of the **whole or part** of a charge of explosive in a shot-hole;"

Let's analyze the given options in light of this official definition:

- (A) Failure of ignition of explosive: This is a cause of a misfire, not the definition of the event itself.
- (B) Improper ignition of explosive: This is also a potential cause.
- (C) Failure to explode entire charge of Explosive: This describes one type of misfire (a complete misfire).
- (D) Failure to explode part of an explosive: This also describes a type of misfire (a partial misfire).

The official definition includes both partial and complete failure.

However, in multiple-choice questions, we must select the best possible answer.

Option (C) describes the most common and clear-cut understanding of a misfire.

While the regulation is more comprehensive, option (C) is the most accurate description of the event among the choices provided, as a partial failure (D) is a subset of the general problem of the charge not exploding as intended.

Given the provided answer key, the intended answer is the failure of the entire charge.

Step 3: Final Answer:

As per the options provided, the meaning of a misfire is the failure to explode the entire charge of explosive.

This corresponds to option (C).

💡 Quick Tip

In a real-world scenario and per the legal definition, a misfire includes both partial and total failures.

When dealing with a misfired hole, one must always assume the entire charge is still live and dangerous.

For exam purposes, understand the nuances, but choose the option that best fits the general understanding if the perfect answer isn't available.

194. The value of Bending factor of rope used in winding as per CMR-2017 is

- (A) 100
- (B) 120
- (C) 80
- (D) 150

Correct Answer: (A) 100

Solution:

Step 1: Understanding the Question:

The question asks for the minimum value of the "Bending Factor" for a winding rope as specified in the Coal Mines Regulations, 2017.

Step 2: Detailed Explanation:

The Bending Factor is the ratio of the diameter of the winding drum or sheave (D) to the diameter of the wire rope (d).

$$\text{Bending Factor} = \frac{D}{d}$$

This ratio is critical for the service life of a wire rope.

If the rope is bent around too small a drum, the bending stresses will be very high, leading to rapid fatigue and premature failure.

To prevent this, regulations specify a minimum value for this ratio.

According to Regulation 84(2) of the Coal Mines Regulations, 2017:

"The diameter of the winding drum or sheave shall not be less than **100** times the diameter of the winding rope, or of such size as may be approved in writing by the Chief Inspector."

Therefore, the minimum value for the bending factor (D/d ratio) is 100.

Step 3: Final Answer:

The value of the Bending factor of rope used in winding as per CMR-2017 is 100.

This corresponds to option (A).

💡 Quick Tip

This is another key number to memorize from the mining regulations.

The D/d ratio is a fundamental concept in wire rope engineering.

A larger ratio means less bending stress and a longer fatigue life for the rope.

100 is the standard minimum for many hoisting applications.

195. The process of setting goals and establishing guide lines to fulfil them is called

- (A) Planning
- (B) Supervising
- (C) Evaluating
- (D) Managing

Correct Answer: (A) Planning

Solution:

Step 1: Understanding the Question:

The question asks for the management term that defines the process of setting goals and creating a roadmap to achieve them.

Step 2: Detailed Explanation:

This is a direct definition of one of the core functions of management:

- **Planning:** This is the fundamental management function that involves deciding in advance what to do, when to do it, how to do it, and who is to do it.

It bridges the gap from where an organization is to where it wants to be.

It involves **setting goals** and **establishing strategies and guidelines** to achieve those goals.

This perfectly matches the question's description.

- **Supervising (or Directing):** This involves guiding, leading, and overseeing the work of subordinates to ensure they are working towards the established goals.

- **Evaluating (or Controlling):** This involves monitoring performance, comparing it with goals, and taking corrective action.

- **Managing:** This is the overall term that encompasses all functions, including planning, organizing, supervising, and controlling.

Step 3: Final Answer:

The process of setting goals and establishing guidelines to fulfil them is called Planning. This corresponds to option (A).

 Quick Tip

Remember the four basic functions of management:

1. **Planning** (Deciding the future) 2. **Organizing** (Arranging resources) 3. **Leading/Directing** (Guiding the people) 4. **Controlling** (Checking the results)

196. In CPM method of network analysis the Critical activities are associated with

- (A) Maximum Float
- (B) Minimum Float
- (C) Negative Float
- (D) Zero Float

Correct Answer: (D) Zero Float

Solution:

Step 1: Understanding the Question:

The question asks about a defining characteristic of "critical activities" within the Critical Path Method (CPM) of project management.

Step 2: Detailed Explanation:

- **Float (or Slack):** In network analysis, float is the amount of time that an activity can be delayed without causing a delay to the overall project completion time.

- **Critical Path:** This is the sequence of activities in a project network that determines the longest path from start to finish.

The length of the critical path is the shortest possible duration for the entire project.

- **Critical Activities:** These are the activities that lie on the critical path.

Because these activities are on the longest path, they have no room for delay.

Any delay in a critical activity will directly delay the completion of the entire project.

Therefore, by definition, critical activities have a float of **zero**.

Non-critical activities have a positive float, meaning they can be delayed to some extent without affecting the project deadline.

Step 3: Final Answer:

Critical activities are associated with zero float.

This corresponds to option (D).

 Quick Tip

The name says it all: the activity is "critical" precisely because it has no "float" or leeway.

It must start and end on time to avoid delaying the project.

Critical Path = Longest Path = Path of Zero Float.

197. In a PERT network for an activity Pessimistic, Most likely, and Optimistic times are 8 days, 6 days and 4 days respectively. The expected duration of the activity is -----

- (A) 6 days
- (B) 2 days
- (C) 8 days
- (D) 9 days

Correct Answer: (A) 6 days

Solution:

Step 1: Understanding the Question:

We need to calculate the "expected duration" of a single activity using the PERT (Program Evaluation and Review Technique) three-point estimate.

Step 2: Key Formula or Approach:

The PERT formula for the expected time (t_e) of an activity is a weighted average of the three estimates:

$$t_e = \frac{t_o + 4t_m + t_p}{6}$$

where: - t_o = Optimistic time (best-case scenario) - t_m = Most likely time - t_p = Pessimistic time (worst-case scenario)

Step 3: Detailed Explanation:

From the question, we have: - Optimistic time, $t_o = 4$ days - Most likely time, $t_m = 6$ days - Pessimistic time, $t_p = 8$ days

Substitute these values into the formula:

$$\begin{aligned} t_e &= \frac{4 + (4 \times 6) + 8}{6} \\ t_e &= \frac{4 + 24 + 8}{6} \\ t_e &= \frac{36}{6} = 6 \text{ days} \end{aligned}$$

Step 4: Final Answer:

The expected duration of the activity is 6 days.

This corresponds to option (A).

💡 Quick Tip

Notice that the "most likely" time is given the most weight (a factor of 4) in the PERT formula.

This pulls the expected time towards the most probable outcome.

In this case, since the time estimates are symmetric around the most likely time ($6-2=4$, $6+2=8$), the expected time is equal to the most likely time.

198. The time by which activity completion time can be delayed without affecting start of succeeding activity

- (A) Duration
- (B) Free float
- (C) Total float
- (D) Interfering float

Correct Answer: (B) Free float

Solution:**Step 1: Understanding the Question:**

The question asks for the specific project management term that defines the amount of delay possible for an activity without impacting the very next activity in the sequence.

Step 2: Detailed Explanation:

Let's define the different types of float:

- **Total Float:** This is the maximum amount of time an activity can be delayed without delaying the entire project's completion date.

- **Free Float:** This is the amount of time an activity can be delayed without delaying the *early start* of any immediately following ("succeeding") activity.

This is a subset of Total Float.

It represents the "local" slack an activity has before it starts affecting its neighbours.

This perfectly matches the definition in the question.

- **Interfering Float:** This is the difference between Total Float and Free Float.

It is the amount of slack that, if used, will not delay the project but will delay the start of a subsequent activity.

Step 3: Final Answer:

The time by which an activity can be delayed without affecting the start of a succeeding activity is called Free Float.

This corresponds to option (B).

 Quick Tip

Think of the floats in terms of their impact:

- **Free Float:** Delay impacts **NO ONE** else. (You are "free" to use it).
- **Total Float:** Delay impacts **SOMEONE** else down the line, but not the final project deadline.
- Zero Float: Delay impacts **EVERYONE**, including the project deadline (Critical Activity).

199. Which one of the following is a principle of TQM ?

- (A) Product-Centred system
- (B) Intermittent improvement
- (C) Customer-focus
- (D) Decision made by top executives

Correct Answer: (C) Customer-focus

Solution:

Step 1: Understanding the Question:

The question asks to identify a core principle of Total Quality Management (TQM).

Step 2: Detailed Explanation:

Total Quality Management (TQM) is a management philosophy that involves all employees in an organization in a continual effort to improve the quality of products, services, and processes.

Some of its core principles are:

- **Customer Focus:** TQM organizations recognize that the customer determines the level of quality.

Understanding and satisfying customer needs is paramount.

This is a fundamental principle.

- **Total Employee Involvement:** All employees participate in working toward common goals.

TQM aims to empower employees, not have decisions made only by top executives.

- **Process-centered:** TQM focuses on the processes that lead to the final product or service.

- **Continuous Improvement (Kaizen):** TQM is a commitment to a continuous, ongoing process of improvement, not intermittent or one-off fixes.

Let's evaluate the options:

(A) Product-Centred system: TQM is customer-centered, not product-centered.

(B) Intermittent improvement: TQM is based on continuous improvement.

(C) Customer-focus: This is a core pillar of TQM.

(D) Decision made by top executives: TQM advocates for total employee involvement and decentralized decision-making.

Step 3: Final Answer:

Customer-focus is a key principle of TQM.

This corresponds to option (C).

💡 Quick Tip

The three words in "Total Quality Management" hint at the principles:

- **Total:** Everyone is involved (not just top executives).
- **Quality:** Defined by the customer (customer-focus).
- **Management:** Through systematic, process-centered, continuous improvement.

200. For a person working in Opencast mine the rate at which leave with wages calculated is

- (A) One day for every 15 days of work performed by him
- (B) One day for every 20 days of work performed by him
- (C) One day for every 10 days of work performed by him
- (D) One day for every 30 days of work performed by him.

Correct Answer: (B) One day for every 20 days of work performed by him

Solution:

Step 1: Understanding the Question:

This is a regulatory question asking for the rate at which a worker in an opencast mine earns paid leave.

Step 2: Detailed Explanation:

The provision for "Leave with wages" for mine workers in India is governed by Section 52 of the Mines Act, 1952.

The act specifies two different rates for earning leave, based on the location of work:

1. For a person employed **below ground**: He shall be entitled to one day of leave for every **15 days** of work performed.
2. For **any other person** (which includes all persons working on the surface and in opencast mines): He shall be entitled to one day of leave for every **20 days** of work performed.

Since the question specifies a person working in an opencast mine, the second rate applies.

Step 3: Final Answer:

A person working in an opencast mine earns leave at the rate of one day for every 20 days of work performed.

This corresponds to option (B).

💡 Quick Tip

Remember the two different rates for leave with wages:

- **Underground** = 1 for 15 (More leave, as the work is more arduous).
- **Surface / Opencast** = 1 for 20.