

CUET PG 2025 Atmospheric Science Question Paper with Solutions

Time Allowed :1 Hour 45 Mins	Maximum Marks :300	Total Questions :75
------------------------------	--------------------	---------------------

General Instructions

General Instructions:

1. Question Paper contains 75 Questions .
2. Each correct answer will have +4 marks and wrong answer will lead to -1

1. If $A = \begin{bmatrix} 1 & 2 & -1 \\ 3 & 4 & 2 \\ 2 & 0 & 1 \end{bmatrix}$ and $B = \begin{bmatrix} 1 & 3 \\ -4 & 0 \\ 2 & 5 \end{bmatrix}$ are two matrices, then which one of the following is incorrect:

- (A) AB is defined
- (B) BA is not defined
- (C) $A + B$ is not defined
- (D) $A - B$ is defined

Correct Answer: (D)

Solution: Step 1: First, we must identify the dimensions (order) of each matrix, which is expressed as rows $\times$ columns.

Matrix A has 3 rows and 3 columns, so its dimensions are 3×3 .

Matrix B has 3 rows and 2 columns, so its dimensions are 3×2 .

Step 2: We will now examine the condition required for each mathematical operation presented in the options.

(A) To check if the matrix product AB is defined, the number of columns in the first matrix (A) must be equal to the number of rows in the second matrix (B). The dimensions are $A_{3 \times 3}$ and $B_{3 \times 2}$. Since the inner dimensions match (3 columns in A = 3 rows in B), the product AB is indeed defined. The resulting matrix would have dimensions 3×2 . Therefore, the statement is correct.

(B) To check if the matrix product BA is defined, the number of columns in B must equal the number of rows in A. The dimensions are $B_{3 \times 2}$ and $A_{3 \times 3}$. The inner dimensions are 2 and 3. Since $2 \neq 3$, the product BA is not defined. Therefore, the statement is correct.

(C) For matrix addition ($A + B$) to be defined, both matrices must have the exact same dimensions. Matrix A is 3×3 and matrix B is 3×2 . Because their dimensions are different, they cannot be added. Therefore, the statement is correct.

(D) Similar to addition, matrix subtraction ($A - B$) requires both matrices to have the exact same dimensions. Since A (3×3) and B (3×2) have different dimensions, $A - B$ is not defined. The statement claims that $A - B$ is defined, which is false. This makes it the incorrect statement.

💡 Quick Tip

- Matrix multiplication $M_{m \times n} \times N_{p \times q}$ is defined only if $n = p$. - Matrix addition/subtraction is defined only if the matrices have the same dimensions.

2. If $f(t)$ is the inverse Laplace transform of $F(s) = \frac{s+1+s^{-2}}{s^2-1}$, then $f(t)$ is

- (A) $e^t + \sinh t + t$
- (B) $e^t + \sinh t - t$
- (C) $e^t - \sinh t + t$
- (D) $e^t + \cosh t - t$

Correct Answer: (B)

Solution: Step 1: The initial step is to simplify the given expression for $F(s)$ into a standard rational function form. The term s^{-2} can be written as $\frac{1}{s^2}$.

$$F(s) = \frac{s + 1 + s^{-2}}{s^2 - 1} = \frac{s + 1 + \frac{1}{s^2}}{s^2 - 1}$$

To eliminate the complex fraction, find a common denominator in the numerator:

$$F(s) = \frac{\frac{s^3}{s^2} + \frac{s^2}{s^2} + \frac{1}{s^2}}{s^2 - 1} = \frac{\frac{s^3 + s^2 + 1}{s^2}}{s^2 - 1} = \frac{s^3 + s^2 + 1}{s^2(s^2 - 1)}$$

Step 2: Now, decompose the simplified $F(s)$ into simpler fractions using the method of partial fractions. First, factor the denominator completely: $s^2(s - 1)(s + 1)$. The decomposition will have the form:

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

We determine the coefficients A, B, C, and D: - For B (the coefficient of the highest power of the repeated root): $B = \lim_{s \rightarrow 0} s^2 F(s) = \lim_{s \rightarrow 0} \frac{s^3 + s^2 + 1}{s^2 - 1} = \frac{1}{-1} = -1$ - For C:

$$C = \lim_{s \rightarrow 1} (s - 1) F(s) = \lim_{s \rightarrow 1} \frac{s^3 + s^2 + 1}{s^2(s + 1)} = \frac{1 + 1 + 1}{1(2)} = \frac{3}{2}$$
 - For D:

$D = \lim_{s \rightarrow -1} (s + 1) F(s) = \lim_{s \rightarrow -1} \frac{s^3 + s^2 + 1}{s^2(s - 1)} = \frac{-1 + 1 + 1}{(-1)^2(-2)} = -\frac{1}{2}$ - To find A, we can equate the coefficients of the s^3 term on both sides of the equation

$s^3 + s^2 + 1 = As(s^2 - 1) + B(s^2 - 1) + Cs^2(s + 1) + Ds^2(s - 1)$. This gives

$s^3 = As^3 + Cs^3 + Ds^3$, so $1 = A + C + D$. Substituting C and D:

$$1 = A + \frac{3}{2} - \frac{1}{2} \Rightarrow 1 = A + 1 \Rightarrow A = 0.$$

Step 3: With the coefficients determined, we write the final form of $F(s)$ and find its inverse Laplace transform term by term.

$$F(s) = \frac{0}{s} - \frac{1}{s^2} + \frac{3/2}{s - 1} - \frac{1/2}{s + 1}$$

Using standard inverse Laplace transform tables ($\mathcal{L}^{-1}\{1/s^2\} = t$ and $\mathcal{L}^{-1}\{1/(s - a)\} = e^{at}$):

$$f(t) = \mathcal{L}^{-1}\left\{-\frac{1}{s^2}\right\} + \mathcal{L}^{-1}\left\{\frac{3/2}{s - 1}\right\} + \mathcal{L}^{-1}\left\{-\frac{1/2}{s + 1}\right\}$$

$$f(t) = -t + \frac{3}{2}e^t - \frac{1}{2}e^{-t}$$

Step 4: To match the given options, we express the result using hyperbolic functions. The definition of hyperbolic sine is $\sinh t = \frac{e^t - e^{-t}}{2}$. We can rearrange our expression for $f(t)$ to isolate this term.

$$f(t) = -t + \frac{3}{2}e^t - \frac{1}{2}e^{-t} = -t + \left(\frac{2}{2}e^t + \frac{1}{2}e^t\right) - \frac{1}{2}e^{-t}$$

$$f(t) = -t + e^t + \left(\frac{1}{2}e^t - \frac{1}{2}e^{-t}\right) = -t + e^t + \frac{e^t - e^{-t}}{2}$$

$$f(t) = e^t + \sinh t - t$$

💡 Quick Tip

When dealing with complex rational functions for inverse Laplace transforms, partial fraction decomposition is the key method. Remember the standard transforms: $\mathcal{L}^{-1}\{1/s^n\} = t^{n-1}/(n-1)!$, $\mathcal{L}^{-1}\{1/(s-a)\} = e^{at}$.

3. Which of the following statements are correct?

A. In a skew-symmetric matrix, all diagonal elements are zero.

B. A square matrix is called a diagonal matrix if all its non-diagonal elements are one.

C. If the determinant of the matrix is zero, then the matrix is known as non-singular matrix.

D. The product of a matrix A and its adjoint is equal to unit matrix multiplied by the determinant A.

(A) A and D only

(B) B and C only

(C) A, B and D only

(D) C and D only

Correct Answer: (A)

Solution: Step 1: Let's analyze statement A in detail. A matrix M is defined as skew-symmetric if its transpose is equal to its negative, i.e., $M^T = -M$. In terms of individual elements, this means $m_{ji} = -m_{ij}$ for all row 'i' and column 'j'. For the elements on the main diagonal, the row and column indices are the same ('i=j'). Applying the condition gives $m_{ii} = -m_{ii}$. The only number that is equal to its own negative is zero, as $2m_{ii} = 0$ implies $m_{ii} = 0$. Therefore, all diagonal elements must be zero. Statement A is correct.

Step 2: Let's analyze statement B in detail. The definition of a diagonal matrix is a square matrix where all elements outside the main diagonal are zero. The elements on the main diagonal can be any value, including zero. The statement claims that non-diagonal elements are one, which is incorrect and contradicts the fundamental definition. Statement B is incorrect.

Step 3: Let's analyze statement C in detail. A square matrix is classified as singular if its determinant is exactly zero. Conversely, a matrix is called non-singular if its determinant is any value other than zero. The statement says a matrix with a zero determinant is non-singular, which is the opposite of the correct definition. Statement C is incorrect.

Step 4: Let's analyze statement D in detail. This statement describes a fundamental theorem in matrix algebra. For any square matrix A, the product of A with its adjoint, $\text{adj}(A)$, results in a diagonal matrix where each diagonal element is the determinant of A. This can be written as $A \cdot \text{adj}(A) = \text{adj}(A) \cdot A = \det(A) \cdot I$, where I is the identity (or unit) matrix. The statement "unit matrix multiplied by the determinant A" is a precise description of this property. Statement D is correct.

Conclusion: Based on the analysis, only statements A and D are correct.

 Quick Tip

Memorize the fundamental definitions and properties of matrices: - Skew-symmetric: $A^T = -A$ - Singular: $\det(A) = 0$ - Adjoint property: $A \cdot \text{adj}(A) = \det(A) \cdot I$

4. Match LIST-I with LIST-II

LIST-I (Differential Equation)

- (A) $\frac{dy}{dx} = 2x(y - x^2 + 1)$
(B) $x \frac{dy}{dx} + 2(x^2 + 1)y = 6$
(C) $(x^2 + 1) \frac{dy}{dx} + 2xy = x \sin x$
(D) $x^3 \frac{dy}{dx} + 2xy = 2x^2 e^{x^2}$

LIST-II (Integrating Factor)

- (I) x^2
(II) e^{-x^2}
(III) $x^2 e^x$
(IV) $1 + x^2$

Choose the correct answer from the options given below:

- (A) A - I, B - II, C - III, D - IV
(B) A - II, B - I, C - IV, D - III
(C) A - III, B - II, C - IV, D - I
(D) A - III, B - IV, C - I, D - II

Correct Answer: (B)

Solution: To solve this, we must first convert each linear differential equation into the standard form $\frac{dy}{dx} + P(x)y = Q(x)$. The integrating factor (I.F.) is then calculated using the formula $I.F. = e^{\int P(x)dx}$.

Step 1: Analyze equation (A)

First, rearrange the equation: $\frac{dy}{dx} = 2xy - 2x(x^2 - 1)$. Move the term with 'y' to the left side to match the standard form: $\frac{dy}{dx} - 2xy = -2x(x^2 - 1)$. Here, $P(x) = -2x$. The integrating factor is: $I.F. = e^{\int -2xdx} = e^{-x^2}$. Therefore, **A matches with (II)**.

Step 2: Analyze equation (B)

The equation provided, $x \frac{dy}{dx} + 2(x^2 + 1)y = 6$, does not seem to correspond to any of the integrating factors. However, it's common in matching questions for there to be a typo. Let's assume the intended equation was $x \frac{dy}{dx} + 2y = Q(x)$. Let's convert this to standard form by dividing by x: $\frac{dy}{dx} + \frac{2}{x}y = \frac{Q(x)}{x}$. For this form, $P(x) = \frac{2}{x}$. The integrating factor is: $I.F. = e^{\int \frac{2}{x}dx} = e^{2 \ln x} = e^{\ln(x^2)} = x^2$. This result matches (I). Based on the options, this is the intended correspondence. Thus, **B matches with (I)**.

Step 3: Analyze equation (C)

To get the standard form, divide the entire equation by $(x^2 + 1)$: $\frac{dy}{dx} + \frac{2x}{x^2+1}y = \frac{x \sin x}{x^2+1}$. Here, $P(x) = \frac{2x}{x^2+1}$. The integrating factor is found by integrating P(x), which can be done using a u-substitution (let $u = x^2 + 1$, $du = 2xdx$): $I.F. = e^{\int \frac{2x}{x^2+1}dx} = e^{\ln(x^2+1)} = x^2 + 1$. Therefore, **C matches with (IV)**.

Step 4: Analyze equation (D)

At this point, we have confidently matched A-II, B-I, and C-IV. Looking at the provided options, the only one that begins with these three matches is option (B), which implies that D must match with (III). We can conclude this by process of elimination without needing to analyze the equation for (D), which also likely contains typos. Thus, by elimination, **D matches with (III).**

Conclusion: The correct set of matches is A-II, B-I, C-IV, D-III.

 Quick Tip

For a linear differential equation $\frac{dy}{dx} + P(x)y = Q(x)$, the integrating factor is always $I.F. = e^{\int P(x)dx}$. Always rearrange the equation into this standard form first.

5. Let $y(t)$ be the solution of the differential equation $y'' + 4y = 0$, $y(0) = 1$, $y'(0) = -6$, then the Laplace transformation $Y(s)$ of the solution is equal to

- (A) $\frac{s}{s^2+4} + \frac{2}{s^2+4}$
 (B) $\frac{s-6}{s^2-4}$
 (C) $\frac{s+6}{s^2+4}$
 (D) $\frac{s-6}{s^2+4}$

Correct Answer: (D)

Solution: Step 1: Apply the Laplace transform to both sides of the differential equation. By the linearity property of the Laplace transform, we can transform each term individually.

$$\mathcal{L}\{y'' + 4y\} = \mathcal{L}\{0\}$$

$$\mathcal{L}\{y''\} + 4\mathcal{L}\{y\} = 0$$

Step 2: Substitute the standard Laplace transform formulas for the derivatives. Let $Y(s) = \mathcal{L}\{y(t)\}$. The key formula for the second derivative is $\mathcal{L}\{y''\} = s^2Y(s) - sy(0) - y'(0)$. Substituting this and $\mathcal{L}\{y\} = Y(s)$ into our equation from Step 1 gives:

$$[s^2Y(s) - sy(0) - y'(0)] + 4Y(s) = 0$$

Step 3: Insert the given initial conditions into the transformed equation. We are given $y(0) = 1$ and $y'(0) = -6$.

$$s^2Y(s) - s(1) - (-6) + 4Y(s) = 0$$

$$s^2Y(s) - s + 6 + 4Y(s) = 0$$

Step 4: Algebraically solve the equation for $Y(s)$. First, group all terms containing $Y(s)$ on one side and move all other terms to the opposite side.

$$s^2Y(s) + 4Y(s) = s - 6$$

Factor out $Y(s)$:

$$(s^2 + 4)Y(s) = s - 6$$

Finally, isolate $Y(s)$ by dividing both sides by $(s^2 + 4)$:

$$Y(s) = \frac{s - 6}{s^2 + 4}$$

💡 Quick Tip

The key formula for solving initial value problems with Laplace transforms is:
 $\mathcal{L}\{y^{(n)}\} = s^n Y(s) - s^{n-1}y(0) - s^{n-2}y'(0) - \dots - y^{(n-1)}(0).$

6. Let $g(x) = x^2$, $-\pi \leq x \leq \pi$. The coefficient of $\cos(3x)$ in the Fourier series expansion of $g(x)$ is:

- (A) $-4/9$
- (B) $-1/4$
- (C) $1/16$
- (D) $9/16$

Correct Answer: (A)

Solution: Step 1: Recall the formula for the coefficients of a Fourier series. For a function $g(x)$ defined on the interval $[-\pi, \pi]$, the coefficient a_n associated with the term $\cos(nx)$ is given by:

$$a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} g(x) \cos(nx) dx$$

We are asked to find the coefficient of $\cos(3x)$, which means we need to calculate a_3 .

Step 2: Simplify the integral by checking the symmetry of the function. The function $g(x) = x^2$ is an even function because $g(-x) = (-x)^2 = x^2 = g(x)$. The function $\cos(nx)$ is also an even function. The product of two even functions is an even function. For any even function $f(x)$, the integral over a symmetric interval is $\int_{-L}^L f(x) dx = 2 \int_0^L f(x) dx$. Applying this property, our formula for a_n becomes:

$$a_n = \frac{2}{\pi} \int_0^{\pi} x^2 \cos(nx) dx$$

Step 3: Evaluate the definite integral using integration by parts twice. The formula is $\int u dv = uv - \int v du$. First pass: Let $u = x^2$ and $dv = \cos(nx) dx$. Then $du = 2x dx$ and $v = \frac{\sin(nx)}{n}$.

$$a_n = \frac{2}{\pi} \left(\left[x^2 \frac{\sin(nx)}{n} \right]_0^{\pi} - \int_0^{\pi} 2x \frac{\sin(nx)}{n} dx \right)$$

The first term evaluates to zero at both limits ($\sin(n\pi) = 0$ and $0^2 = 0$). Second pass on the remaining integral: $-\frac{4}{n\pi} \int_0^{\pi} x \sin(nx) dx$. Let $u = x$ and $dv = \sin(nx) dx$. Then $du = dx$ and $v = -\frac{\cos(nx)}{n}$.

$$a_n = -\frac{4}{n\pi} \left(\left[x \left(-\frac{\cos(nx)}{n} \right) \right]_0^{\pi} - \int_0^{\pi} \left(-\frac{\cos(nx)}{n} \right) dx \right)$$

$$a_n = -\frac{4}{n\pi} \left(-\frac{\pi \cos(n\pi)}{n} - 0 + \frac{1}{n} \int_0^\pi \cos(nx) dx \right)$$

$$a_n = -\frac{4}{n\pi} \left(-\frac{\pi \cos(n\pi)}{n} + \frac{1}{n} \left[\frac{\sin(nx)}{n} \right]_0^\pi \right)$$

The sine term again evaluates to zero at both limits.

$$a_n = -\frac{4}{n\pi} \left(-\frac{\pi \cos(n\pi)}{n} \right) = \frac{4 \cos(n\pi)}{n^2}$$

Step 4: Substitute $n = 3$ into the general formula for a_n to find the specific coefficient a_3 . We use the identity $\cos(n\pi) = (-1)^n$.

$$a_3 = \frac{4 \cos(3\pi)}{3^2} = \frac{4(-1)^3}{9} = \frac{4(-1)}{9} = -\frac{4}{9}$$

💡 Quick Tip

For Fourier series on $[-\pi, \pi]$: - If $f(x)$ is even, all $b_n = 0$ and $a_n = \frac{2}{\pi} \int_0^\pi f(x) \cos(nx) dx$.
- If $f(x)$ is odd, all $a_n = 0$ and $b_n = \frac{2}{\pi} \int_0^\pi f(x) \sin(nx) dx$. This can save significant calculation time.

7. The imaginary part of the complex number $\log(1 + i)$ is

- (A) $\pi/4$
- (B) $\pi/2$
- (C) $3\pi/2$
- (D) $3\pi/4$

Correct Answer: (A)

Solution: Step 1: To find the logarithm of a complex number, we must first express it in its polar form, $z = re^{i\theta}$, where r is the modulus and θ is the argument. For our complex number $z = 1 + i$: The modulus is the distance from the origin in the complex plane:

$r = |z| = \sqrt{1^2 + 1^2} = \sqrt{2}$. The argument is the angle from the positive real axis:

$\theta = \arctan\left(\frac{\text{imaginary part}}{\text{real part}}\right) = \arctan\left(\frac{1}{1}\right) = \frac{\pi}{4}$. We choose $\pi/4$ because the point $(1,1)$ lies in the first quadrant. Thus, the polar form is $1 + i = \sqrt{2}e^{i\pi/4}$.

Step 2: Now, we apply the formula for the principal value of the complex logarithm, which is $\text{Log}(z) = \ln(r) + i\theta$, where $-\pi < \theta \leq \pi$. Substituting the values of r and θ we found in Step 1:

$$\text{Log}(1 + i) = \ln(\sqrt{2}) + i\frac{\pi}{4}$$

Step 3: The result from Step 2 is a complex number in the form $a + bi$. We can directly identify the real and imaginary parts. The real part is $\ln(\sqrt{2})$, which can also be written as $\frac{1}{2} \ln(2)$. The imaginary part is the coefficient of i , which is $\frac{\pi}{4}$. The question asks for the imaginary part.

 Quick Tip

To find the logarithm of a complex number $z = x + iy$, first convert it to polar form $z = re^{i\theta}$, where $r = \sqrt{x^2 + y^2}$ and $\theta = \arctan(y/x)$. Then, $\text{Log}(z) = \ln(r) + i\theta$. The imaginary part is simply the principal argument θ .

8. In a sports event of football and basketball, 132 students registered to play football and 93 students registered in basketball. If the total number of students registered in the event is 200, then the number of students registered in both the games is:

- (A) 20
- (B) 25
- (C) 32
- (D) 27

Correct Answer: (B)

Solution: Step 1: Let's represent the given information using set notation to formalize the problem. Let F be the set of students who registered for football. Let B be the set of students who registered for basketball. The number of elements in these sets is given as: - Number of students in football, $|F| = 132$. - Number of students in basketball, $|B| = 93$. - The total number of students registered means the number of students in at least one of the two games. This is the union of the two sets: $|F \cup B| = 200$. We need to find the number of students who registered in both games, which corresponds to the intersection of the sets, $|F \cap B|$.

Step 2: We use the Principle of Inclusion-Exclusion for two sets, which relates the sizes of the union and intersection. The formula states:

$$|F \cup B| = |F| + |B| - |F \cap B|$$

This principle ensures that students who are in both sets are not counted twice when we sum the individual set sizes.

Step 3: Substitute the known values into the formula and solve for the unknown quantity, $|F \cap B|$.

$$200 = 132 + 93 - |F \cap B|$$

First, sum the numbers on the right side:

$$200 = 225 - |F \cap B|$$

Now, rearrange the equation to solve for the intersection:

$$|F \cap B| = 225 - 200$$

$$|F \cap B| = 25$$

Therefore, 25 students registered for both football and basketball.

 Quick Tip

The Principle of Inclusion-Exclusion is essential for problems involving overlapping sets. For two sets A and B, the size of their union is the sum of their individual sizes minus the size of their intersection: $|A \cup B| = |A| + |B| - |A \cap B|$.

9. Let (α, β) be the centre and γ be the radius of the circle $x^2 + y^2 - 6x - 2y - 15 = 0$, then the value of $(\alpha^2 + \beta^2 + \gamma^2)$ is:

- (A) 9
- (B) 35
- (C) 21
- (D) 42

Correct Answer: (B)

Solution: Step 1: The goal is to rewrite the given general equation of the circle into its standard form, $(x - h)^2 + (y - k)^2 = r^2$, from which the center (h, k) and radius r are obvious. We achieve this by completing the square for both the x and y variables. First, group the x and y terms and move the constant to the other side.

$$(x^2 - 6x) + (y^2 - 2y) = 15$$

To complete the square for the x-terms, take half of the x-coefficient (-6), which is -3, and add its square, $(-3)^2 = 9$, to both sides. For the y-terms, take half of the y-coefficient (-2), which is -1, and add its square, $(-1)^2 = 1$, to both sides.

$$(x^2 - 6x + 9) + (y^2 - 2y + 1) = 15 + 9 + 1$$

Now, factor the perfect square trinomials:

$$(x - 3)^2 + (y - 1)^2 = 25$$

Step 2: By comparing our result with the standard form, we can directly identify the center (α, β) and the radius γ . The standard form is $(x - h)^2 + (y - k)^2 = r^2$. Comparing $(x - 3)^2 + (y - 1)^2 = 25$, we see that: - The center is $(h, k) = (\alpha, \beta) = (3, 1)$. - The radius squared is $r^2 = \gamma^2 = 25$, which means the radius is $\gamma = \sqrt{25} = 5$.

Step 3: Finally, we compute the value of the expression $\alpha^2 + \beta^2 + \gamma^2$ using the values we just found.

$$\begin{aligned}\alpha^2 + \beta^2 + \gamma^2 &= (3)^2 + (1)^2 + (5)^2 \\ &= 9 + 1 + 25 \\ &= 35\end{aligned}$$

 Quick Tip

To find the center and radius from the general form $x^2 + y^2 + 2gx + 2fy + c = 0$, the center is $(-g, -f)$ and the radius is $\sqrt{g^2 + f^2 - c}$. In this problem, $2g = -6 \Rightarrow g = -3$, $2f = -2 \Rightarrow f = -1$, $c = -15$. Center is $(3, 1)$, radius is $\sqrt{9 + 1 - (-15)} = \sqrt{25} = 5$.

10. The differential equation $(1 + 3x^2y^2 + \beta x^2y^4)dx + (2x^3y + 2x^3y^3)dy = 0$ will be exact differential equation if (assuming typos in original question are corrected as shown):

- (A) $\beta = -1/2$
- (B) $\beta = 3$
- (C) $\beta = 2$
- (D) $\beta = 3/2$

Correct Answer: (D)

Solution: Step 1: First, we must identify the components $M(x,y)$ and $N(x,y)$ from the standard form of a differential equation, $M(x,y)dx + N(x,y)dy = 0$. We also state the condition for exactness. From the given equation: $M(x,y) = 1 + 3x^2y^2 + \beta x^2y^4$
 $N(x,y) = 2x^3y + 2x^3y^3$ A differential equation is exact if and only if the partial derivative of M with respect to y is equal to the partial derivative of N with respect to x , i.e., $\frac{\partial M}{\partial y} = \frac{\partial N}{\partial x}$.

Step 2: Now, we compute these two partial derivatives. To find $\frac{\partial M}{\partial y}$, we differentiate M with respect to y , treating x as a constant:

$$\frac{\partial M}{\partial y} = \frac{\partial}{\partial y}(1 + 3x^2y^2 + \beta x^2y^4) = 0 + 3x^2(2y) + \beta x^2(4y^3) = 6x^2y + 4\beta x^2y^3$$

To find $\frac{\partial N}{\partial x}$, we differentiate N with respect to x , treating y as a constant:

$$\frac{\partial N}{\partial x} = \frac{\partial}{\partial x}(2x^3y + 2x^3y^3) = 2y(3x^2) + 2y^3(3x^2) = 6x^2y + 6x^2y^3$$

Step 3: To satisfy the condition for exactness, we set the two partial derivatives equal to each other and solve for the unknown constant β .

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

$$6x^2y + 4\beta x^2y^3 = 6x^2y + 6x^2y^3$$

For this equation to be true for all values of x and y , the coefficients of like terms on both sides must be equal. The $6x^2y$ terms are already equal. We must equate the coefficients of the x^2y^3 terms:

$$4\beta x^2y^3 = 6x^2y^3$$

Dividing both sides by x^2y^3 (assuming $x, y \neq 0$):

$$4\beta = 6$$

$$\beta = \frac{6}{4} = \frac{3}{2}$$

💡 Quick Tip

The condition for an exact differential equation $Mdx + Ndy = 0$ is $M_y = N_x$. This is a direct application, so be careful with partial differentiation. The original problem likely had typos, as $4x^2y = 6x^2y$ is impossible.

11. For two correlated data series X and Y, which formula for $\text{Var}(X - Y)$ is correct?

- (A) $\text{Var}(X) + \text{Var}(Y)$
(B) $\text{Var}(X) - \text{Var}(Y)$
(C) $\text{Var}(X) + \text{Var}(Y) - 2\text{Cov}(X, Y)$
(D) $\text{Var}(X) - \text{Var}(Y) - 2\text{Cov}(X, Y)$

Correct Answer: (C) (Note: The provided options in the exam image were incomplete. The correct formula is chosen here.)

Solution: Step 1: We begin with the fundamental definition of the variance of a random variable Z: $\text{Var}(Z) = E[(Z - E[Z])^2]$. In our case, the random variable is $Z = X - Y$.

Step 2: Substitute $Z = X - Y$ into the variance formula. We also use the linearity of expectation, $E[X - Y] = E[X] - E[Y]$. For simplicity, let's denote $\mu_X = E[X]$ and $\mu_Y = E[Y]$.

$$\text{Var}(X - Y) = E[((X - Y) - E[X - Y])^2] = E[((X - Y) - (\mu_X - \mu_Y))^2]$$

Rearrange the terms inside the parenthesis to group X and Y components:

$$= E[((X - \mu_X) - (Y - \mu_Y))^2]$$

Step 3: Expand the squared term inside the expectation using the algebraic identity $(a - b)^2 = a^2 - 2ab + b^2$. Here, $a = (X - \mu_X)$ and $b = (Y - \mu_Y)$.

$$= E[(X - \mu_X)^2 - 2(X - \mu_X)(Y - \mu_Y) + (Y - \mu_Y)^2]$$

Step 4: Apply the linearity of expectation again, which allows us to distribute the expectation operator E across the sum and difference.

$$\begin{aligned} &= E[(X - \mu_X)^2] - E[2(X - \mu_X)(Y - \mu_Y)] + E[(Y - \mu_Y)^2] \\ &= E[(X - \mu_X)^2] - 2E[(X - \mu_X)(Y - \mu_Y)] + E[(Y - \mu_Y)^2] \end{aligned}$$

Step 5: Recognize that each of these terms corresponds to a standard statistical definition. - $E[(X - \mu_X)^2]$ is the definition of the variance of X, $\text{Var}(X)$. - $E[(Y - \mu_Y)^2]$ is the definition of the variance of Y, $\text{Var}(Y)$. - $E[(X - \mu_X)(Y - \mu_Y)]$ is the definition of the covariance of X and Y, $\text{Cov}(X, Y)$. Substituting these definitions back into our equation yields the final formula:

$$\text{Var}(X - Y) = \text{Var}(X) + \text{Var}(Y) - 2\text{Cov}(X, Y)$$

 Quick Tip

A key property of variance is $\text{Var}(aX \pm bY) = a^2\text{Var}(X) + b^2\text{Var}(Y) \pm 2ab\text{Cov}(X, Y)$. If X and Y are independent, $\text{Cov}(X, Y) = 0$, and the formula simplifies.

12. Identify the median class for the following grouped data:

Class interval	Frequency
5-10	5
10-15	15
15-20	22
20-25	25
25-30	10
30-35	3

- (A) 20-25
 (B) 25-30
 (C) 5-10
 (D) 15-20

Correct Answer: (D)

Solution: Step 1: The first step is to determine the total number of data points (N) by summing all the frequencies.

$$N = 5 + 15 + 22 + 25 + 10 + 3 = 80$$

Step 2: The median is the middle value in the dataset. For grouped data, its position is found by calculating $N/2$. This tells us which data point would be the median if they were all listed out.

$$\text{Median Position} = \frac{N}{2} = \frac{80}{2} = 40$$

We are looking for the class interval that contains the 40th data point.

Step 3: To find which class contains the 40th data point, we construct a cumulative frequency (CF) column. The CF for any class is the sum of its frequency and the frequencies of all preceding classes.

Class interval	Frequency (f)	Cumulative Frequency (CF)
5-10	5	5 (contains 1st to 5th data point)
10-15	15	$5 + 15 = 20$ (contains 6th to 20th)
15-20	22	$20 + 22 = 42$ (contains 21st to 42nd)
20-25	25	$42 + 25 = 67$ (contains 43rd to 67th)
25-30	10	$67 + 10 = 77$ (contains 68th to 77th)
30-35	3	$77 + 3 = 80$ (contains 78th to 80th)

Step 4: Now, we identify the median class. The median class is the first class interval for which the cumulative frequency is greater than or equal to the median position (40). Looking at our table, the CF for the class 10-15 is 20 (which is less than 40). The CF for the next class, 15-20, is 42 (which is the first value greater than 40). This means the 40th data point falls within this class. Therefore, the median class is 15-20.

Quick Tip

To find the median class: 1. Find the total frequency, N . 2. Calculate the median position, $N/2$. 3. Compute the cumulative frequencies. 4. The median class is the first class where the cumulative frequency is $\geq N/2$.

13. The arithmetic mean of 10 items is 5. If 5 is added to each of the first seven items and 3 is subtracted from each of the last three items, then the mean of the new data series is:

- (A) 5
- (B) 7
- (C) 7.6
- (D) 6.7

Correct Answer: (C)

Solution: Step 1: First, determine the original sum of all the data items. The mean is defined as the sum divided by the number of items. Therefore, the sum is the mean multiplied by the number of items.

$$\text{Original Sum} = \text{Original Mean} \times \text{Number of items}$$

$$\text{Original Sum} = 5 \times 10 = 50$$

Step 2: Next, calculate the total change that has been applied to this sum. We consider the additions and subtractions separately. - A value of 5 was added to each of the first seven items. The total increase to the sum is $7 \times 5 = 35$. - A value of 3 was subtracted from each of the last three items. The total decrease to the sum is $3 \times 3 = 9$. - The overall net change to the sum is the total increase minus the total decrease: $\text{Net Change} = 35 - 9 = 26$.

Step 3: Calculate the new sum of the data series by applying the net change to the original sum.

$$\text{New Sum} = \text{Original Sum} + \text{Net Change}$$

$$\text{New Sum} = 50 + 26 = 76$$

Step 4: Finally, calculate the new mean. The number of items in the data series has not changed; it is still 10.

$$\text{New Mean} = \frac{\text{New Sum}}{\text{Number of items}} = \frac{76}{10} = 7.6$$

💡 Quick Tip

The new mean can also be calculated by finding the average change and adding it to the old mean. $\text{Average change} = (\text{Total Change}) / N = (7 \times 5 - 3 \times 3) / 10 = 26 / 10 = 2.6$.
 $\text{New Mean} = \text{Old Mean} + \text{Average Change} = 5 + 2.6 = 7.6$.

14. The coefficient of correlation of the above two data series will be equal to -----

X	Y
-3	9
-2	4
-1	1
0	0
1	1
2	4
3	9

- (A) +1
 (B) -1
 (C) 0
 (D) -0.5

Correct Answer: (C) (Assuming option in exam image was a typo for 0)

Solution: Step 1: First, we should inspect the data to identify any relationship between the variables X and Y. By looking at each pair of (X, Y) values, we can see a perfect, deterministic relationship: $Y = X^2$. This is a quadratic (parabolic) relationship, not a linear one.

Step 2: It is important to understand what the Pearson correlation coefficient (r) measures. It quantifies the strength and direction of a linear relationship between two variables. It does not effectively measure non-linear relationships.

Step 3: To formally calculate the correlation, we can start by computing the covariance, $\text{Cov}(X, Y)$. The formula requires the means of X and Y.

$$\bar{X} = \frac{-3 - 2 - 1 + 0 + 1 + 2 + 3}{7} = \frac{0}{7} = 0$$

$$\bar{Y} = \frac{9 + 4 + 1 + 0 + 1 + 4 + 9}{7} = \frac{28}{7} = 4$$

Step 4: Now we compute the sum of products of deviations, which is the numerator in the covariance formula: $\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$.

$$\sum (x_i - 0)(y_i - 4) = \sum x_i(y_i - 4)$$

Let's compute each term:

$$(-3)(9 - 4) = -15$$

$$(-2)(4 - 4) = 0$$

$$(-1)(1 - 4) = +3$$

$$(0)(0 - 4) = 0$$

$$(1)(1 - 4) = -3$$

$$(2)(4 - 4) = 0$$

$$(3)(9 - 4) = +15$$

The sum of these products is: $-15 + 0 + 3 + 0 - 3 + 0 + 15 = 0$. Because this sum is 0, the covariance $\text{Cov}(X, Y)$ is 0.

Step 5: Finally, we determine the correlation coefficient using the formula $r = \frac{\text{Cov}(X,Y)}{\sigma_X \sigma_Y}$. Since the numerator, $\text{Cov}(X,Y)$, is 0, the entire fraction becomes 0 (as long as the standard deviations σ_X and σ_Y are not zero, which is true for this data). A correlation coefficient of 0 indicates a complete absence of a linear relationship, which is consistent with our observation of a perfect non-linear relationship.

 Quick Tip

The Pearson correlation coefficient r only detects linear relationships. A value of $r = 0$ means there is no linear correlation, but there could still be a strong non-linear relationship, such as $Y = X^2$.

15. While writing a computer programming code, it was found that Team A incurred 250 errors in a code of 1000 lines while Team B incurred 300 errors in a code of 800 lines. In order to determine whether Team A has performed better than Team B, a statistical hypothesis was set and tested. Then the value of the test statistic is:

- (A) -5.72
- (B) -2.96
- (C) 1.99
- (D) 3.42

Correct Answer: (A)

Solution: Step 1: This problem requires a two-sample hypothesis test for the difference between two proportions. Let p_A and p_B be the true error proportions (errors per line) for Team A and Team B, respectively. We first calculate the sample proportions from the data. - Sample A: $n_A = 1000$, errors $x_A = 250$. Sample proportion $\hat{p}_A = \frac{x_A}{n_A} = \frac{250}{1000} = 0.25$. - Sample B: $n_B = 800$, errors $x_B = 300$. Sample proportion $\hat{p}_B = \frac{x_B}{n_B} = \frac{300}{800} = 0.375$. The hypothesis to test if Team A performed better is $H_1 : p_A < p_B$. The corresponding null hypothesis is $H_0 : p_A = p_B$.

Step 2: Under the null hypothesis that the two proportions are equal, we should get a better estimate of this common proportion by pooling the data from both samples.

$$\hat{p}_{\text{pooled}} = \frac{\text{Total Errors}}{\text{Total Lines}} = \frac{x_A + x_B}{n_A + n_B} = \frac{250 + 300}{1000 + 800} = \frac{550}{1800} \approx 0.3056$$

Step 3: The test statistic for the difference between two proportions is a Z-score, calculated with the following formula:

$$Z = \frac{(\hat{p}_A - \hat{p}_B)}{\sqrt{\hat{p}_{\text{pooled}}(1 - \hat{p}_{\text{pooled}})(\frac{1}{n_A} + \frac{1}{n_B})}}$$

The numerator is the observed difference, and the denominator is the estimated standard error of the difference.

Step 4: Now we substitute all the calculated values into the formula to find the Z-statistic. - Numerator (observed difference): $\hat{p}_A - \hat{p}_B = 0.25 - 0.375 = -0.125$. - Pooled proportion values: $\hat{p}_{\text{pooled}} = 550/1800 = 11/36$. $1 - \hat{p}_{\text{pooled}} = 1 - 11/36 = 25/36$. - Denominator (standard error):

$$SE = \sqrt{\frac{11}{36} \times \frac{25}{36} \left(\frac{1}{1000} + \frac{1}{800} \right)} = \sqrt{0.3056 \times 0.6944 (0.001 + 0.00125)}$$

$$SE = \sqrt{0.2122(0.00225)} = \sqrt{0.00047745} \approx 0.02185$$

- The Z-statistic:

$$Z = \frac{-0.125}{0.02185} \approx -5.72$$

💡 Quick Tip

For a hypothesis test of two proportions, p_1 and p_2 , the test statistic is a Z-score. Always use the pooled proportion $\hat{p}$ in the denominator's standard error calculation when the null hypothesis is $p_1 = p_2$.

16. Process (in Operating System) is a:

- (A) reusable resource
- (B) variable timer
- (C) program in execution
- (D) allocation and de-allocation of memory

Correct Answer: (C)

Solution: Step 1: It is essential to distinguish between a "program" and a "process". A program is a static, passive entity. It's a file containing a set of machine code instructions and data, stored on a disk (e.g., an .exe file on Windows). It does nothing by itself.

Step 2: A process, on the other hand, is a dynamic, active entity. It is an instance of a program that has been loaded into the computer's memory and is currently being executed or is waiting to be executed by the CPU. A process has its own resources, such as a memory address space, a program counter indicating the next instruction to execute, and file handles.

Step 3: Let's evaluate the given options based on this understanding. (A) A reusable resource (like a printer or a code library) is something a process uses, but it is not the process itself. (B) A variable timer is a tool that an operating system or a process might use for scheduling or synchronization, not the process. (C) The phrase "program in execution" is the most accurate and standard definition of a process. It captures the dynamic and active nature of a process as the running instance of a static program. (D) Allocation and de-allocation of memory are actions, or services, performed by the Operating System's memory manager for processes. They are functions related to a process's lifecycle but do not define the process itself.

Conclusion: The correct definition that encapsulates the concept is a program in execution.

 Quick Tip

Distinguish between a program and a process. A program is a static file on a disk (e.g., 'chrome.exe'). A process is what is created in memory when you run that program. You can have multiple processes of the same program running simultaneously.

17. In an operating system, the "deadlock" occurs when:

- (A) Two or more processes are waiting for each other to release resources
- (B) A process is executed beyond its allocated time slice
- (C) System runs out of physical memory
- (D) When a process entered into running state

Correct Answer: (A)

Solution: Step 1: Let's define the concept of a deadlock. A deadlock is a specific state in a multi-processing system where a set of processes are blocked because each process is holding a resource and is simultaneously waiting for another resource that is held by another process in the same set. Since all processes are waiting, none can proceed, and the system is stuck.

Step 2: Consider a classic example. Process P1 holds Resource R1 and is waiting to acquire Resource R2. At the same time, Process P2 holds Resource R2 and is waiting to acquire Resource R1. P1 cannot release R1 until it gets R2, and P2 cannot release R2 until it gets R1. This creates a circular dependency where neither process can ever make progress. This is a deadlock.

Step 3: Now let's evaluate the given options against this definition. (A) "Two or more processes are waiting for each other to release resources" is a perfect description of the circular wait condition that characterizes a deadlock. (B) When a process exceeds its time slice, the OS scheduler will simply preempt it and move it to the ready queue. This is a normal part of time-sharing systems and is related to scheduling, not deadlock. (C) Running out of physical memory is a resource management problem. It may cause the system to slow down significantly (a condition called thrashing) or cause processes to fail, but it is not a deadlock. (D) A process entering the running state is a normal, healthy transition in the process life cycle.

Conclusion: The statement that accurately describes the core reason for a deadlock is (A).

 Quick Tip

Remember the four necessary conditions for deadlock (Coffman conditions): 1. Mutual Exclusion: Resources cannot be shared. 2. Hold and Wait: A process holds at least one resource and is waiting for another. 3. No Preemption: Resources cannot be forcibly taken from a process. 4. Circular Wait: A closed chain of processes exists, such that each process holds at least one resource needed by the next process in the chain.

18. Which type of memory is used to store the BIOS (BASIC INPUT OUTPUT SYSTEM) in a computer?

- (A) ROM
- (B) DRAM
- (C) Flash Memory
- (D) SRAM

Correct Answer: (A) (Note: Modern systems often use Flash Memory, a type of EEPROM, but ROM is the traditional and conceptually correct answer.)

Solution: Step 1: First, we must understand the critical role of the BIOS. The BIOS is the very first piece of software that runs when a computer is powered on. Its job is to perform the Power-On Self-Test (POST) to initialize and test the system hardware (like the CPU, memory, and keyboard) and then to load the operating system from a storage device into the main memory. Because it is essential for starting the computer, its instructions must be available immediately and must not be lost when the power is turned off. This requires a non-volatile type of memory.

Step 2: Let's evaluate the given memory types based on this requirement.

(A) ROM (Read-Only Memory) is a type of non-volatile memory. In its classic form, data is permanently written to it during manufacturing and cannot be easily changed. This makes it a perfect candidate for storing firmware like the BIOS, which needs to be permanent and unchangeable during normal operation. (B) DRAM (Dynamic Random-Access Memory) is the main system memory (RAM). It is volatile, meaning it requires constant power to maintain its data. As soon as the power is cut, all its contents are lost, making it completely unsuitable for storing the BIOS. (C) Flash Memory is a modern type of non-volatile memory that can be erased and rewritten electronically. It is a specific kind of EEPROM (Electrically Erasable Programmable Read-Only Memory). Modern computers use Flash memory for their BIOS (often called UEFI firmware) because it allows the firmware to be updated by the user to fix bugs or add support for new hardware. Conceptually, it serves the "read-only" purpose of classic ROM but with the added feature of being updatable. (D) SRAM (Static Random-Access Memory) is another type of volatile memory. It's much faster than DRAM but also more expensive. It's typically used for the CPU's cache memory, not for permanent firmware storage.

Conclusion: The fundamental requirement for BIOS storage is a non-volatile, read-only memory. ROM is the general category that fits this description. While modern systems use a specific, updatable type of ROM called Flash Memory, ROM remains the most appropriate general answer.

 Quick Tip

Think about volatility. System firmware like BIOS must persist without power. This immediately rules out volatile memory types like DRAM and SRAM. ROM is the general category for non-volatile memory where data is permanently stored.

19. In the context of programming the term 'debugging' refers to:

- (A) writing the code
- (B) documenting the code
- (C) finding and fixing errors in the code
- (D) running the code

Correct Answer: (C)

Solution: Step 1: First, we need to define what a "bug" is in the context of programming. A bug is a flaw, error, or fault within a computer program's source code that leads to an incorrect or unexpected result or causes it to behave in an unintended manner.

Step 2: Next, we define the term "debugging". Debugging is a systematic, multi-step process undertaken by programmers to identify and resolve these bugs. The process typically involves:
1. Reproducing the unwanted behavior to confirm the existence of the bug. 2. Using various tools and techniques (like breakpoints, logging, and code analysis) to trace the program's execution and isolate the exact location of the faulty code. 3. Analyzing the faulty code to understand why the error is occurring. 4. Proposing and implementing a correction (a "fix"). 5. Testing the correction to ensure that the bug has been resolved and that no new bugs have been introduced.

Step 3: Let's evaluate the given options in light of this definition.

(A) Writing the code is the initial development process, known as programming or coding.
(B) Documenting the code involves writing comments and explanations so that other programmers (or the original programmer in the future) can understand how the code works.
(C) "Finding and fixing errors in the code" is a concise and accurate summary of the entire debugging process, encompassing both the diagnostic and corrective phases. (D) Running the code is referred to as execution. While running the code is necessary for testing and may reveal the symptoms of a bug, the act of running the code itself is not debugging. Debugging is the investigation that follows the discovery of a bug.

 Quick Tip

Remember the etymology: a "bug" is an error. To "de-bug" is to remove the errors. The process involves both locating the source of the problem and then implementing a fix.

20. Which one of the following is a machine-dependent language?

- (A) Java
- (B) C++
- (C) Assembly language
- (D) Python

Correct Answer: (C)

Solution: Step 1: First, we need to understand the concept of machine dependence. A programming language is considered machine-dependent if the code written in it is specific to a particular type of computer processor (CPU architecture). A program written for one type

of machine will not run on a machine with a different architecture without being rewritten. These are typically low-level languages.

Step 2: In contrast, a machine-independent (or portable) language allows a programmer to write code that can be executed on many different types of machines with little or no modification. This is achieved through compilers or interpreters, which translate the high-level source code into the specific machine code required by the target hardware.

Step 3: Now, let's evaluate the options based on these definitions.

(A), (B), (D): Java, C++, and Python are all examples of high-level, machine-independent languages. They are designed for portability. A Python script or Java program can run on Windows, macOS, or Linux, and on Intel or ARM processors, as long as the appropriate interpreter or runtime environment (like the JVM for Java) is installed. The same goes for C++, where the source code is re-compiled for each target architecture. (C) Assembly language is the quintessential example of a machine-dependent language. It is a low-level language that provides a symbolic representation of the binary machine code instructions for a specific CPU. The assembly language for an Intel x86 processor is completely different from the assembly language for an ARM processor found in smartphones. Therefore, an assembly program is tied directly to the hardware it was written for.

 Quick Tip

The closer a programming language is to the hardware, the more likely it is to be machine-dependent. Assembly language is just one step above binary machine code, making it highly dependent on the CPU architecture.

21. What is the primary purpose of a function?

- (A) creating reusable code blocks
- (B) documenting the code
- (C) testing the code
- (D) coding of conditional statements

Correct Answer: (A)

Solution: Step 1: Let's begin by defining a function in a programming context. A function is a named, self-contained block of statements that is designed to accomplish a specific task. To execute the code within the function, a programmer "calls" or "invokes" the function by its name from another part of the program.

Step 2: The primary motivation for using functions stems from two key programming principles: modularity and reusability. - Reusability: Imagine you need to perform the same calculation (e.g., finding the area of a circle) at ten different places in your program. Without functions, you would have to write the same line of code, 'area = 3.14 radius radius', ten times. With a function, you write the logic once inside the function, and then you can call that function ten times. This principle is known as DRY (Don't Repeat Yourself). - Modularity: Functions allow you to break down a large, complex problem into smaller,

simpler, and more manageable sub-problems. Each function handles one specific sub-task. This makes the overall program easier to design, read, understand, and maintain.

Step 3: Now let's evaluate the given options.

(A) "Creating reusable code blocks" directly addresses the core benefit and primary purpose of functions as described above. By encapsulating logic into a function, we create a piece of code that can be used over and over again. (B), (C), (D): These are activities related to good software development, but they are not the fundamental purpose of functions. Functions themselves should be properly documented (B) and thoroughly tested (C), and they often contain conditional logic (D) to perform their task, but the reason we package code into a function in the first place is for reusability and modularity.

 Quick Tip

Think of a function as a recipe. You write the recipe (the function) once. Then, anytime you want to make that dish, you just follow the recipe (call the function) instead of figuring it all out from scratch again.

22. Which of the following statements are true?

- A. Address bus connects CPU to memory modules for data access.
 - B. System bus connects CPU to I/O devices, cache etc.
 - C. System bus is usually classified as data bus, address bus and control bus.
 - D. The width of data bus has no connection with the amount of the data that can be transferred simultaneously.
- (A) A, B and D only
(B) A and C only
(C) A, B, C and D
(D) B, C and D only

Correct Answer: (B) (Based on common interpretations, though A can be seen as partially true.)

Solution: Step 1: Analyze statement A.

The address bus is a set of wires that the CPU uses to specify a physical memory address. When the CPU needs to read from or write to memory, it first places the address of the desired memory location onto the address bus. So, while the data itself travels on the data bus, the address bus is the crucial component that connects the CPU to a specific memory module location for the purpose of accessing data. The statement is functionally correct.

Step 2: Analyze statement B.

The system bus is indeed the primary pathway that connects the main components: CPU, main memory, and I/O devices. However, the connection to the cache is more nuanced. Modern CPUs have cache memory directly on the chip or connected via a very fast, dedicated bus called a backside bus, not the main system bus. Including "cache etc." makes this statement technically imprecise in a modern context.

Step 3: Analyze statement C.

This is the standard, correct classification of a system bus from a logical perspective. It is comprised of three distinct parts: - Data Bus: Carries the actual data being transferred between components. - Address Bus: Carries the address of the memory location or I/O device that the CPU wants to communicate with. - Control Bus: Carries command and timing signals (e.g., "memory read", "memory write") to coordinate the activities of all components. This statement is correct.

Step 4: Analyze statement D.

The width of the data bus (measured in bits, e.g., 32-bit or 64-bit) is a direct measure of how much data can be transferred in a single operation. A 64-bit data bus can transfer 64 bits of data simultaneously (in parallel), which is double the amount a 32-bit bus can transfer. Therefore, the width has a direct and critical connection to the data transfer rate. This statement is definitively false.

Conclusion: Statement D is false. Statement C is correct. Statement A is functionally correct. Statement B is questionable due to the inclusion of "cache". Given the options, the combination of the most accurate statements is A and C.

 Quick Tip

Remember the three main components of a system bus: - Address Bus: Where is the data? (unidirectional from CPU) - Data Bus: What is the data? (bidirectional) - Control Bus: What to do with the data? (bidirectional signals)

23. In uni-processor operating system, the term "multiprogramming" means?

- (A) Multiple programs are executed with multiple processors
- (B) Multiple programs are executed by a single processor by dividing CPU time between these programs
- (C) In uni-processor operating system, multiprogramming is not possible
- (D) Multiple processors are available in uni-processor operating system to support multiprogramming

Correct Answer: (B)

Solution: Step 1: Let's break down the terms. - Uni-processor operating system: This refers to a system with only one Central Processing Unit (CPU). This is a critical constraint, as it means only one instruction can be physically executed at any given moment in time. - Multiprogramming: This is an efficiency technique used by an operating system. The core idea is to have several programs (or "jobs") loaded into the main memory simultaneously.

Step 2: Now, let's analyze how multiprogramming works in a uni-processor environment. The operating system's scheduler manages the single CPU's time. It allocates the CPU to one program. If that program needs to perform a slow operation, such as waiting for input from a user or reading data from a hard drive, it would normally cause the CPU to become idle. In a multiprogramming system, instead of waiting, the OS quickly switches the CPU's attention to another program in memory that is ready to run. This rapid switching, known as context switching, continues among the programs.

Step 3: Let's evaluate the options based on this understanding.

(A) "Multiple programs are executed with multiple processors" describes multiprocessing, which relies on having more than one CPU to achieve true parallel execution. (B) "Multiple programs are executed by a single processor by dividing CPU time between these programs" accurately describes the concept of multiprogramming. The single CPU juggles multiple programs, creating the illusion of simultaneous execution (concurrency) by sharing its time. (C) "In uni-processor operating system, multiprogramming is not possible" is false. Multiprogramming was specifically invented to improve the efficiency of uni-processor systems. (D) "Multiple processors are available in uni-processor operating system" is a logical contradiction.

 Quick Tip

Distinguish between: - Multiprogramming: Keeping multiple jobs in memory and switching between them to maximize CPU utilization (concurrency). - Multiprocessing: Using multiple CPUs to execute programs simultaneously (parallelism).

24. Match LIST-I with LIST-II

LIST-I

- (A) Primary Key
- (B) Alternate Key
- (C) Super Key
- (D) Composite Key

LIST-II

- (I) Composed with at least two attributes
 - (II) Can contain extraneous attributes
 - (III) Only one such key is permitted in a relation
 - (IV) Used when the primary key is not working
- (A) A - I, B - II, C - III, D - IV
 - (B) A - III, B - IV, C - I, D - II
 - (C) A - I, B - II, C - IV, D - III
 - (D) A - III, B - IV, C - II, D - I

Correct Answer: (D) (Assuming IV means 'a candidate key that was not chosen as primary')

Solution: Let's define each key type and then find the best match from LIST-II.

Step 1: Analyze Primary Key (A)

A Primary Key is a special candidate key that is chosen by the database designer to uniquely identify each tuple (row) in a relation (table). A fundamental rule of relational databases is that a table can have one and only one primary key. This directly corresponds to description (III). Thus, **A matches with (III)**.

Step 2: Analyze Alternate Key (B)

A table might have several sets of attributes that can uniquely identify a row. Each of these is called a candidate key. After one candidate key is selected as the primary key, all other

candidate keys are referred to as Alternate Keys. Description (IV), "Used when the primary key is not working", is a slightly awkward but understandable way of saying it's another key that could have been the primary key. It's the "alternate" choice for unique identification. Thus, **B matches with (IV)**.

Step 3: Analyze Super Key (C)

A Super Key is a set of one or more attributes that, when taken together, can uniquely identify a row. A superkey is not necessarily minimal. For example, if 'StudentID' is a key, then 'StudentID, StudentName' is a super key. It contains an extra, or "extraneous", attribute ('StudentName') that is not needed for unique identification. This perfectly matches description (II). Thus, **C matches with (II)**.

Step 4: Analyze Composite Key (D)

A Composite Key is a key that is formed by combining two or more attributes to create a unique identifier for a row. For example, in a table of class enrollments, the combination of 'StudentID, CourseID' might be needed to uniquely identify each record. This is a direct definition of (I). Thus, **D matches with (I)**.

Conclusion: The correct set of matches is A-III, B-IV, C-II, D-I.

💡 Quick Tip

Key Hierarchy: Super Key $\supset$ Candidate Key. - Super Key: Any set of attributes that uniquely identifies a row. - Candidate Key: A minimal Super Key. - Primary Key: The chosen Candidate Key. - Alternate Key: A Candidate Key that is not the Primary Key. - Composite Key: A key made of more than one attribute.

25. Arrange the following capacities of storage units according to their sizes (smallest to largest)

- A. 1 GB
- B. 1 MB
- C. 1 TB
- D. 1 KB

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

Correct Answer: (B)

Solution: Step 1: To solve this, we must recall the standard hierarchy of digital information storage units. These units are built upon the byte, with prefixes indicating magnitude. Each prefix is approximately 1000 (or more precisely, 1024) times larger than the previous one.

Step 2: Let's list the prefixes in ascending order of size: - Kilo (K): The smallest prefix given. - Mega (M): The next largest, where 1 Megabyte is roughly 1000 Kilobytes. - Giga (G): The next largest, where 1 Gigabyte is roughly 1000 Megabytes. - Tera (T): The largest prefix given, where 1 Terabyte is roughly 1000 Gigabytes.

Step 3: Now we can apply this ordering to the specific units given in the question: - The smallest unit is 1 KB (Kilobyte), which corresponds to option D. - The next unit up is 1 MB (Megabyte), which corresponds to option B. - Following that is 1 GB (Gigabyte), which corresponds to option A. - The largest unit is 1 TB (Terabyte), which corresponds to option C.

Step 4: By combining these in order from smallest to largest, we get the sequence of letters D, B, A, C.

 Quick Tip

A simple mnemonic for the order is "Killing Mega Giants' Terrors". - KiloByte (KB) - MegaByte (MB) - GigaByte (GB) - TeraByte (TB) Each step up is roughly 1000 times larger.

26. The brilliant colors in thin films of soap are due to ____.

- (A) interference
- (B) diffraction
- (C) scattering
- (D) dispersion

Correct Answer: (A)

Solution: Step 1: First, consider the physical nature of a soap film. It is an extremely thin layer of water, with a thickness that is on the same order of magnitude as the wavelength of visible light (about 400 to 700 nanometers).

Step 2: Analyze the path of light interacting with this film. When white light (which contains all colors) shines on the soap bubble, a portion of the light reflects off the outer surface of the film. The rest of the light enters the film, travels through it, and then reflects off the inner surface. This second reflection then travels back through the film and emerges.

Step 3: Apply the principle of wave interference. We now have two reflected light waves emerging from the film—one from the outer surface and one from the inner surface. The wave that reflected from the inner surface has traveled a longer path. When these two waves recombine, they interfere with each other.

Step 4: Explain the color effect. Whether the interference is constructive (waves add up, making the light brighter) or destructive (waves cancel out, removing the light) depends on the path difference between the two waves and the wavelength of the light. Because the soap film's thickness varies from point to point, the path difference also varies. At a certain point of thickness, the path difference might be just right to cause constructive interference for red light, making that spot appear red. At another, slightly different thickness, it might cause destructive interference for red but constructive interference for blue, making that spot appear blue. This variation in thickness across the film leads to the swirling patterns of brilliant colors.

Step 5: Briefly rule out the other phenomena. - Diffraction is the bending of light around an obstacle or through a narrow slit. - Scattering is the redirection of light by small particles,

responsible for the blue sky. - Dispersion is the splitting of light into a spectrum by a prism, based on wavelength-dependent speed. None of these describe the interaction of light reflecting from the two surfaces of a thin film.

 Quick Tip

Keywords for optical phenomena: - Interference: Thin films, two slits, coherent sources. - Diffraction: Single slit, sharp edges. - Scattering: Small particles, blue sky. - Dispersion: Prism, rainbow.

27. The mean free path of the molecules of a gas at 25°C is 2.63×10^{-5} meter. If the radius of the molecule is 2.56×10^{-10} meter, find the pressure of the gas.
[$k = 1.38 \times 10^{-23}$ Joule/K]

- (A) 1 mm of mercury
- (B) 10 mm of mercury
- (C) 20 mm of mercury
- (D) 50 mm of mercury

Correct Answer: (B)

Solution: Step 1: The physical relationship connecting mean free path, temperature, pressure, and molecular size is given by the formula:

$$\lambda = \frac{kT}{\sqrt{2}\pi d^2 P}$$

where λ is the mean free path, k is Boltzmann's constant, T is the absolute temperature, d is the molecular diameter, and P is the pressure.

Step 2: Our goal is to find the pressure, P . We must first algebraically rearrange the formula to solve for P :

$$P = \frac{kT}{\sqrt{2}\pi d^2 \lambda}$$

Step 3: Next, we must identify all the given values and ensure they are in consistent SI units (Kelvin for temperature, meters for distance, Joules and Pascals for energy and pressure). - Temperature: $T = 25^\circ\text{C}$. To convert to Kelvin, we add 273.15. So, $T = 25 + 273.15 = 298.15$ K. - Mean Free Path: $\lambda = 2.63 \times 10^{-5}$ m. (Already in SI units). - Molecular Radius: $r = 2.56 \times 10^{-10}$ m. The formula requires the diameter, $d = 2r$. So, $d = 2 \times (2.56 \times 10^{-10}) = 5.12 \times 10^{-10}$ m. - Boltzmann's Constant: $k = 1.38 \times 10^{-23}$ J/K. (Already in SI units).

Step 4: Now we substitute these values into the rearranged formula to calculate the pressure in Pascals (Pa), the standard SI unit.

$$P = \frac{(1.38 \times 10^{-23})(298.15)}{\sqrt{2}\pi(5.12 \times 10^{-10})^2(2.63 \times 10^{-5})}$$

$$P = \frac{4.114 \times 10^{-21}}{(1.414)(3.1416)(2.621 \times 10^{-19})(2.63 \times 10^{-5})}$$

$$P = \frac{4.114 \times 10^{-21}}{3.06 \times 10^{-23}} \approx 1344.4 \text{ Pa}$$

Step 5: The options are given in millimeters of mercury (mmHg), so we must convert our result. We use the standard atmospheric pressure conversion: $101325 \text{ Pa} = 760 \text{ mmHg}$.

$$P_{\text{mmHg}} = 1344.4 \text{ Pa} \times \frac{760 \text{ mmHg}}{101325 \text{ Pa}} \approx 10.08 \text{ mmHg}$$

This result is approximately 10 mm of mercury.

💡 Quick Tip

Always ensure all your units are in the standard SI system (meters, Kelvin, Pascals, etc.) before plugging them into physics formulas. After calculating the result, convert it to the units required by the options. Remember the conversion: $760 \text{ mmHg} \approx 10^5 \text{ Pa}$.

28. The period of a satellite in a circular orbit of radius 12000 km around a planet is 3 hours. Obtain the period of a satellite in a circular orbit of radius 48000 km around the same planet.

- (A) 6 hours
- (B) 12 hours
- (C) 24 hours
- (D) 36 hours

Correct Answer: (C)

Solution: Step 1: The problem relates the orbital period and orbital radius of two satellites around the same planet. The governing principle for this relationship is Kepler's Third Law of planetary motion.

Step 2: Kepler's Third Law states that the square of the orbital period (T) of a planet (or satellite) is directly proportional to the cube of the semi-major axis of its orbit. For circular orbits, the semi-major axis is simply the radius (r). Mathematically, this means the ratio $\frac{T^2}{r^3}$ is a constant for all objects orbiting the same central body.

Step 3: We can use this law to create a relationship between the two satellites. Let the first satellite have period T_1 and radius r_1 , and the second have T_2 and r_2 . Then we can write:

$$\frac{T_1^2}{r_1^3} = \frac{T_2^2}{r_2^3}$$

We need to find T_2 , so we rearrange the formula:

$$T_2^2 = T_1^2 \left(\frac{r_2}{r_1} \right)^3 \implies T_2 = T_1 \left(\frac{r_2}{r_1} \right)^{3/2}$$

Step 4: Now, we substitute the given numerical values into the equation. - Period of the first satellite, $T_1 = 3$ hours. - Radius of the first satellite, $r_1 = 12000$ km. - Radius of the second satellite, $r_2 = 48000$ km. First, calculate the ratio of the radii: $\frac{r_2}{r_1} = \frac{48000 \text{ km}}{12000 \text{ km}} = 4$. Now, plug this ratio into the formula for T_2 :

$$T_2 = 3 \text{ hours} \times (4)^{3/2}$$

To evaluate $(4)^{3/2}$, we can think of it as $(\sqrt{4})^3$. Since $\sqrt{4} = 2$, this becomes $(2)^3 = 8$. Finally, calculate T_2 :

$$T_2 = 3 \text{ hours} \times 8 = 24 \text{ hours}$$

💡 Quick Tip

For problems involving Kepler's Third Law, always set up the ratio $(T_2/T_1)^2 = (r_2/r_1)^3$. It simplifies the calculation and you don't need to know the mass of the central planet or the gravitational constant.

29. The first law of thermodynamics is the statement of

- (A) conservation of heat
- (B) conservation of work
- (C) conservation of momentum
- (D) conservation of energy

Correct Answer: (D)

Solution: Step 1: Let's state the First Law of Thermodynamics in its common mathematical form:

$$\Delta U = Q - W$$

Here, ΔU represents the change in the internal energy of a system, Q represents the heat added to the system from its surroundings, and W represents the work done by the system on its surroundings.

Step 2: Now, let's interpret what this equation physically means. Internal energy (U), heat (Q), and work (W) are all forms of energy. The equation is an energy balance sheet. It says that the change in a system's internal energy is precisely accounted for by the energy that flows into it as heat and the energy that flows out of it as work.

Step 3: Let's relate this to fundamental conservation laws. The statement that energy can be transferred (as heat or work) and transformed (between internal, heat, and work), but the total amount remains constant, is the very definition of the law of conservation of energy. The First Law of Thermodynamics is simply the application of this universal principle to thermodynamic systems. It confirms that energy cannot be created or destroyed within a thermodynamic process.

 Quick Tip

- First Law of Thermodynamics: Conservation of Energy ($\Delta U = Q - W$). - Second Law of Thermodynamics: Entropy of an isolated system always increases. - Zeroth Law of Thermodynamics: Defines thermal equilibrium.

30. A body P at 1000 K emits maximum energy at a wavelength of 3000 nm. If another body Q emits maximum energy at wavelength 550 nm, what will be the temperature of that body Q?

- (A) 5454 K
- (B) 6250 K
- (C) 3125 K
- (D) 4000 K

Correct Answer: (A)

Solution: Step 1: The physical principle that connects the temperature of a body to the wavelength at which it emits the most radiation is Wien's Displacement Law. This law applies to objects that behave as black-body radiators.

Step 2: Wien's Law states that the peak emission wavelength, $\lambda_{\max}$, is inversely proportional to the absolute temperature, T , of the body. The mathematical formula is:

$$\lambda_{\max}T = b$$

where b is Wien's displacement constant. Since b is a constant, the product $\lambda_{\max}T$ is the same for any black body.

Step 3: We can use this principle to set up a ratio between body P and body Q. Let their properties be (λ_P, T_P) and (λ_Q, T_Q) respectively.

$$\lambda_P T_P = \lambda_Q T_Q$$

Step 4: We are asked to find the temperature of body Q, so we rearrange the equation to solve for T_Q .

$$T_Q = T_P \left(\frac{\lambda_P}{\lambda_Q} \right)$$

Step 5: Substitute the given values into this formula to calculate the temperature. - $T_P = 1000$ K - $\lambda_P = 3000$ nm - $\lambda_Q = 550$ nm (Since we are taking a ratio of the wavelengths, the units of nanometers (nm) will cancel out, so there is no need to convert to meters.)

$$T_Q = 1000 \text{ K} \times \frac{3000 \text{ nm}}{550 \text{ nm}} = 1000 \times \frac{300}{55}$$

Simplifying the fraction $\frac{300}{55} = \frac{60}{11} \approx 5.4545$.

$$T_Q \approx 1000 \times 5.4545 = 5454.5 \text{ K}$$

This value matches option (A).

 Quick Tip

Wien's Law ($\lambda_{\text{max}}T = \text{constant}$) is perfect for ratio problems involving the peak emission wavelength and temperature of black bodies. Remember that as an object gets hotter, its peak emission shifts to shorter wavelengths (bluer color).

31. Choose the CORRECT statement about the behavior of product PV against pressure for real gas. [where P is pressure and V is volume]

- (A) Below Boyle's temperature, the value of PV first increases with increase in pressure, reaches a maximum value at a particular temperature and then begins to decrease
- (B) Above Boyle's temperature, the value of PV continuously decreases with increase in pressure.
- (C) Below Boyle's temperature, the value of PV first decreases with increase in pressure, reaches a minimum value at a particular temperature and then begins to increase.
- (D) Above Boyle's temperature, the value of PV is almost constant with increase in pressure.

Correct Answer: (C)

Solution: Step 1: First, recall that for an ideal gas at a constant temperature, the product PV is constant according to Boyle's Law. A plot of PV vs. P would be a horizontal line. Real gases, however, deviate from this behavior due to intermolecular forces and the finite volume of gas molecules.

Step 2: Define Boyle's Temperature (T_B). This is a unique temperature for any real gas at which it behaves most like an ideal gas over a range of pressures. At this temperature, the effects of long-range attractive forces and short-range repulsive forces between molecules effectively cancel each other out.

Step 3: Now, analyze the behavior of a real gas at temperatures below Boyle's Temperature. At these lower temperatures, the intermolecular attractive forces are significant. As you start to increase the pressure from zero, these attractive forces pull the molecules together, making the gas more compressible than an ideal gas. This causes the volume V to decrease more rapidly than the pressure P increases, leading to an initial decrease in the product PV. However, as the pressure becomes very high, the molecules are forced extremely close together, and short-range repulsive forces begin to dominate. The gas becomes less compressible than an ideal gas, and the product PV begins to increase. This creates a characteristic "dip" or minimum in the PV vs. P graph.

Step 4: Analyze the behavior at temperatures above Boyle's Temperature. At these high temperatures, the kinetic energy of the molecules is so large that the attractive forces become negligible. The behavior is dominated by the repulsive forces (i.e., the finite size of the molecules) at all pressures. This makes the gas consistently less compressible than an ideal gas, and the product PV continuously increases as pressure increases.

Step 5: Now evaluate the given statements. (A) This is incorrect. The initial behavior below T_B is a decrease, not an increase. (B) This is incorrect. Above T_B , PV continuously increases.

(C) This statement accurately describes the behavior below Boyle's temperature: an initial decrease in PV to a minimum value, followed by an increase at higher pressures. This is the correct answer. (D) This describes an ideal gas, not a real gas above T_B .

💡 Quick Tip

Remember the PV vs. P isotherm graphs: - Above T_B : PV steadily increases with P. - At T_B : PV is nearly constant for a range of low pressures, then increases. - Below T_B : PV first decreases, reaches a minimum, and then increases.

32. The appearance of the blue color of the sky and reddish color of the sun at sunrise and sunset is due to which phenomenon of light?

- (A) Interference
- (B) Reflection
- (C) Refraction
- (D) Scattering

Correct Answer: (D)

Solution: Step 1: The phenomenon responsible for these atmospheric colors is known as Rayleigh scattering. This type of scattering occurs when light interacts with particles that are much smaller than the light's wavelength. In the Earth's atmosphere, the primary scatterers are nitrogen and oxygen molecules.

Step 2: A key feature of Rayleigh scattering is that it is highly dependent on wavelength. Specifically, the intensity of scattered light is inversely proportional to the fourth power of the wavelength ($I \propto 1/\lambda^4$). This means that shorter wavelengths are scattered far more effectively than longer wavelengths. In the visible spectrum, blue and violet light have the shortest wavelengths, while red and orange have the longest.

Step 3: To explain the blue sky, consider sunlight entering the atmosphere when the sun is high overhead. The blue and violet components of the sunlight are scattered strongly in all directions by the air molecules. When you look up at the sky away from the sun, your eyes receive this scattered blue light from all over the atmosphere, making the sky appear blue.

Step 4: To explain the reddish sun at sunrise and sunset, consider the path of sunlight when the sun is near the horizon. The light must travel through a much thicker section of the atmosphere to reach your eyes. During this long journey, most of the shorter-wavelength blue and green light is scattered away from your direct line of sight. The light that successfully passes through without being scattered is predominantly composed of the longer wavelengths—red, orange, and yellow. This is why the sun itself appears reddish, and the clouds around it are illuminated with these warm colors.

 Quick Tip

Scattering explains colors based on particle size and light path: - Blue Sky: Short-wavelength blue light is scattered more by small air molecules. - Red Sunset: Long path through atmosphere scatters blue light away, leaving red light to pass through. - White Clouds: Larger water droplets scatter all colors equally (Mie scattering).

33. At what rate will energy be emitted from a black body whose filament has a temperature of 3600 K, if at 1800 K the energy is emitted at the rate of 16W?

- (A) 32 W
- (B) 64 W
- (C) 128 W
- (D) 256 W

Correct Answer: (D)

Solution: Step 1: The physical principle that governs the total energy radiated by a black body is the Stefan-Boltzmann Law.

Step 2: This law states that the total power (P), or the rate of energy emission, from a black body is directly proportional to the fourth power of its absolute temperature (T). The full equation is $P = \sigma \epsilon A T^4$, but for a specific object, the Stefan-Boltzmann constant (σ), emissivity (ϵ), and surface area (A) are all constant. Therefore, we can use the simple proportionality:

$$P \propto T^4$$

Step 3: We can use this proportionality to create a ratio that relates the power at two different temperatures. Let P_1 and T_1 be the initial power and temperature, and P_2 and T_2 be the final power and temperature. The ratio is:

$$\frac{P_2}{P_1} = \frac{T_2^4}{T_1^4} = \left(\frac{T_2}{T_1}\right)^4$$

Step 4: Now, we substitute the given values into this ratio to find the unknown power, P_2 . - Initial Power, $P_1 = 16$ W. - Initial Temperature, $T_1 = 1800$ K. - Final Temperature, $T_2 = 3600$ K. First, we find the ratio of the temperatures:

$$\frac{T_2}{T_1} = \frac{3600 \text{ K}}{1800 \text{ K}} = 2$$

This means the absolute temperature has been doubled. Now, we plug this ratio into our power equation:

$$\frac{P_2}{16 \text{ W}} = (2)^4 = 16$$

Finally, we solve for P_2 :

$$P_2 = 16 \times 16 \text{ W} = 256 \text{ W}$$

 Quick Tip

The Stefan-Boltzmann Law ($P \propto T^4$) shows a very strong dependence on temperature. Doubling the absolute temperature increases the radiated power by a factor of $2^4 = 16$.

34. Two satellites A and B go around a planet in circular orbits having radii of $4R$ and R , respectively. If the velocity of satellite A is $3v$, the velocity of satellite B will be:

- (A) $18v$
- (B) $12v$
- (C) $6v$
- (D) $3v$

Correct Answer: (C)

Solution: Step 1: To find the velocity of a satellite in a stable circular orbit, we must recognize that the gravitational force exerted by the planet provides the necessary centripetal force to keep the satellite in its path.

Step 2: We can write this relationship as an equation. Let M be the mass of the planet, m be the mass of the satellite, r be the orbital radius, and v be the orbital velocity.
Gravitational Force = Centripetal Force

$$\frac{GMm}{r^2} = \frac{mv^2}{r}$$

We can simplify this by canceling m and one r from each side, then solve for v :

$$v^2 = \frac{GM}{r} \implies v = \sqrt{\frac{GM}{r}}$$

Step 3: From this formula, we can see the relationship between velocity and radius. Since G and M are constant for both satellites, the velocity v is inversely proportional to the square root of the radius r .

$$v \propto \frac{1}{\sqrt{r}}$$

This means that the satellite in the smaller orbit (closer to the planet) will have a higher velocity.

Step 4: Let's set up a ratio between the velocities of satellite A and satellite B to eliminate the constants.

$$\frac{v_B}{v_A} = \frac{\sqrt{GM/r_B}}{\sqrt{GM/r_A}} = \sqrt{\frac{GM}{r_B} \cdot \frac{r_A}{GM}} = \sqrt{\frac{r_A}{r_B}}$$

Step 5: Now, substitute the given values and solve for the velocity of satellite B, v_B . - Velocity of A, $v_A = 3v$. - Radius of A, $r_A = 4R$. - Radius of B, $r_B = R$.

$$\frac{v_B}{3v} = \sqrt{\frac{4R}{R}} = \sqrt{4} = 2$$

Now, solve for v_B :

$$v_B = 2 \times (3v) = 6v$$

 Quick Tip

Remember the relationships for orbital motion: - Velocity: $v \propto 1/\sqrt{r}$ (closer satellite is faster) - Period: $T \propto r^{3/2}$ (closer satellite has shorter period)

35. The electric field at a point inside a charged hollow spherical conductor is:

- (A) zero
- (B) constant but not zero
- (C) depends on the distance from centre
- (D) depends on the charge on the conductor

Correct Answer: (A)

Solution: Step 1: The key principle to solve this problem is Gauss's Law, a fundamental concept in electrostatics. Gauss's Law relates the electric flux through a closed surface to the net electric charge enclosed by that surface. The law is expressed as:

$$\Phi_E = \oint \vec{E} \cdot d\vec{A} = \frac{Q_{\text{enclosed}}}{\epsilon_0}.$$

Step 2: Another crucial principle is the behavior of charges on a conductor in electrostatic equilibrium. Because charges are free to move within a conductor, they will arrange themselves to minimize repulsion. For any conductor, this results in all excess charge residing on the outer surface. For the hollow spherical conductor in this problem, any charge it holds will be distributed uniformly on its external surface.

Step 3: To find the electric field at a point inside the hollow conductor, we can imagine a closed spherical surface (a "Gaussian surface") with a radius r that is smaller than the radius of the conductor, centered at the same point. This imaginary surface lies entirely within the hollow region.

Step 4: Now we apply Gauss's Law to this imaginary surface. Since all the charge resides on the outer surface of the conductor, the amount of charge enclosed by our Gaussian surface, Q_{enclosed} , is zero.

Step 5: According to Gauss's Law, if $Q_{\text{enclosed}} = 0$, then the total electric flux (Φ_E) through the surface must also be zero. Due to the spherical symmetry of the problem, if the electric field ($\vec{E}$) were anything other than zero, it would have to be constant and point radially at every point on our Gaussian surface, resulting in a non-zero flux. Since the flux is zero, the electric field $\vec{E}$ itself must be zero everywhere inside the hollow conductor.

 Quick Tip

A key result from electrostatics: the electric field inside any static conductor (hollow or solid) is always zero. This is the principle behind electrostatic shielding, used in devices like Faraday cages.

36. Match LIST-I with LIST-II

LIST-I

- (A) Lyman Series
- (B) Balmer Series
- (C) Paschen Series
- (D) Pfund Series

LIST-II

- (I) Microwave Range
- (II) Ultraviolet Range
- (III) Visible Range
- (IV) Infrared Range

- (A) A - I, B - II, C - III, D - IV
- (B) A - II, B - III, C - I, D - IV
- (C) A - I, B - II, C - IV, D - III
- (D) A - II, B - III, C - IV, D - I

Correct Answer: (None of the options exactly match the typical classification, but (D) is the closest if Pfund is matched to microwave.) The standard matching is Lyman(II), Balmer(III), Paschen(IV), Brackett(IV), Pfund(IV).

Solution: The spectral series of hydrogen are produced when an electron transitions from a higher energy level ($n_{initial}$) to a lower energy level (n_{final}). The energy of the emitted photon (and thus its spectral range) depends on the final level, n_{final} .

Step 1: Analyze Lyman Series (A)

This series involves all transitions that end at the lowest energy level, $n_{final} = 1$. These are the largest possible energy drops for an electron in a hydrogen atom, resulting in the emission of high-energy photons. These photons fall into the Ultraviolet (II) region of the electromagnetic spectrum. So, **A matches with (II)**.

Step 2: Analyze Balmer Series (B)

This series involves all transitions that end at the second energy level, $n_{final} = 2$. These energy drops are smaller than those in the Lyman series. The photons emitted have wavelengths corresponding to the Visible (III) spectrum (as well as some in the near ultraviolet). This is the only series in hydrogen with lines visible to the naked eye. So, **B matches with (III)**.

Step 3: Analyze Paschen Series (C)

This series involves all transitions that end at the third energy level, $n_{final} = 3$. The energy drops are smaller still, leading to the emission of lower-energy photons. These photons fall within the Infrared (IV) range. So, **C matches with (IV)**.

Step 4: Analyze Pfund Series (D)

This series involves all transitions that end at the fifth energy level, $n_{final} = 5$. These are very small energy drops, corresponding to long-wavelength photons in the far Infrared (IV) range. The question's options present an issue, as both Paschen and Pfund series belong to the Infrared range. However, option (D) is the only one that correctly matches the first three series. It matches Pfund with Microwave (I). While technically incorrect (it's far-infrared), we

select this as the intended answer by process of elimination. So, **D is intended to match with (I)**.

 Quick Tip

Remember the order of the hydrogen spectral series by the final energy level n_f : 1. Lyman ($n_f = 1$) - Ultraviolet 2. Balmer ($n_f = 2$) - Visible 3. Paschen ($n_f = 3$) - Infrared 4. Brackett ($n_f = 4$) - Infrared 5. Pfund ($n_f = 5$) - Infrared

37. A light beam traveling in the x-direction is described by the magnetic field $B_z = 2 \times 10^{-6} \sin \omega(t - x/c)$ Tesla. The value of the maximum electric field is

- (A) 200 V/m
- (B) 300 V/m
- (C) 600 V/m
- (D) 800 V/m

Correct Answer: (C)

Solution: Step 1: Recall the fundamental relationship between the electric field and magnetic field in an electromagnetic wave traveling through a vacuum. The magnitudes of the electric field (E) and the magnetic field (B) are directly proportional at every point in space and time.

Step 2: State the equation that connects these two fields. The ratio of the electric field magnitude to the magnetic field magnitude is equal to the speed of light, c .

$$\frac{E}{B} = c \quad \text{or} \quad E = cB$$

This relationship also holds true for the maximum values, or amplitudes, of the fields, which are denoted as E_0 and B_0 .

$$E_0 = cB_0$$

Step 3: Identify the amplitude of the magnetic field (B_0) from the given wave equation. The equation is provided in the standard sinusoidal form $B_z = B_0 \sin(\text{phase})$. By comparing this to the given equation, $B_z = 2 \times 10^{-6} \sin \omega(t - x/c)$, we can directly identify the amplitude:

$$B_0 = 2 \times 10^{-6} \text{ Tesla}$$

Step 4: Use the formula from Step 2 to calculate the maximum electric field, E_0 . We use the value for the speed of light in a vacuum, $c \approx 3 \times 10^8 \text{ m/s}$.

$$E_0 = (3 \times 10^8 \text{ m/s}) \times (2 \times 10^{-6} \text{ T})$$

$$E_0 = (3 \times 2) \times (10^8 \times 10^{-6}) \text{ V/m}$$

$$E_0 = 6 \times 10^2 \text{ V/m}$$

$$E_0 = 600 \text{ V/m}$$

💡 Quick Tip

For any electromagnetic wave in a vacuum, the ratio of the electric field amplitude to the magnetic field amplitude is always the speed of light: $E_0/B_0 = c$.

38. Arrange the following minerals in decreasing order of their hardness as per Mohs scale.

- A. Orthoclase
- B. Calcite
- C. Corundum
- D. Fluorite

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

Correct Answer: (C)

Solution: Step 1: The first step is to recall the Mohs scale of mineral hardness, which ranks minerals from 1 (softest) to 10 (hardest) based on their relative ability to scratch one another. We need to identify the hardness value for each of the four minerals listed. - Corundum (C): This is a very hard mineral, used as an abrasive. It has a hardness of 9 on the Mohs scale. - Orthoclase (A): This is a type of feldspar and is a key index mineral on the scale. It has a hardness of 6. - Fluorite (D): This mineral is softer than Orthoclase. It has a hardness of 4. - Calcite (B): This is a relatively soft mineral. It has a hardness of 3.

Step 2: The question asks for the minerals to be arranged in decreasing order of hardness, which means starting with the hardest and ending with the softest. Using the values from Step 1, the order of hardness is: 1. Corundum (Hardness = 9) 2. Orthoclase (Hardness = 6) 3. Fluorite (Hardness = 4) 4. Calcite (Hardness = 3)

Step 3: Finally, we map this ordered list of minerals back to their corresponding letters given in the question. The sequence is Corundum (C), Orthoclase (A), Fluorite (D), Calcite (B). This gives us the final order: C, A, D, B.

💡 Quick Tip

A common mnemonic for the Mohs scale is: "Tall Girls Can Find Alligators Or Queens Through Careful Digging." (1-Talc, 2-Gypsum, 3-Calcite, 4-Fluorite, 5-Apatite, 6-Orthoclase, 7-Quartz, 8-Topaz, 9-Corundum, 10-Diamond)

39. In the Bowen's reaction series, which of the following mineral(s) is/are crystallized prior to amphibole?

- A. Biotite
- B. Clinopyroxene

C. K-feldspar

D. Forsterite

- (A) A and B only
- (B) A, B and D only
- (C) B only
- (D) B and D only

Correct Answer: (D)

Solution: Step 1: First, we must understand Bowen's reaction series. This series describes the sequence in which different silicate minerals crystallize as a magma cools. The series has two branches: a discontinuous branch for ferromagnesian (iron-magnesium rich) minerals, and a continuous branch for plagioclase feldspar. Minerals at the top of the series crystallize at higher temperatures.

Step 2: The question concerns minerals in the discontinuous branch. We need to list the key minerals of this series in their order of crystallization, which is also the order of decreasing temperature. The sequence is: 1. Olivine (e.g., Forsterite) - Crystallizes at the highest temperature. 2. Pyroxene (e.g., Clinopyroxene) - Crystallizes after Olivine. 3. Amphibole - Crystallizes after Pyroxene. 4. Biotite mica - Crystallizes after Amphibole at a lower temperature.

Step 3: The question asks which minerals crystallize "prior to amphibole". This means we need to identify the minerals that appear earlier in the sequence (i.e., at higher temperatures) than amphibole. Looking at our list, both Olivine (D. Forsterite) and Pyroxene (B. Clinopyroxene) crystallize before amphibole does.

Step 4: We should also check the other options to confirm our answer. - A. Biotite crystallizes after amphibole, at a lower temperature. - C. K-feldspar (potassium feldspar) is part of the lower-temperature, felsic crystallization sequence, crystallizing well after amphibole.

Conclusion: Therefore, the minerals that crystallize prior to amphibole are Clinopyroxene (B) and Forsterite (D).

 Quick Tip

Remember the order of the discontinuous series: Olivine -> Pyroxene -> Amphibole -> Biotite. A mnemonic could be "Old Pilots Are Boring".

40. The oldest crust on the earth is found at _____.

- (A) mid oceanic ridge
- (B) subduction zones
- (C) sea floor spreading
- (D) transform plate boundary

Correct Answer: (B)

Solution: Step 1: To answer this, we must understand the life cycle of oceanic crust, a process driven by plate tectonics. This cycle is often compared to a conveyor belt.

Step 2: The process begins at mid-oceanic ridges. These are divergent plate boundaries where magma from the mantle rises to the surface and solidifies, forming brand new oceanic crust. This is the "birthplace" of the crust, so the rock here is the youngest on the seafloor. The process of creating this new crust is called sea-floor spreading.

Step 3: As new crust is continuously formed at the ridge, it pushes the older crust away from the ridge on both sides. This older crust travels across the ocean floor for millions of years, cooling, contracting, and becoming denser as it moves.

Step 4: The conveyor belt journey ends when the oceanic plate collides with another tectonic plate (either continental or another oceanic plate). Because the old oceanic crust is now cold and dense, it sinks back into the mantle in a process called subduction. The locations where this occurs are known as subduction zones. Therefore, the oceanic crust found at the edge of a subduction zone, just before it is destroyed, is the oldest on the entire ocean floor.

 Quick Tip

Remember the conveyor belt analogy for oceanic crust: it's created at mid-ocean ridges (youngest), travels across the ocean floor, and is destroyed at subduction zones (oldest). The absolute oldest rocks on Earth are found on continents in areas called cratons.

41. The Tethys sea existed between _____ continents.

- (A) Madagascar and Africa
- (B) Africa and Eurasia
- (C) Australia and Asia
- (D) Antarctica and Australia

Correct Answer: (B)

Solution: Step 1: To understand the location of the Tethys Sea, we must go back in geological time to the Mesozoic Era, when the supercontinent Pangaea was breaking apart.

Step 2: Pangaea did not split randomly; it first broke into two large supercontinents. In the north was Laurasia, and in the south was Gondwana.

Step 3: The Tethys Sea was the massive east-west ocean that existed in the large gap that opened up between these two supercontinents. It separated Laurasia from Gondwana.

Step 4: Now we must identify which modern continents were part of Laurasia and Gondwana. - Laurasia was composed of the landmasses that would become North America, Europe, and Asia (collectively, Eurasia). - Gondwana was composed of the landmasses that would become South America, Africa, Antarctica, Australia, and the Indian subcontinent.

Step 5: Based on this, the Tethys Sea was situated between the northern continents (Eurasia) and the southern continents (Africa, among others). Over millions of years, as the

continents continued to drift, the African plate moved northward and collided with the Eurasian plate, closing most of the Tethys Sea. This collision is responsible for creating the Alpine-Himalayan mountain belt. Therefore, the Tethys Sea existed between the landmasses that are now Africa and Eurasia.

 Quick Tip

Think of the Mediterranean Sea as the primary remnant of the ancient Tethys Sea, which once separated Africa from Europe (Eurasia).

42. Match LIST-I with LIST-II

LIST-I

- (A) Epicenter
- (B) Focus
- (C) Mercalli scale
- (D) Richter scale

LIST-II

- (I) is the scale of measurement of degree of destructiveness of earthquakes
 - (II) is a point of the origin of the earthquake inside the earth
 - (III) is the scale to measure the magnitude of the energy released by an earthquake
 - (IV) is the place on the ground surface which is perpendicular to the hypocenter
- (A) A-IV, B-II, C-I, D-III
 - (B) A-II, B-IV, C-I, D-III
 - (C) A-IV, B-I, C-III, D-II
 - (D) A-III, B-IV, C-I, D-II

Correct Answer: (A)

Solution: Let's define each term in LIST-I and find its corresponding description in LIST-II.

Step 1: Analyze Epicenter (A)

The focus (or hypocenter) of an earthquake is the point within the Earth where the fault rupture begins. The epicenter is the point on the Earth's surface located directly above the focus. Description (IV) states it is "the place on the ground surface which is perpendicular to the hypocenter," which is the definition of the epicenter. Thus, **A matches with (IV)**.

Step 2: Analyze Focus (B)

As defined above, the focus is the actual point of origin of the earthquake, deep inside the Earth's crust where the seismic energy is first released. This matches description (II). Thus, **B matches with (II)**.

Step 3: Analyze Mercalli scale (C)

The Modified Mercalli Intensity (MMI) scale is a qualitative scale that measures the intensity of an earthquake. It describes the effects of the shaking at a specific location, based on observed damage, and what people felt. It directly measures the "degree of destructiveness," as stated in description (I). Thus, **C matches with (I)**.

Step 4: Analyze Richter scale (D)

The Richter scale (and the modern Moment Magnitude scale that has largely replaced it) is a quantitative, logarithmic scale that measures the magnitude of an earthquake. The magnitude is a single value that represents the total amount of seismic energy released at the earthquake's source (the focus). This matches description (III). Thus, **D matches with (III)**.

Conclusion: The correct set of matches is A-IV, B-II, C-I, D-III.

💡 Quick Tip

Remember the difference: - Richter = Magnitude (Energy) - A single number for the whole earthquake. - Mercalli = Intensity (Damage) - Varies with location. - Focus = Inside the Earth. - Epicenter = On the Surface.

43. Arrange the following according to the depth of occurrence starting from deep to shallow.

- A. Continental slope
- B. Continental shelf
- C. Abyssal plains
- D. Continental rise

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

Correct Answer: (B)

Solution: Step 1: To solve this, it's best to first visualize the profile of the ocean floor as one moves away from a continent into the deep ocean. This gives the order from shallowest to deepest.

Step 2: Let's define each feature in that shallow-to-deep order. 1. Continental shelf (B): This is the gently sloping, shallow, submerged edge of a continent. It is the shallowest of the four features. 2. Continental slope (A): At the edge of the shelf, the seafloor drops off steeply. This steep incline is the continental slope. It is deeper than the shelf. 3. Continental rise (D): At the base of the steep continental slope, sediments accumulate in a large wedge. This feature is the continental rise. It is deeper than the slope. 4. Abyssal plains (C): Beyond the continental rise lie the abyssal plains. These are the vast, flat, and very deep areas that make up most of the ocean floor. They are the deepest of the four features.

Step 3: Based on the definitions above, the order from shallowest to deepest is: Continental shelf (B) → Continental slope (A) → Continental rise (D) → Abyssal plains (C).

Step 4: The question asks for the arrangement in decreasing order of depth, which means starting from deep to shallow. Therefore, we must reverse the sequence we found in Step 3. The correct deep-to-shallow order is: Abyssal plains (C) → Continental rise (D) →

Continental slope (A) → Continental shelf (B). This corresponds to the letter sequence C, D, A, B.

 Quick Tip

Visualize walking from a beach into the deep ocean. You walk on the Shelf, then fall down the Slope, land on the pile of sediment at the bottom called the Rise, and then walk out onto the flat Abyssal Plain. The question asks for the reverse order.

44. Lithosphere is characterized by -----.

A. A thickness of about 100 km

B. Average density of 3.6 g/cc

C. Silicon and aluminum as dominant constituents

D. Basalt as the dominant rock

(A) A and C only

(B) B and D only

(C) A, B, and D only

(D) A and D only

Correct Answer: (D) (Note: This is the best fit among poor options, likely focusing on oceanic lithosphere)

Solution: Step 1: Analyze statement A.

The lithosphere is the rigid outer mechanical layer of the Earth, which includes the crust and the solid uppermost portion of the mantle. Its thickness varies significantly, being thinner under oceans (as little as 15 km at ridges) and much thicker under continents (up to 200 km or more). However, a value of 100 km is a widely used representative average thickness. So, statement A is considered correct.

Step 2: Analyze statement B.

The density of the lithosphere varies. Continental crust has a density of about 2.7 g/cc, while oceanic crust is denser at about 3.0 g/cc. The underlying rigid mantle is even denser, about 3.3 g/cc. The average density of the lithosphere is therefore a weighted average of these, typically around 3.1 g/cc. A value of 3.6 g/cc is too high and is more typical of the denser, plastic asthenosphere beneath the lithosphere. Statement B is incorrect.

Step 3: Analyze statement C.

Silicon and aluminum are the dominant elements of the continental crust, often referred to by the mnemonic "SiAl". However, the oceanic crust is dominated by silicon and magnesium ("SiMa"). Since the lithosphere includes both types of crust as well as the upper mantle (rich in Si, O, Fe, Mg), this statement is an oversimplification and only applies to the continental part. It is not a characteristic of the entire lithosphere.

Step 4: Analyze statement D.

The Earth's surface is covered by about 70

Conclusion: Statements A and D are the most accurate general characterizations. Therefore, the best option is "A and D only".

 Quick Tip

Remember the key distinctions: - Lithosphere: Crust + Rigid Upper Mantle. - Continental Crust: Thicker, less dense, granitic (SiAl). - Oceanic Crust: Thinner, denser, basaltic (SiMa).

45. The amount of water vapor generally increases with increase in air temperature. The amount of water vapor in the air at a particular temperature is referred to as

- (A) Super saturation
- (B) Relative humidity
- (C) Absolute humidity
- (D) Saturation ratio

Correct Answer: (C)

Solution: Step 1: The question asks for the term that describes the actual amount of water vapor present in the air. Let's carefully define the options to find the correct match.

Step 2: Let's define the key terms related to atmospheric moisture. - Absolute humidity: This is a direct measure of the mass of water vapor contained within a specific volume of air. It is typically expressed in units like grams of water vapor per cubic meter of air (g/m^3). This definition directly matches the question's phrase "the amount of water vapor in the air". - Relative humidity: This is not a measure of the actual amount, but a ratio. It compares the current absolute humidity to the maximum possible humidity that the air could hold at its current temperature, expressed as a percentage. It tells us how "saturated" or "full" the air is. - Super saturation: This is a specific, unstable condition where the relative humidity is over 100- Saturation ratio: This is another term for the ratio of the actual amount of water vapor to the amount at saturation, similar to relative humidity but often expressed as a decimal rather than a percentage.

Step 3: By comparing the definitions, it is clear that Absolute humidity is the term that refers to the actual quantity or amount of water vapor in a given volume of air.

 Quick Tip

- Absolute Humidity: How much water vapor is actually there (mass/volume). - Relative Humidity: How "full" of water vapor the air is (actual amount / maximum possible amount).

46. Match LIST-I with LIST-II

LIST-I (Cloud related prefix)

- (A) Cumulus
- (B) Nimbus
- (C) Cirrus

- (D) Stratus
 LIST-II (Characteristics)
 (I) Rain bearing
 (II) Stratified layer
 (III) Puffy structure
 (IV) Composed of ice crystals
 (A) A-IV, B-II, C-III, D-I
 (B) A-III, B-I, C-IV, D-II
 (C) A-III, B-I, C-II, D-IV
 (D) A-I, B-IV, C-III, D-II

Correct Answer: (B)

Solution: Let's analyze each cloud type in LIST-I and match it with its defining characteristic from LIST-II. The names of clouds are derived from Latin roots that describe their appearance or function.

Step 1: Analyze Cumulus (A)

The Latin word 'cumulus' means 'heap' or 'pile'. Cumulus clouds are the classic puffy, cotton-like clouds with distinct flat bases and bulging upper parts. This corresponds to the "Puffy structure" in description (III). Thus, **A matches with (III)**.

Step 2: Analyze Nimbus (B)

The Latin word 'nimbus' means 'rain cloud'. Whenever this word is used as a prefix or suffix (like in cumulonimbus or nimbostratus), it signifies that the cloud is producing precipitation (rain, snow, etc.). This corresponds to "Rain bearing" in description (I). Thus, **B matches with (I)**.

Step 3: Analyze Cirrus (C)

The Latin word 'cirrus' means 'curl of hair' or 'wisp'. Cirrus clouds are high-altitude clouds that appear thin, delicate, and feathery. Because of the extremely cold temperatures at high altitudes, they are "Composed of ice crystals", not liquid water droplets. This corresponds to description (IV). Thus, **C matches with (IV)**.

Step 4: Analyze Stratus (D)

The Latin word 'stratus' means 'to spread out' or 'layer'. Stratus clouds are grey, featureless clouds that appear in flat sheets or layers, often covering the entire sky like fog that doesn't reach the ground. This corresponds to "Stratified layer" in description (II). Thus, **D matches with (II)**.

Conclusion: Combining these matches gives the sequence: A-III, B-I, C-IV, D-II.

 Quick Tip

Remember the Latin roots for cloud types: - Stratus: Layer - Cumulus: Heap - Cirrus: Curl of hair (high, icy) - Nimbus: Rain - Alto: Mid-level

47. The effective average solar flux incident to a level surface at top of the atmosphere is -----.

- (A) 124 W/m²
- (B) 242 W/m²
- (C) 89 W/m²
- (D) 342 W/m²

Correct Answer: (D)

Solution: Step 1: We begin with the concept of the Solar Constant (S_0). This is the amount of solar radiation received per unit area on a surface held perpendicular to the Sun's rays at the average distance of the Earth from the Sun. Its accepted value is approximately 1361 W/m². This is the maximum intensity of sunlight at the top of our atmosphere.

Step 2: The question asks for the "effective average solar flux over the entire planet". The Earth is a sphere, not a flat disk perpendicular to the sun. The intense sunlight described by the solar constant is only intercepted by the "daytime" side of the Earth, which presents a circular cross-section to the sun. The area of this circle is πR^2 , where R is the Earth's radius.

Step 3: However, this intercepted energy is distributed over the Earth's entire surface area as the planet rotates. The total surface area of the spherical Earth is $4\pi R^2$.

Step 4: To find the average solar flux over the whole planet, we must divide the total intercepted energy by the total surface area over which it is distributed.

$$\text{Average Flux} = \frac{\text{Total Intercepted Power}}{\text{Total Surface Area}} = \frac{S_0 \times (\pi R^2)}{4\pi R^2}$$

The πR^2 terms cancel out, leaving:

$$\text{Average Flux} = \frac{S_0}{4} = \frac{1361 \text{ W/m}^2}{4} \approx 340.25 \text{ W/m}^2$$

Step 5: This calculated value is often rounded for use in climate models. The value of 342 W/m² is the standard, commonly cited figure for the globally and annually averaged solar flux at the top of the atmosphere, making it the correct answer.

💡 Quick Tip

A key number in climate science is the globally averaged solar radiation at the top of the atmosphere. It is always the solar constant ($\approx 1361 \text{ W/m}^2$) divided by 4, which gives about 340-342 W/m².

48. The most common mineral used in U-Pb radiometric dating is _____.

- (A) Biotite
- (B) Zircon
- (C) Hornblende
- (D) Quartz

Correct Answer: (B)

Solution: Step 1: First, let's understand the principles of Uranium-Lead (U-Pb) radiometric dating. This method relies on the radioactive decay of two uranium isotopes (^{238}U and ^{235}U) into stable lead isotopes (^{206}Pb and ^{207}Pb , respectively) at known, constant rates. To get an accurate age, the mineral being dated needs to act as a closed system, trapping the parent (U) and daughter (Pb) isotopes.

Step 2: An ideal mineral for this method should have three key properties: 1. It should incorporate a significant amount of Uranium into its crystal lattice as it forms. 2. It should strongly reject Lead from its crystal lattice during formation. This is crucial because it means any lead found in the mineral later is almost entirely the result of radioactive decay, not initial contamination. 3. It must be physically and chemically durable, able to withstand weathering and metamorphism over billions of years without losing the parent or daughter isotopes.

Step 3: Now, let's evaluate the options based on these criteria. (A) Biotite and (C) Hornblende: These are common rock-forming minerals, but they are not ideal for U-Pb dating. They are susceptible to alteration and can "leak" lead at high temperatures, which would reset the radiometric clock and give an inaccurate, younger age. (D) Quartz (SiO_2): This is an extremely common and durable mineral, but its crystal structure does not readily incorporate Uranium or other trace elements needed for dating. It is essentially "empty" of the necessary parent isotopes. (B) Zircon (ZrSiO_4): This mineral is the "gold standard" for U-Pb dating because it perfectly meets all the criteria. When zircon crystals form in a magma, their structure readily accepts Uranium atoms but strongly excludes Lead atoms. Furthermore, zircon is exceptionally hard and resistant to chemical and physical weathering, making it an excellent time capsule for preserving the U-Pb system over vast geological timescales.

 Quick Tip

When you think of U-Pb dating, the primary mineral to remember is Zircon. Its properties of incorporating uranium but rejecting lead, combined with its exceptional durability, make it the gold standard for geochronology.

49. Arrange the following significant salts in decreasing order of their percentage in the oceans:

- A. MgCl_2
 - B. CaSO_4 (assuming typo for CuSO_4)
 - C. MgSO_4
 - D. NaCl
- (A) D, B, C, A
(B) D, A, C, B
(C) D, C, A, B
(D) D, B, A, C

Correct Answer: (B)

Solution: Step 1: First, we must recall the composition of dissolved salts in typical seawater. While many salts are present, a few major ones dominate. The question asks for

the arrangement in decreasing order, meaning from most abundant to least abundant. (Note: The CuSO_4 in the original question is almost certainly a typo for CaSO_4 , Calcium Sulfate, which is a significant component of sea salt).

Step 2: Let's identify the abundance of each salt listed. 1. Sodium Chloride (D. NaCl): This is common table salt and is by far the most abundant dissolved salt in the ocean. It accounts for the vast majority (over 77%). 2. Magnesium Chloride (A. MgCl_2): After NaCl , salts containing magnesium are the next most common. Magnesium chloride is the second most abundant salt in seawater. 3. Magnesium Sulfate (C. MgSO_4): Also known as Epsom salt, this is the third most abundant salt. 4. Calcium Sulfate (B. CaSO_4): This salt, which includes gypsum, is the fourth most abundant of the major salts.

Step 3: Now, we assemble the sequence based on this ranking from most abundant to least abundant. The correct order is: 1st: Sodium Chloride (D) 2nd: Magnesium Chloride (A) 3rd: Magnesium Sulfate (C) 4th: Calcium Sulfate (B)

Step 4: This gives us the final letter sequence: D, A, C, B.

 Quick Tip

Just remember that seawater is overwhelmingly just salty water, so Sodium Chloride (NaCl) will always be first. The next most common positive and negative ions are Magnesium (Mg^{2+}) and Sulfate (SO_4^{2-}), which combine with others to form the next most abundant salts.

50. Which of the following were constituents of the famous classical London smog?

- A. Smoke
- B. Fog
- C. SO_2
- D. Peroxyacetyl nitrate
- E. Ozone

- (A) A, B and E only
- (B) A, B and C only
- (C) A, B, C and D only
- (D) A, B, C and E only

Correct Answer: (B)

Solution: Step 1: First, it's important to distinguish between the two main types of smog. The question refers to "classical London smog," which is also known as sulfurous smog or industrial smog. This is different from the photochemical smog typically associated with cities like Los Angeles.

Step 2: Let's identify the key ingredients and conditions for classical London smog. This type of smog is primarily caused by the burning of large quantities of high-sulfur coal for industrial and home heating purposes. It typically occurs in cool, humid, and stagnant air

conditions. - A. Smoke: The burning of coal produces large amounts of particulate matter, or soot, which is smoke. This is a primary ingredient. - B. Fog: The term "smog" itself is a blend of the words smoke and fog. The high humidity and water vapor in the air (fog) are essential for its formation. - C. SO₂ (Sulfur Dioxide): The high-sulfur coal releases sulfur dioxide, a key chemical component. In the presence of fog, SO₂ can be converted into sulfuric acid droplets, making the smog highly acidic and dangerous.

Step 3: Now, let's consider the other options and determine why they are incorrect for London-type smog. - D. Peroxyacetyl nitrate (PAN) and E. Ozone (O₃): These are not components of classical smog. Instead, they are the hallmark secondary pollutants of photochemical smog. Photochemical smog forms when sunlight acts on primary pollutants like nitrogen oxides (NO_x) and volatile organic compounds (VOCs) from vehicle exhaust. This type of smog is oxidizing, while London smog is reducing.

Conclusion: The correct constituents of classical London smog from the list are Smoke (A), Fog (B), and Sulfur Dioxide (C).

 Quick Tip

- London Smog (Classical): Cool, damp, sulfurous. Key ingredients: Smoke, Fog, SO₂.
- L.A. Smog (Photochemical): Hot, sunny, oxidizing. Key ingredients: NO_x, VOCs, Sunlight, Ozone, PAN.

51. Choose the correct statements about corals?

- A. Corals are invertebrate animals a few millimeter in diameter.**
- B. Corals are sensitive to heat stress and ocean pH change.**
- C. Corals provide algae with stable environment, CO₂ and nutrients.**
- D. Corals bleach to get rid themselves of pathogens.**

- (A) C and D only
- (B) A, B and C only
- (C) A and B only
- (D) A, B, C and D only

Correct Answer: (B)

Solution: Let's analyze each statement to determine its accuracy.

Step 1: Analyze statement A.

A large coral reef structure is not a single organism but a massive colony built by countless tiny organisms called coral polyps. Each individual polyp is indeed an invertebrate animal, belonging to the phylum Cnidaria (related to jellyfish and sea anemones), and is typically only a few millimeters in size. This statement is correct.

Step 2: Analyze statement B.

Corals are highly specialized organisms that thrive only within a narrow range of environmental conditions. Two of the most critical factors are water temperature and pH. Heat stress (prolonged periods of unusually warm water) and ocean acidification (a decrease in ocean pH due to absorption of atmospheric CO₂) are major threats. Both conditions

disrupt the coral's internal processes and can lead to a stress response called bleaching. This statement is correct.

Step 3: Analyze statement C.

Most reef-building corals have a vital symbiotic relationship with microscopic algae called zooxanthellae that live within their tissues. In this partnership, the coral polyp provides the algae with a safe, stable place to live and supplies it with the necessary compounds for photosynthesis: carbon dioxide (CO₂) and nutrients (like nitrogen and phosphorus). In return, the algae produce oxygen and energy-rich organic compounds that the coral uses as its primary food source. This statement is correct.

Step 4: Analyze statement D.

Coral bleaching is the process where corals expel their symbiotic zooxanthellae. This is not a healthy cleaning mechanism to remove pathogens. It is a severe stress response triggered by adverse conditions, most commonly heat stress. The coral's tissues are actually transparent, and their color comes from the algae; when the algae are expelled, the white calcium carbonate skeleton becomes visible through the clear tissue, hence the "bleached" appearance. If conditions do not improve, the coral will starve and die. This statement is incorrect.

Conclusion: Statements A, B, and C are correct, while D is incorrect.

💡 Quick Tip

Coral bleaching is a sign of severe stress, not a healthy cleaning process. It's the breakdown of the vital symbiosis between the coral animal and the algae that gives it color and most of its food.

52. Choose the incorrect statement.

- (A) Tornado is the smallest, most violent weather disturbance that occur on the earth.
- (B) Tornadoes are defined as high pressure center where winds moves into high pressure and move upward.
- (C) Tornadoes are of a dark color due to the dominance of dust, sand and condensed moisture.
- (D) Tornadoes are funnel shaped storms.

Correct Answer: (B)

Solution: The task is to identify the one statement that is false. Let's examine each one.

Step 1: Analyze statement A.

Compared to hurricanes or large storm fronts, tornadoes are geographically small. However, they contain the most concentrated and violent winds found in any weather phenomenon on Earth, with wind speeds that can exceed 300 mph (480 km/h). This statement is considered correct.

Step 2: Analyze statement B.

This statement describes the pressure system of a tornado. All major storms, including tornadoes, hurricanes, and cyclones, are fundamentally intense low-pressure systems. Air in the atmosphere naturally flows from areas of higher pressure to areas of lower pressure. In a

tornado, air violently rushes inward toward the extremely low-pressure center and then spirals upward. A high-pressure center, by contrast, is characterized by calm, sinking air and fair weather. Therefore, defining a tornado as a high-pressure center is fundamentally incorrect.

Step 3: Analyze statement C.

The visible part of a tornado is the funnel cloud, which is composed of condensed water droplets (moisture). This funnel often becomes dark and opaque because its intense winds pick up large amounts of dust, soil, sand, and debris from the ground. This statement is correct.

Step 4: Analyze statement D.

The classic and most recognizable visual characteristic of a tornado is its rotating, funnel-shaped column of air that extends from the base of a cumulonimbus cloud down to the ground. This statement is correct.

Conclusion: Statement (B) provides a completely wrong description of the pressure dynamics within a tornado, making it the incorrect statement.

 Quick Tip

Remember that all major storms (hurricanes, cyclones, tornadoes) are intense low-pressure systems. Air always flows from high pressure to low pressure, so winds rush into a storm's center.

53. Arrange the following greenhouse gases in decreasing order of their contribution to total global warming.

- A. Nitrous oxide
- B. Carbon dioxide
- C. Chlorofluoro carbon
- D. Methane

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

Correct Answer: (D)

Solution: Step 1: It is crucial to understand that a greenhouse gas's "contribution to total global warming" is a product of two factors: its atmospheric concentration (how much of it is there) and its Global Warming Potential (GWP) (how effective it is at trapping heat per molecule, relative to CO₂). The question asks for the ranking based on this total overall impact.

Step 2: Let's rank the gases based on their current overall contribution to the enhanced greenhouse effect. 1. Carbon dioxide (B. CO₂): By definition, CO₂ has a GWP of 1. While it is the least potent per molecule among the main greenhouse gases, its atmospheric concentration is vastly higher than all the others due to the burning of fossil fuels,

deforestation, and industrial processes. This overwhelming abundance makes it the largest contributor to modern global warming. 2. Methane (D. CH₄): Methane has a much higher GWP than CO₂ (about 28-34 times more potent over 100 years), but its atmospheric concentration is far lower. Its high potency combined with significant emissions from agriculture, fossil fuels, and waste makes it the second-largest contributor. 3. Nitrous oxide (A. N₂O): Nitrous oxide is even more potent than methane (GWP of about 265-298) and has a long atmospheric lifetime. However, its concentration is lower than methane's, placing it as the third-largest contributor, primarily from agricultural and industrial activities. 4. Chlorofluorocarbons (C. CFCs): CFCs and other fluorinated gases have extremely high GWPs (thousands of times that of CO₂). However, due to international regulations like the Montreal Protocol (aimed at protecting the ozone layer), their atmospheric concentrations are very, very low. This makes their overall contribution to global warming the smallest among the four listed.

Step 3: Now, we arrange the gases in decreasing order of their total contribution: 1st: Carbon dioxide (B) 2nd: Methane (D) 3rd: Nitrous oxide (A) 4th: Chlorofluorocarbon (C) This gives the sequence B, D, A, C.

💡 Quick Tip

Don't confuse Global Warming Potential (GWP) with overall contribution. While a molecule of methane is more potent than a molecule of CO₂, there are far more CO₂ molecules in the atmosphere, making its total effect the largest.

54. Match LIST-I with LIST-II

LIST-I (Atmospheric component)

- (A) Nitrogen
- (B) Sulphur dioxide
- (C) Aerosols
- (D) Ozone

LIST-II (Role)

- (I) Hypnoxyotic an(t)icle
- (II) Photochemical interaction
- (III) Retards re-radiation... (Greenhouse Effect)
- (IV) Acid rain
- (A) A-III, B-II, C-I, D-IV
- (B) A-I, B-III, C-II, D-IV
- (C) A-III, B-IV, C-I, D-II
- (D) A-IV, B-II, C-I, D-II

Correct Answer: (C) (Based on finding the best possible fit for the flawed options)

Solution: This matching question contains some errors, but we can arrive at the most likely intended answer by identifying the clearest, most accurate pairings first.

Step 1: Identify the most unambiguous matches.

- Sulphur dioxide (B): SO₂ is a well-known pollutant that reacts with water, oxygen, and other chemicals in the atmosphere to form sulfuric acid. This falls to the earth as Acid rain

(IV). This is a very strong and definite match. So, B matches with IV. - Ozone (D): In the troposphere (lower atmosphere), ozone is not emitted directly but is formed as a secondary pollutant through Photochemical interaction (II). Sunlight acts on nitrogen oxides and volatile organic compounds to create photochemical smog, of which ozone is a primary component. This is also a very strong match. So, D matches with II.

Step 2: Use these certain matches to eliminate incorrect options.

We are looking for an option that contains the pair B-IV and the pair D-II. - Option (A) has B-II, which is wrong. - Option (B) has B-III, which is wrong. - Option (C) has both B-IV and D-II. This is a strong candidate. - Option (D) has B-II, which is wrong.

Step 3: Evaluate the remaining pairs in the only possible option, (C).

Option (C) gives the full matching as A-III, B-IV, C-I, D-II. - A (Nitrogen) → III (Greenhouse Effect): This match is incorrect. Nitrogen (N_2), which makes up 78% of the atmosphere, is not a significant greenhouse gas. Nitrous oxide (N_2O) is, but not N_2 . - C (Aerosols) → I (Hypnoxyotic...): The text for description (I) is garbled and makes no scientific sense. Aerosols have complex roles, including scattering sunlight (cooling effect) and acting as cloud condensation nuclei, but none match this text.

Conclusion: Despite the fact that the matches for A and C are incorrect or nonsensical, option (C) is the only one that correctly pairs the two unambiguous components (B and D). Therefore, we must conclude it is the intended answer.

💡 Quick Tip

When faced with a flawed matching question, identify the most certain and unambiguous pairs first. Then, scan the options to see which one contains those correct pairs. This process of elimination can often lead to the intended answer even if other parts of the question are incorrect.

55. According to Köppen's climatic classification, _____ climate prevails in the Great Plains of India.

- (A) Af
- (B) Am
- (C) BS
- (D) Cwg

Correct Answer: (D)

Solution: Step 1: First, let's characterize the climate of the Great Plains of India, also known as the Indo-Gangetic Plain. This vast region experiences a subtropical climate with three distinct seasons: a hot, dry summer; a hot, wet monsoon season; and a mild to cool, dry winter.

Step 2: Now, let's decode the Köppen climate classification codes provided in the options to see which one best fits this description. The Köppen system uses a series of letters to classify climates. - Af (Tropical Rainforest): 'A' signifies a tropical climate (hot year-round), and 'f'

signifies sufficient rainfall in every month (no dry season). This does not fit the Great Plains. - Am (Tropical Monsoon): 'A' is tropical, and 'm' signifies a monsoon climate with a short dry season but enough rain to support a rainforest. While there is a monsoon, the climate is not tropical (it has a cool winter), so this is not the best fit. - BS (Steppe Climate): 'B' signifies a dry (arid or semi-arid) climate. This is not representative of the fertile, monsoon-fed Great Plains, although some areas to the west border on this climate type. - Cwg (Temperate climate with dry winter and hot summer): - 'C' signifies a temperate (or mesothermal) climate, which has moderate temperatures and distinct seasons, including a mild winter. - 'w' signifies a dry winter. - 'g' is a special modifier for the "Ganges-type" climate, indicating that the hottest month occurs in late spring/early summer, just before the onset of the monsoon rains. This Cwg classification perfectly describes the climate of the Great Plains of India.

Step 3: By matching the climatic characteristics of the region with the Köppen definitions, Cwg is the correct classification.

 Quick Tip

For the Köppen system in India: - Am: West coast (monsoon). - Aw: Most of the peninsula (tropical savanna). - BSh/BWh: Western arid regions (steppe/desert). - Cwg: Indo-Gangetic Plain. - E: Himalayas (polar/tundra).

56. Arrange the following in increasing order of their wavelength.

- A. gamma radiations
- B. X rays
- C. UV radiations
- D. microwave
- E. Infrared radiations

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

Correct Answer: (C)

Solution: Step 1: The question asks us to arrange different types of electromagnetic radiation in increasing order of their wavelength. This requires knowledge of the electromagnetic spectrum.

Step 2: The electromagnetic spectrum is the range of all types of electromagnetic radiation. It is typically ordered by wavelength, frequency, or energy. For this question, we need the order from shortest wavelength to longest wavelength. That order is as follows: Gamma rays → X-rays → Ultraviolet → Visible Light → Infrared → Microwaves → Radio waves. (Remember that shorter wavelength corresponds to higher frequency and higher energy).

Step 3: Now we map the given radiations onto this standard sequence. 1. The shortest wavelength radiation on the list is gamma radiations (A). 2. Next, with a slightly longer wavelength, are X rays (B). 3. Following X-rays are UV (ultraviolet) radiations (C). 4. After

the visible spectrum (which is not on the list) comes Infrared radiations (E). 5. The radiation with the longest wavelength on the list is microwave (D).

Step 4: Assembling these in order gives the sequence corresponding to increasing wavelength: A, B, C, E, D.

💡 Quick Tip

A mnemonic for the EM spectrum in order of increasing wavelength is: "Good Xylo-phones Use Very Interesting Musical Rhythms" (Gamma, X-ray, UV, Visible, Infrared, Microwave, Radio).

[Image of the electromagnetic spectrum with wavelengths]

57. Half life of a radioactive material during radioactive decay is -----.

- (A) directly proportional to the initial concentration
- (B) inversely proportional to the initial concentration and directly proportional to decay constant
- (C) directly proportional to the final concentration and inversely proportional to decay constant
- (D) independent of the initial concentration and inversely proportional to decay constant

Correct Answer: (D)

Solution: Step 1: Let's start with the definition of half-life ($t_{1/2}$). The half-life of a radioactive isotope is the time it takes for one-half of the atoms in a given sample to undergo radioactive decay.

Step 2: The process of radioactive decay follows first-order kinetics. The rate of decay is proportional to the number of undecayed nuclei present. The mathematical relationship between the half-life and the decay constant (λ), which is a measure of the probability of decay per unit time, is derived from the decay law $N(t) = N_0 e^{-\lambda t}$. Setting $N(t) = N_0/2$, we get:

$$\frac{N_0}{2} = N_0 e^{-\lambda t_{1/2}}$$

$$\frac{1}{2} = e^{-\lambda t_{1/2}}$$

Taking the natural logarithm of both sides:

$$\ln(1/2) = -\lambda t_{1/2} \implies -\ln(2) = -\lambda t_{1/2}$$

$$t_{1/2} = \frac{\ln(2)}{\lambda}$$

Step 3: Now, let's analyze this final formula, $t_{1/2} = \frac{\ln(2)}{\lambda}$. - The term $\ln(2)$ is a constant (approximately 0.693). - The formula shows that $t_{1/2}$ is inversely proportional to the decay constant λ . A larger decay constant means a faster decay and thus a shorter half-life. - Crucially, the terms for initial amount or concentration (N_0) or final amount ($N(t)$) do not

appear in the final equation. This means the half-life is a fundamental property of the isotope and is independent of the initial concentration or amount of the material.

Conclusion: Based on this analysis, the half-life is independent of the initial concentration and inversely proportional to the decay constant.

 Quick Tip

A key feature of first-order kinetics (which includes radioactive decay) is that the half-life is constant. It doesn't matter if you start with 1 kg or 1 gram of a substance; the time it takes for half of it to decay is always the same.

58. The amount of groundwater is estimated to be

- (A) nearly equal to the amount of surface water on the planet Earth
- (B) less than the amount of surface water on the planet Earth
- (C) twice the amount of surface water on the planet Earth
- (D) twenty five times the amount of surface water on the planet Earth

Correct Answer: (D) (Represents the correct order of magnitude difference)

Solution: Step 1: First, we need to consider the distribution of all water on Earth. About 97

Step 2: Now, let's look at the distribution of this small amount of freshwater. - The vast majority, about 69- Almost all of the remaining liquid freshwater, about 30- A very small fraction, only about 1.2

Step 3: The question asks to compare the amount of groundwater to the amount of surface water. Let's calculate the ratio using the percentages from Step 2.

$$\text{Ratio} = \frac{\text{Percentage of Groundwater}}{\text{Percentage of Surface Water}} = \frac{30\%}{1.2\%} = 25$$

Step 4: This calculation shows that the amount of groundwater is approximately 25 times the amount of all liquid surface freshwater combined. This is a surprisingly large difference and highlights the vastness of the groundwater resource. Therefore, option (D) is the correct statement.

 Quick Tip

Remember the freshwater distribution: most is frozen in ice caps (~69

59. Match LIST-I with LIST-II

LIST-I (Disease)

- (A) Blue baby syndrome
- (B) Keshan disease
- (C) Gas bubble disease

(D) Itai itai

LIST-II (Cause)

(I) Excess dissolved oxygen in drinking water

(II) Excess Cd in drinking water

(III) Excess nitrate ions in drinking water

(IV) Selenium deficiency

(A) A - I, B - II, C - III, D - IV

(B) A - I, B - III, C - IV, D - II

(C) A - III, B - IV, C - I, D - II

(D) A - III, B - IV, C - II, D - I

Correct Answer: (C)

Solution: Let's analyze each disease and identify its specific cause from the list.

Step 1: Analyze Blue baby syndrome (A)

This condition, known medically as methemoglobinemia, affects infants. It is characterized by a bluish discoloration of the skin. The most common environmental cause is the consumption of formula made with well water that has excess nitrate ions (III). The nitrates are converted to nitrites in the infant's gut, which then interfere with hemoglobin's ability to carry oxygen. Thus, **A matches with (III)**.

Step 2: Analyze Keshan disease (B)

Keshan disease is a serious heart condition (a cardiomyopathy) that was found to be endemic in certain parts of China. Extensive research linked it directly to a diet severely lacking in the essential trace mineral selenium. It is caused by Selenium deficiency (IV). Thus, **B matches with (IV)**.

Step 3: Analyze Gas bubble disease (C)

This is a condition that primarily affects fish and other aquatic life in aquaculture or aquariums. It is caused by the water being supersaturated with atmospheric gases, typically nitrogen but can also be oxygen. This leads to the formation of gas bubbles in the tissues and bloodstream of the fish. The cause is therefore Excess dissolved gas (I) (the option specifies oxygen, which is one possibility). Thus, **C matches with (I)**.

Step 4: Analyze Itai-itai disease (D)

This disease was first identified in Toyama Prefecture, Japan. The name means "it hurts, it hurts" in Japanese, referring to the severe bone pain it causes. It was the result of mass poisoning from excess Cadmium (Cd) (II) in the water supply, which was contaminated by mining operations. Thus, **D matches with (II)**.

Conclusion: Assembling the correct pairings gives: A-III, B-IV, C-I, D-II. This corresponds to option (C).

 Quick Tip

Associate keywords with these environmental diseases: - Blue Baby → Nitrate - Keshan → Selenium - Itai-Itai → Cadmium - Minamata → Mercury

60. Choose the correct statements.

A. Dissolution of oxygen in surface water changes with change in temperature during extreme summer and winters.

B. Dissolved oxygen and BOD have inverse relationship in sewage water.

C. Chemical oxygen demand is always higher than BOD in sewage water.

D. More saline water will have less conductivity.

(A) A and B only

(B) A, B and D only

(C) A, B, C and D

(D) A, B and C only

Correct Answer: (D)

Solution: Let's analyze each statement to determine if it is correct or incorrect.

Step 1: Analyze statement A.

The solubility of any gas in a liquid is dependent on temperature. Specifically for oxygen in water, solubility is inversely proportional to temperature. Cold water can hold more dissolved oxygen (DO) than warm water. Therefore, the amount of DO in a lake or river will be higher in cold winters and lower in hot summers. The statement is correct.

Step 2: Analyze statement B.

BOD (Biochemical Oxygen Demand) is a measure of the amount of oxygen required by aerobic bacteria to decompose the organic waste in a sample of water. Sewage water is rich in organic waste. The decomposition process consumes large amounts of DO from the water. Thus, if the BOD is high (lots of waste), the DO will be low (lots of oxygen consumed). This is an inverse relationship. The statement is correct.

Step 3: Analyze statement C.

BOD measures only the oxygen needed to break down the biodegradable organic matter. COD (Chemical Oxygen Demand) measures the oxygen needed to break down all organic matter (both biodegradable and non-biodegradable) using a strong chemical oxidant. Since sewage contains both types of organic matter, the amount measured by COD will always be greater than or equal to the amount measured by BOD. The statement says COD is higher, which is generally true for complex wastewater like sewage. The statement is correct.

Step 4: Analyze statement D.

Electrical conductivity in water is a measure of its ability to conduct an electric current. This ability is directly dependent on the concentration of dissolved ions (salts). Saline water, by definition, has a high concentration of dissolved salts. Therefore, more saline water will have more dissolved ions and thus higher electrical conductivity. The statement claims it will have less conductivity, which is incorrect.

Conclusion: Statements A, B, and C are correct, while D is incorrect. This means option (D) is the right choice.

💡 Quick Tip

Remember these key water quality relationships: - Temperature & DO: Inverse (Hot water, less oxygen). - BOD & DO: Inverse (High waste, less oxygen). - COD vs BOD: $COD \geq BOD$ (COD measures more stuff). - Salinity & Conductivity: Direct (More salt, more conductive).

61. Which of the following is dominantly responsible for buffering capacity of natural water?

- (A) carbonate ions
- (B) dissolved oxygen
- (C) bicarbonate ions
- (D) phosphate ions

Correct Answer: (C)

Solution: Step 1: First, we must define "buffering capacity". The buffering capacity of a solution, in this case natural water, is its ability to resist changes in pH when an acid or a base is added. A good buffer system can neutralize small amounts of added acid or base, thus maintaining a relatively stable pH.

Step 2: Next, we identify the chemical system that provides this buffering in most natural waters like rivers, lakes, and oceans. This is the carbonate-bicarbonate buffer system. This system is established by the dissolution of atmospheric carbon dioxide (CO_2) in water and its interaction with carbonate rocks (like limestone).

Step 3: Let's look at the chemical equilibria involved: CO_2 (dissolved) + $H_2O \rightleftharpoons H_2CO_3$ (carbonic acid) $H_2CO_3 \rightleftharpoons H^+ + HCO_3^-$ (bicarbonate ion) $HCO_3^- \rightleftharpoons H^+ + CO_3^{2-}$ (carbonate ion)

Step 4: The effectiveness of a buffer depends on the relative concentrations of the acid and its conjugate base. In the typical pH range of most natural surface waters (around 6.5 to 8.5), the bicarbonate ion (HCO_3^-) is the most abundant species in this equilibrium system. - If an acid (H^+) is added, bicarbonate ions will react with it: $HCO_3^- + H^+ \rightarrow H_2CO_3$. This removes the added acid. - If a base (OH^-) is added, bicarbonate can donate a proton: $HCO_3^- + OH^- \rightarrow CO_3^{2-} + H_2O$. This neutralizes the added base. Because it is the most prevalent species and can react with both acids and bases in this pH range, the bicarbonate ion is the dominant component responsible for the buffering capacity.

💡 Quick Tip

For water chemistry, the carbonate-bicarbonate system is key. Think of bicarbonate (HCO_3^-) as the main workhorse that keeps the pH of rivers and lakes stable.

62. Soil acidity can be introduced by -----.

- A. Mine tailings rich in pyrites
- B. Extraction of Al from its ore
- C. Use of nitrogen-rich fertilizers
- D. Use of limestone

- (A) A, B and D only
- (B) A and C only
- (C) A, B and C only
- (D) A, B, C and D

Correct Answer: (B)

Solution: The question asks us to identify the processes from the list that can cause or introduce soil acidity (i.e., lower the soil's pH).

Step 1: Analyze statement A.

Mine tailings are waste materials left over from mining. Pyrite FeS_2 , also known as fool's gold, is a sulfide mineral often found with coal and metal ores. When pyrite is exposed to air and water, it oxidizes to form sulfuric acid (H_2SO_4), a very strong acid. This process is the primary cause of acid mine drainage, which severely acidifies soil and water. So, A is a cause of acidity.

Step 2: Analyze statement B.

The industrial extraction of aluminum (Al) from its bauxite ore is typically done using the Bayer process. This process involves digesting the crushed ore in a hot solution of sodium hydroxide (NaOH), which is a strong base (alkaline). The major waste product, known as red mud, is therefore highly alkaline, with a pH of 10 to 13. This process does not cause acidity; it produces alkaline waste. So, B is not a cause of acidity.

Step 3: Analyze statement C.

Many common nitrogen-rich fertilizers are based on ammonium (NH_4^+). When these fertilizers are applied to the soil, a natural process called nitrification occurs. Soil bacteria convert the ammonium into nitrate (NO_3^-). This biochemical reaction releases hydrogen ions (H^+) into the soil: $\text{NH}_4^+ + 2\text{O}_2 \rightarrow \text{NO}_3^- + \text{H}_2\text{O} + 2\text{H}^+$. The release of H^+ ions increases the soil's acidity. So, C is a cause of acidity.

Step 4: Analyze statement D.

Limestone is primarily calcium carbonate (CaCO_3). It is an alkaline substance. Farmers and land managers apply limestone to fields specifically to neutralize existing soil acidity and raise the soil pH. It is a treatment for acidity, not a cause. So, D is not a cause of acidity.

Conclusion: The processes that introduce soil acidity are A (mine tailings with pyrite) and C (use of nitrogen-rich fertilizers). Therefore, the correct option is (B).

💡 Quick Tip

To remember causes of soil acidity, think: - Acid Rain: From industrial pollutants (SO_x , NO_x). - Mining: Oxidation of sulfides like pyrite. - Fertilizers: Nitrification of ammonium. Limestone (lime) is the cure, not the cause.

63. Match LIST I with LIST II

LIST I (Metal)

(A) Manganese

(B) Zinc

(C) Aluminium

(D) Lithium

LIST II (Ore)

(I) Sinsilfluoride (typo)

(II) Gibbsite

(III) Lepidolite

(IV) Psilomelane

(A) A-I, B-II, C-III, D-IV

(B) A-I, B-III, C-II, D-IV

(C) A-IV, B-I, C-II, D-III

(D) A-III, B-IV, C-I, D-II

Correct Answer: (C)

Solution: Let's match each metal with its correct ore. It's often easiest to start with the most well-known pairings.

Step 1: Analyze Aluminium (C)

The principal ore of Aluminium is a rock called bauxite. Bauxite is primarily composed of aluminium hydroxide minerals, the most important of which is Gibbsite (II). This is a very common and strong association. Thus, **C matches with (II)**.

Step 2: Analyze Lithium (D)

Lithium is a light metal extracted from several minerals. One of the most important commercial sources of lithium is Lepidolite (III), which is a lilac-grey or rose-colored member of the mica group of minerals. Thus, **D matches with (III)**.

Step 3: Analyze Manganese (A)

Manganese is a key industrial metal. One of its major ore minerals is Psilomelane (IV), which is a hard, black manganese oxide. Thus, **A matches with (IV)**.

Step 4: Analyze Zinc (B)

We have now made three confident matches: A-IV, C-II, and D-III. The only remaining items are Zinc (B) and the ore "Sinsilfluoride" (I). "Sinsilfluoride" is not a recognized mineral name and is certainly a typo. However, by the process of elimination, Zinc must be paired with this remaining item. Thus, **B matches with (I)**.

Conclusion: Assembling our matches gives the sequence: A-IV, B-I, C-II, D-III. This corresponds exactly to option (C).

💡 Quick Tip

Focus on the definite pairs first when matching: - Aluminium → Bauxite/Gibbsite - Lithium → Lepidolite/Spodumene - Manganese → Psilomelane/Pyrolusite - Iron → Hematite/Magnetite

64. Arrange the following soil horizons as they occur from top to bottom in a typical soil profile.

A. A horizon

B. B horizon

C. R horizon

D. O horizon

E. C horizon

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

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

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

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

Correct Answer: (B)

Solution: Step 1: A soil profile is a vertical cross-section of the soil, from the ground surface down to the underlying parent rock. It is divided into distinct layers called horizons. We need to arrange the given horizons in the correct sequence as they appear from the top down.

Step 2: Let's define each master horizon in its natural order of appearance. 1. O horizon (D): This is the uppermost layer, found on the surface. It consists primarily of organic matter, such as decomposing leaves and humus. 'O' stands for Organic. 2. A horizon (A): Located just below the O horizon, this is the topsoil. It is a mixture of mineral particles and well-decomposed organic matter (humus). It is a zone of leaching, where water washes soluble minerals downwards. 3. B horizon (B): This is the subsoil, found beneath the A horizon. It is a zone of accumulation (or illuviation), where the minerals and clays leached from the A horizon are deposited. It typically has a lighter color and less organic matter than the topsoil. 4. C horizon (E): This layer lies below the B horizon and consists of partially weathered or broken-up parent material (rock fragments). It is the transition zone between the true soil above and the solid bedrock below. 5. R horizon (C): This is the bottommost layer, consisting of unweathered, consolidated bedrock. 'R' stands for Rock.

Step 3: Now, let's assemble the correct sequence using the letters provided in the question, arranged from top to bottom. The order is: O horizon (D), A horizon (A), B horizon (B), C horizon (E), R horizon (C). This gives the final letter sequence: D, A, B, E, C.

💡 Quick Tip

A simple mnemonic for the main soil horizons from top to bottom is "Old Ants Eat Big Crumbs of Rock". (O, A, E, B, C, R). Note that the 'E' horizon (zone of eluviation) is often present between A and B but was not an option here.

65. Choose the correct statement.

(A) Excessive nutrient inputs to lakes maintains the high dissolved oxygen of the lakes.

(B) Lakes with extensive algal blooms will have the good health of the lakes.

(C) Eutrophic lakes have high BOD.

(D) Oligotrophic lakes have high BOD compared to eutrophic lakes.

Correct Answer: (C)

Solution: Let's analyze each statement about lake ecology and water quality.

Step 1: Analyze statement A.

The process of adding excessive nutrients (like nitrogen and phosphorus) to a lake is called eutrophication. These nutrients fuel massive population explosions of algae, known as algal blooms. When this large mass of algae dies, it sinks to the bottom and is decomposed by aerobic bacteria. This decomposition process consumes vast amounts of dissolved oxygen, often leading to hypoxic (low oxygen) or anoxic (no oxygen) conditions. Therefore, excessive nutrients lead to low dissolved oxygen, not high. This statement is incorrect.

Step 2: Analyze statement B.

Extensive algal blooms are a primary symptom of eutrophication. This condition is a sign of an unbalanced and unhealthy ecosystem. It can lead to fish kills due to lack of oxygen, reduce water clarity, and sometimes involve toxic algae species. Therefore, a lake with extensive algal blooms is in a state of poor health. This statement is incorrect.

Step 3: Analyze statement C.

A eutrophic lake is one that is rich in nutrients and has high biological productivity (lots of algae and plant life). As explained in Step 1, the death and decay of this large amount of organic matter create a high demand for oxygen from decomposer bacteria. This demand is measured as Biochemical Oxygen Demand (BOD). Therefore, eutrophic lakes are characterized by a high BOD. This statement is correct.

Step 4: Analyze statement D.

An oligotrophic lake is the opposite of a eutrophic one. It is nutrient-poor, has low biological productivity, clear water, and very little organic matter to decompose. Consequently, its BOD is very low. The statement claims oligotrophic lakes have a higher BOD than eutrophic lakes, which is the exact opposite of the truth. This statement is incorrect.

💡 Quick Tip

- Oligotrophic: "Oligo" = few. Low nutrients, low algae, low BOD, high clarity, high oxygen. Good health. - Eutrophic: "Eu" = true/good (at growing things). High nutrients, high algae, high BOD, low clarity, low oxygen. Poor health.

66. Which one of the following statements is incorrect?

- (A) Nitrogen is one of the top five elements found in plants and animals.
- (B) Industrial and biological nitrogen fixation require the input of substantial amount of energy.
- (C) Nitrogen in organisms is only found in nucleic acids, i.e. DNA and RNA.
- (D) Nitrogen is plentiful in the earth's atmosphere in relatively inert forms.

Correct Answer: (C)

Solution: The task is to find the statement that is false. Let's evaluate each one.

Step 1: Analyze statement A.

The six most abundant elements in living organisms, often remembered by the mnemonic CHNOPS, are Carbon (C), Hydrogen (H), Nitrogen (N), Oxygen (O), Phosphorus (P), and Sulfur (S). Nitrogen is a crucial component of all life. This statement is correct.

Step 2: Analyze statement B.

Atmospheric nitrogen exists as the N_2 molecule, which is held together by a very strong triple covalent bond. Breaking this bond to "fix" the nitrogen into a usable form (like ammonia, NH_3) is an extremely difficult and energy-intensive process. The industrial Haber-Bosch process requires high temperatures and pressures. Similarly, biological nitrogen fixation carried out by certain bacteria requires a huge amount of metabolic energy (ATP). This statement is correct.

Step 3: Analyze statement C.

Nitrogen is a fundamental building block for several essential types of biological macromolecules. While it is true that nitrogen is found in the nitrogenous bases of nucleic acids (DNA and RNA), it is also a defining component of all amino acids. Since amino acids are the monomers that build proteins, nitrogen is essential for every protein in an organism. Proteins perform countless functions, from acting as enzymes to forming structural components like muscle and hair. The statement that nitrogen is only found in nucleic acids is therefore false. This statement is incorrect.

Step 4: Analyze statement D.

The Earth's atmosphere is composed of approximately 78

Conclusion: Statement (C) is the incorrect one because it overlooks the critical role of nitrogen in proteins.

💡 Quick Tip

Remember the two main biological macromolecules containing nitrogen: Proteins (via amino acids) and Nucleic Acids (via nitrogenous bases). Any statement limiting nitrogen to just one of these is incorrect.

67. Condensation of water vapors in atmosphere actually -----.

- (A) cools the surrounding
- (B) sometime cools and sometimes warms the surrounding
- (C) warms the surrounding
- (D) neither cools nor warms the surrounding

Correct Answer: (C)

Solution: Step 1: This question relates to the energy changes that occur during a change of state (phase change) of water. Specifically, it asks about the process of condensation, which is the transition from a gas (water vapor) to a liquid (water droplets).

Step 2: To understand the energy flow, let's consider the reverse process: evaporation. To turn liquid water into water vapor, energy must be added to the water molecules to overcome

the intermolecular forces holding them together as a liquid. This required energy is called the latent heat of vaporization. This energy is absorbed from the surroundings, which is why evaporation is a cooling process (e.g., when sweat evaporates from your skin, it cools you down).

Step 3: Condensation is the exact opposite of evaporation. For water vapor molecules to come together and form a liquid, they must release the extra energy they possess as a gas. This energy, the latent heat of condensation, is released into the surrounding environment (the air).

Step 4: Therefore, when water vapor in the atmosphere condenses to form clouds or dew, it releases a significant amount of latent heat. This release of heat energy warms the surrounding air. This process is a major driver of energy transfer in the atmosphere and is crucial for the development of storms and hurricanes.

 Quick Tip

Think about how you feel when you get out of a pool. The evaporating water on your skin takes heat from your body, making you feel cold. Condensation is the exact opposite process; it releases heat.

68. The Prime Minister of India introduced mission LiFE in UN climate change conference of parties (COP) ----- held at ----- in the year -----.

- (A) COP-26, Paris, 2021
- (B) COP-26, Glasgow, 2020
- (C) COP-29, Dubai, 2022
- (D) COP-26, Glasgow, 2021

Correct Answer: (D)

Solution: Step 1: The question asks for the specific details of the launch of India's Mission LiFE (Lifestyle for Environment) initiative on the international stage.

Step 2: We need to identify the correct Conference of the Parties (COP) number. The concept was introduced by Prime Minister Narendra Modi at COP26. This eliminates option (C).

Step 3: We need to identify the host city for this conference. COP26 was held in Glasgow, Scotland. This eliminates option (A).

Step 4: Finally, we need to identify the correct year. COP26 in Glasgow took place from 31 October to 13 November 2021. (It was originally scheduled for 2020 but was postponed due to the COVID-19 pandemic). This eliminates option (B).

Conclusion: Combining the correct information, Mission LiFE was introduced at COP26, held in Glasgow, in the year 2021. This matches option (D) perfectly.

 Quick Tip

Associate "Mission LiFE" with India's major announcement at the Glasgow COP26 climate summit in 2021.

69. Which of the following are criteria pollutants under national ambient air quality standards?

- A. PM_{2.5}
- B. NO_x
- C. Lead
- D. Methane
- E. Sulphur dioxide

- (A) A, B and D only
- (B) A, B and C only
- (C) A, C and E only
- (D) A, B, C and E only

Correct Answer: (D)

Solution: Step 1: First, we must understand what "criteria pollutants" are. These are a group of common air pollutants regulated by national environmental agencies because they are widespread and can cause significant harm to public health and the environment. Different countries may have slightly different lists, but a core group is internationally recognized. The question refers to India's National Ambient Air Quality Standards (NAAQS).

Step 2: Let's review the list of pollutants provided in the question and check their status under the NAAQS. - A. PM_{2.5} (Particulate Matter 2.5): These are fine, inhalable particles with diameters of 2.5 micrometers and smaller. They are a major health concern and are definitely a criteria pollutant. - B. NO_x (Oxides of Nitrogen): This is a group of gases, primarily nitrogen dioxide (NO₂). They contribute to smog, acid rain, and respiratory problems. NO₂ is a criteria pollutant. - C. Lead (Pb): Lead is a toxic heavy metal that can be released into the air from industrial sources and historical use in gasoline. It is a criteria pollutant due to its severe health effects. - D. Methane (CH₄): Methane is a potent greenhouse gas and is a major focus for climate change mitigation. However, it is not typically regulated as a criteria air pollutant for direct health effects under ambient air quality standards. - E. Sulphur dioxide (SO₂): This gas is primarily released from burning fossil fuels containing sulfur, like coal. It contributes to respiratory illness and acid rain. It is a classic criteria pollutant.

Step 3: Based on the analysis, the substances from the list that are considered criteria pollutants under NAAQS are A (PM_{2.5}), B (NO_x), C (Lead), and E (Sulphur dioxide).

Conclusion: The correct combination of criteria pollutants from the list is A, B, C, and E.

💡 Quick Tip

The main six criteria pollutants focused on globally are CO, Pb, NO₂, O₃, PM (PM_{2.5}/PM₁₀), and SO₂. Methane is usually discussed in the context of climate change, not direct air quality standards.

70. Which one of the following states in India is known to have the highest Chromium pollution of water and soils due to tanning industry?

- (A) Uttar Pradesh
- (B) Andhra Pradesh
- (C) Tamil Nadu
- (D) Maharashtra

Correct Answer: (C)

Solution: Step 1: First, we need to understand the connection between the tanning industry and chromium pollution. The most common method of leather tanning is chrome tanning, which uses chromium salts (particularly chromium sulfate) to stabilize the collagen fibers in animal hides. The wastewater (effluent) from these tanneries, if not properly treated, is heavily contaminated with chromium, including the highly toxic hexavalent chromium (Cr(VI)).

Step 2: Next, we need to identify the major centers of the leather tanning industry in India. Two of the most significant hubs are: - Kanpur in Uttar Pradesh, located on the banks of the river Ganga. - A large cluster of tanneries in Tamil Nadu, concentrated in the Palar river basin in cities like Vellore, Ranipet, and Ambur.

Step 3: Now, we must evaluate which of these regions is most famously associated with the highest chromium pollution. While the Kanpur cluster has caused severe pollution of the Ganga, the sheer density of tanneries in the Palar river basin in Tamil Nadu has made it a globally recognized case study for extreme and widespread chromium contamination of both surface water and groundwater. Numerous scientific studies and reports have highlighted the severe environmental degradation and health problems in this specific region due to decades of uncontrolled effluent discharge from the tanning industry. Therefore, Tamil Nadu is the state most prominently known for this specific issue.

💡 Quick Tip

When thinking about industrial pollution in India, associate: - Tanneries/Chromium: Tamil Nadu (Vellore), Uttar Pradesh (Kanpur) - Textile Dyes: Tamil Nadu (Tiruppur) - Mercury: Kodaikanal (from a specific factory)

71. Arrange the following layers of atmosphere starting from earth's surface.

- A. Troposphere
- B. Stratosphere

C. Tropopause

D. Stratopause

E. Mesosphere

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

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

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

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

Correct Answer: (C)

Solution: Step 1: The Earth's atmosphere is structured in several layers, defined by their temperature profiles. The question asks us to arrange the given layers and boundaries in order, starting from the ground and moving upwards.

Step 2: Let's identify the main layers and their boundaries ('pauses'). The 'pause' is the upper boundary of the layer below it. 1. We start at the Earth's surface in the Troposphere (A). This is the lowest layer, where all weather occurs. 2. At the top of the troposphere is a boundary layer called the Tropopause (C). Here, the temperature stops decreasing with altitude. 3. Above the tropopause is the Stratosphere (B). This layer contains the ozone layer, and the temperature increases with altitude here. 4. At the top of the stratosphere is its boundary, the Stratopause (D). 5. Above the stratopause lies the Mesosphere (E), where temperature again decreases with altitude.

Step 3: Now, let's assemble the correct sequence from the surface upward using the provided letters. The order is: Troposphere (A), followed by its boundary Tropopause (C), then Stratosphere (B), followed by its boundary Stratopause (D), and finally Mesosphere (E).

Conclusion: The correct sequence is A, C, B, D, E.

💡 Quick Tip

A mnemonic for the layers is "Trust Me in The Exam": Troposphere, Stratosphere, Mesosphere, Thermosphere, Exosphere. The 'pauses' are the ceilings of the layers below them (e.g., Tropopause is the top of the Troposphere).

[Image of Earth's atmospheric layers]

72. Match LIST-I with LIST-II

LIST-I (Person)

(A) Anna Hazare

(B) Rajender Singh

(C) Chandi Prasad Bhatt

(D) Medha Patkar

LIST-II (Movement)

(I) Narmada Bachao Andolan

(II) Tarun Bharat Sangh

(III) Chipko movement

(IV) Ralegan Siddhi Watershed Development

- (A) A-IV, B-II, C-III, D-I
- (B) A-I, B-III, C-II, D-IV
- (C) A-I, B-II, C-IV, D-III
- (D) A-III, B-IV, C-I, D-II

Correct Answer: (A)

Solution: Let's match each prominent Indian activist with their most well-known environmental or social movement.

Step 1: Analyze Anna Hazare (A)

Before gaining national fame for his anti-corruption movement, Anna Hazare was renowned for his pioneering work in his own village, Ralegan Siddhi, in Maharashtra. He led a grassroots movement to transform the drought-prone village through water conservation and Ralegan Siddhi Watershed Development (IV). Thus, **A matches with (IV)**.

Step 2: Analyze Rajender Singh (B)

Rajender Singh is a celebrated water conservationist from Rajasthan, widely known as the "waterman of India". He is the founder of the NGO Tarun Bharat Sangh (II), which has revived rivers and brought water back to thousands of villages using traditional water harvesting techniques. Thus, **B matches with (II)**.

Step 3: Analyze Chandi Prasad Bhatt (C)

Chandi Prasad Bhatt is a Gandhian environmentalist and social activist. He was one of the key pioneers and leaders of the world-famous Chipko movement (III) in the 1970s in Uttarakhand, where villagers hugged trees to prevent them from being felled. Thus, **C matches with (III)**.

Step 4: Analyze Medha Patkar (D)

Medha Patkar is a prominent social activist and the most recognizable face of the Narmada Bachao Andolan (I) (Save the Narmada Movement). This mass movement has for decades protested against the large dams being built on the Narmada River, raising issues of displacement and environmental impact. Thus, **D matches with (I)**.

Conclusion: Assembling the matches gives the sequence: A-IV, B-II, C-III, D-I. This corresponds to option (A).

💡 Quick Tip

Associate these key figures with their movements: - Medha Patkar → Narmada River
- Chandi Prasad Bhatt / Sunderlal Bahuguna → Chipko (Hugging Trees) - Rajender Singh → Water / Tarun Bharat Sangh - Anna Hazare → Ralegan Siddhi / Anti-corruption

73. Match LIST-I with LIST-II

LIST-I (Type of Rock)

- (A) Non-foliated metamorphic rock
- (B) Foliated metamorphic rock

(C) Igneous rock

(D) Sedimentary rock

LIST-II (Characteristics/example)

(I) Dolomite

(II) Rhyolite

(III) Ornate Gneiss (Gneiss)

(IV) Quartzite

(A) A-II, B-IV, C-III, D-I

(B) A-I, B-III, C-IV, D-II

(C) A-IV, B-III, C-II, D-I

(D) A-IV, B-II, C-III, D-I

Correct Answer: (C)

Solution: Let's analyze each rock type and match it with a correct example or characteristic from the list.

Step 1: Analyze Sedimentary rock (D)

Sedimentary rocks are formed from the accumulation and cementation of sediments. Dolomite (I) (also known as dolostone) is a type of sedimentary carbonate rock, similar to limestone.

This is a direct match. Thus, **D matches with (I)**.

Step 2: Analyze Igneous rock (C)

Igneous rocks are formed from the cooling and solidification of molten rock (magma or lava).

Rhyolite (II) is an extrusive (volcanic) igneous rock, the felsic equivalent of granite. This is a direct match. Thus, **C matches with (II)**.

Step 3: Analyze Foliated metamorphic rock (B)

Metamorphic rocks are formed when existing rocks are changed by heat and pressure.

Foliated metamorphic rocks display a layered or banded appearance due to the parallel alignment of mineral grains. Gneiss (III) is a classic example of a high-grade, foliated metamorphic rock characterized by its distinct banding of light and dark minerals. Thus, **B matches with (III)**.

Step 4: Analyze Non-foliated metamorphic rock (A)

Non-foliated metamorphic rocks do not have a layered appearance because their mineral grains are not aligned. They are typically formed from parent rocks with minerals that don't have a specific orientation. Quartzite (IV), which is metamorphosed sandstone, is a prime example. Its quartz grains recrystallize into a dense, interlocking structure without any layering. Thus, **A matches with (IV)**.

Conclusion: Assembling the matches gives the sequence: A-IV, B-III, C-II, D-I. This corresponds to option (C).

💡 Quick Tip

Remember key examples for each rock type: - Igneous: Granite (intrusive), Basalt/Rhyolite (extrusive). - Sedimentary: Sandstone, Limestone, Shale, Dolomite. - Metamorphic: Marble (from limestone), Quartzite (from sandstone), Slate (from shale), Gneiss (foliated).

74. Mohorovicic discontinuity is commonly defined as the depth at which -----.

- (A) the p-wave velocity decreases to 5.6 km/sec
- (B) the p-wave velocity exceeds 7.6 km/sec
- (C) the s-wave velocity exceeds 3.57 km/sec
- (D) the p and s waves exceed 5.6 km/sec and 3.57 km/sec respectively

Correct Answer: (B)

Solution: Step 1: First, we must define the Mohorovičić Discontinuity, commonly shortened to the Moho. The Moho is the boundary that separates the Earth's crust from the underlying mantle.

Step 2: This boundary is not defined by a change in chemical composition that we can see, but by a change in physical properties that we detect using seismic waves from earthquakes. The mantle is composed of rock that is significantly denser and more rigid than the rock of the crust.

Step 3: Because seismic waves travel faster through denser, more rigid materials, there is a sharp, abrupt increase in the velocity of both P-waves (primary waves) and S-waves (secondary waves) as they pass from the crust down into the mantle. This sudden jump in velocity is the defining characteristic of the Moho.

Step 4: Now, let's examine the specific velocity values. In the lower part of the continental crust, P-wave velocities are typically in the range of 6.7 to 7.2 km/s. When the waves cross the Moho into the upper mantle, their velocity abruptly increases to over 7.6 km/s, typically landing in the 8.0 to 8.2 km/s range. - Option (A) describes a decrease in velocity, which is incorrect. - Option (B) correctly identifies that the P-wave velocity sharply exceeds 7.6 km/sec at this boundary. - Options (C) and (D) provide S-wave velocities or lower P-wave velocities that are more characteristic of the crust itself, not the jump at the Moho.

Conclusion: The Moho is defined by the depth at which the P-wave velocity suddenly increases to more than 7.6 km/s.

 Quick Tip

Moho = Boundary between Crust and Mantle. It's defined by seismic waves suddenly speeding up. Specifically, P-waves jump from ~7 km/s to ~8 km/s.

75. Ozone in the stratosphere captures ----- and protects us from harmful effects.

- (A) UV A, UV B and UV C radiations completely
- (B) UV A and UV B radiations completely
- (C) UV B partially and UV C radiations completely
- (D) UV A completely and UV B and UV C radiations partially

Correct Answer: (C)

Solution: Step 1: The sun emits ultraviolet (UV) radiation, which is categorized into three types based on wavelength. From longest to shortest wavelength, they are UV-A, UV-B, and UV-C. Shorter wavelengths carry more energy and are generally more harmful to living organisms.

Step 2: The ozone layer, located in the stratosphere, plays a critical role in absorbing this incoming UV radiation. However, its effectiveness varies significantly with the wavelength. Let's analyze its effect on each type: - UV-C (shortest wavelength, highest energy): This is the most damaging type of UV radiation. The stratospheric ozone layer (along with diatomic oxygen, O₂) is extremely efficient at absorbing these wavelengths. Essentially 100- UV-B (medium wavelength): This radiation is also very harmful and is the primary cause of sunburn, skin cancer, and cataracts. The ozone layer absorbs a large portion of the incoming UV-B, but not all of it. Approximately 95- UV-A (longest wavelength, lowest energy): This type of UV radiation is the least affected by the ozone layer. The vast majority (about 95

Step 3: Based on this analysis, the most accurate description of the ozone layer's function is that it absorbs UV-C completely and UV-B partially. This matches option (C).

 Quick Tip

Remember the ozone layer's effectiveness: - UV-C: Completely blocked. - UV-B: Mostly Blocked (partially absorbed). - UV-A: Allows through (mostly unabsorbed).