

CUET PG 2025 Physics 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
3. Use of any electronic appliances is strictly prohibited.

1. The lattice constant of a simple cubic lattice having interplanar spacing 3Å for (002) plane is:

- (1) 4.2Å
- (2) 6.0Å
- (3) 6.2Å
- (4) 4.0Å

Correct Answer: (2) 6.0Å

Solution: Step 1: To begin, we need to utilize the established mathematical relationship that connects the interplanar spacing, denoted as d_{hkl} , to the lattice constant, a , for a cubic crystal structure. This relationship is articulated by the following formula:

$$d_{hkl} = \frac{a}{\sqrt{h^2 + k^2 + l^2}}$$

In this equation, the variables (h, k, l) represent the Miller indices which uniquely identify a specific family of planes within the crystal lattice.

Step 2: Next, we must carefully identify and list all the numerical values provided in the problem statement. - The interplanar spacing for the specified plane is given as $d = 3\text{Å}$. - The Miller indices that define the crystallographic plane in question are $(hkl) = (002)$.

Step 3: With the formula and the given values in hand, we can now proceed to calculate the lattice constant, a . We substitute the known values of d , h , k , and l into the interplanar spacing formula:

$$3\text{Å} = \frac{a}{\sqrt{0^2 + 0^2 + 2^2}}$$

Performing the calculation within the square root gives:

$$3\text{Å} = \frac{a}{\sqrt{4}}$$

Simplifying the denominator yields:

$$3\text{Å} = \frac{a}{2}$$

Finally, to isolate a , we rearrange the equation by multiplying both sides by 2:

$$a = 2 \times 3\text{Å} = 6\text{Å}$$

From this calculation, we can confidently conclude that the lattice constant of the simple cubic lattice is 6.0Å .

Quick Tip

For any cubic crystal system (Simple, BCC, or FCC), the formula for interplanar spacing $d = a/\sqrt{h^2 + k^2 + l^2}$ remains the same. Memorizing this single formula is sufficient for all cubic lattice problems involving Miller indices and lattice parameters.

2. In a semiconductor, intrinsic concentration of charge carriers varies with:

- (1) $T^{1/2}$

- (2) T
- (3) $T^{3/2}$
- (4) $T^{-1/2}$

Correct Answer: (3) $T^{3/2}$

Solution: Step 1: To address this question, we must first bring to mind the standard theoretical formula used to describe the intrinsic carrier concentration, denoted by n_i , in a semiconductor. This concentration is fundamentally dependent on the absolute temperature, T , and the material's energy bandgap, E_g . The well-established equation is:

$$n_i = AT^{3/2} \exp\left(-\frac{E_g}{2k_B T}\right)$$

Here, A represents a constant that is specific to the semiconductor material, and k_B is the universal Boltzmann constant.

Step 2: Now, let's carefully dissect how the intrinsic concentration n_i is influenced by temperature according to this formula. The dependence on temperature T manifests in two distinct parts of the equation: firstly, through the pre-exponential power-law term, $T^{3/2}$, and secondly, through the exponential term, $\exp(-E_g/2k_B T)$. While the exponential term accounts for the most dramatic and dominant change in carrier concentration as temperature varies, the question specifically asks for the nature of the variation of the concentration with T . The $T^{3/2}$ factor, which arises from the temperature dependence of the effective density of states in the conduction and valence bands, is an integral and inseparable part of this overall relationship.

Step 3: The final step is to compare the mathematical form of the temperature dependence found in our formula with the choices provided in the question. The options are presented as various power-law relationships with temperature T . Upon inspection of the intrinsic carrier concentration equation, we can see that the pre-exponential factor, which describes a key aspect of the variation, is precisely $T^{3/2}$. This form perfectly matches the one presented in option (3).

Quick Tip

While the exponential term $\exp(-E_g/2k_B T)$ causes the most significant change in carrier concentration with temperature, the pre-exponential $T^{3/2}$ term is also a fundamental part of the relationship derived from the density of states. Always check if this term is among the options.

3. Brillouin zone is:

- A. Wigner-Seitz cell of reciprocal lattice
- B. Primitive unit cell
- C. The locus of all k-values in the reciprocal lattice which are Bragg reflected.
- D. Wigner-Seitz cell of direct lattice

The correct statements are:

- (1) A, B and D only

- (2) A, B and C only
- (3) A, B, C and D
- (4) B, C and D only

Correct Answer: (2) A, B and C only

Solution: Step 1: Let's start by establishing a clear definition of the First Brillouin Zone. This is a conceptually crucial construct within the field of solid-state physics. It serves as a uniquely defined primitive cell in the reciprocal lattice, which is essential for understanding the behavior of electron waves as they propagate through the periodic potential of a crystal.

Step 2: Now, we will systematically assess the validity of each of the four statements provided. - **A. Wigner-Seitz cell of reciprocal lattice:** This statement presents the formal and most precise definition of the First Brillouin Zone. The Wigner-Seitz cell is constructed by taking a lattice point and identifying the region of space that is closer to that point than to any other lattice point. When this construction procedure is applied to the reciprocal lattice, the resulting cell is, by definition, the First Brillouin Zone. Therefore, this statement is correct. - **B. Primitive unit cell:** The First Brillouin Zone, due to its construction as a Wigner-Seitz cell of the reciprocal lattice, inherently possesses the properties of a primitive unit cell. This means it is a minimum-volume cell that, when translated by all reciprocal lattice vectors, completely fills the reciprocal space without any overlap or gaps. As such, this statement is correct. - **C. The locus of all k -values...which are Bragg reflected:** The boundaries of the Brillouin Zone are formed by planes that perpendicularly bisect the reciprocal lattice vectors connecting the origin to its nearest neighbors. These boundary planes represent the specific set of wave vectors (k -values) that satisfy the Laue condition for Bragg diffraction, which is given by $2\mathbf{k} \cdot \mathbf{G} = |\mathbf{G}|^2$. Consequently, the zone's boundaries define where Bragg reflection first occurs. The statement, by linking the zone's definition to the condition of Bragg reflection, is conceptually sound and therefore correct. - **D. Wigner-Seitz cell of direct lattice:** This statement is fundamentally incorrect. The Wigner-Seitz cell of the *direct* (or real space) lattice is a primitive cell located in real physical space. The Brillouin Zone, in stark contrast, is a concept that exists exclusively in the abstract mathematical space known as the reciprocal lattice (or k -space).

Step 3: After evaluating each statement, we can form a conclusion about which ones are accurate. Statements A, B, and C all provide correct descriptions or properties of the Brillouin zone. Statement D, however, is incorrect. Based on this analysis, the option that correctly groups only the true statements is the one containing A, B, and C.

💡 Quick Tip

The key to understanding the Brillouin zone is to remember that it lives in "reciprocal space" or " k -space," not real space. It is the Wigner-Seitz cell of the reciprocal lattice, and its boundaries are directly related to the condition for Bragg diffraction.

4. Arrange the following crystal structures in ascending order of their coordination number.

- A. Diamond
- B. Sodium Chloride

C. Cesium Chloride

D. Zinc with hexagonal closed packed structure

Choose the **CORRECT** answer from the options given below:

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

Correct Answer: (1) A, B, C, D

Solution: Step 1: The initial task is to identify the coordination number associated with each of the specified crystal structures. The coordination number is defined as the count of the immediate nearest atomic or ionic neighbors surrounding a central atom or ion within the crystal lattice. - **A. Diamond:** In the diamond cubic lattice structure, every carbon atom is covalently bonded to four other carbon atoms. These neighbors are positioned at the vertices of a tetrahedron, with the central atom at its center. Consequently, the coordination number for the diamond structure is 4. - **B. Sodium Chloride (NaCl):** The NaCl crystal adopts the rock salt structure. In this arrangement, any given ion (for instance, a Na^+ ion) is surrounded by six ions of the opposite charge (Cl^-). These six neighbors are located at the vertices of an octahedron. Therefore, the coordination number is 6. - **C. Cesium Chloride (CsCl):** In the CsCl crystal structure, each ion (e.g., Cs^+) is situated at the center of a cubic unit cell and is surrounded by eight ions of the opposite charge (e.g., Cl^-) located at the corners of that cube. This results in a coordination number of 8. - **D. Zinc (HCP):** Zinc crystallizes in a Hexagonal Close-Packed (HCP) structure. For any close-packed arrangement, including HCP and Face-Centered Cubic (FCC), each atom is in direct contact with a total of 12 other atoms. These neighbors consist of 6 atoms within its own plane, 3 atoms in the plane directly above, and 3 atoms in the plane directly below. Thus, the coordination number for an HCP structure is 12.

Step 2: Having determined the coordination number for each structure, we can now arrange them in an ascending sequence based on these values. The coordination numbers we found are: - Diamond (A): 4 - Sodium Chloride (B): 6 - Cesium Chloride (C): 8 - Zinc (HCP) (D): 12. Placing these in order from the smallest to the largest value gives the following sequence: A (with a coordination number of 4) is less than B (6), which is less than C (8), which is less than D (12). The correct ascending order is therefore A, B, C, D.

 Quick Tip

Memorize the coordination numbers for common crystal structures: Diamond (4), NaCl (6), BCC/CsCl (8), and HCP/FCC (12). This is a frequent topic in solid-state physics questions.

5. Subtract $(29.A)_{16}$ from $(4F.B)_{16}$

- (1) $(26.1)_{16}$
- (2) $(26.A)_{16}$
- (3) $(4F.A)_{16}$

(4) $(16.1)_{16}$

Correct Answer: (1) $(26.1)_{16}$

Solution: Step 1: The first step is to correctly align the two hexadecimal numbers for subtraction, ensuring that the hexadecimal points are lined up vertically. The problem is to compute the difference:

$$\begin{array}{r} 4F. \\ B \\ - 29. \\ \hline A \end{array}$$

Step 2: We begin the subtraction process from the rightmost column, which represents the fractional part of the numbers. We need to compute B minus A. In the decimal system, the hexadecimal digit B is equivalent to 11, and A is equivalent to 10. The subtraction is straightforward: $11 - 10 = 1$. The result for this column is 1_{16} .

$$\begin{array}{r} 4F. \\ B \\ - 29. \\ \hline A \\ 1 \end{array}$$

Step 3: Next, we move to the left of the hexadecimal point and subtract the integer parts, proceeding column by column from right to left. First, we address the units column (16^0): F minus 9. In decimal, F is 15. So, the calculation is $15 - 9 = 6$. The result for this column is 6_{16} .

$$\begin{array}{r} 4F. \\ B \\ - 29. \\ \hline A \\ 6. \\ 1 \end{array}$$

Then, we proceed to the next column to the left (the 16^1 s column): 4 minus 2. This is a simple subtraction: $4 - 2 = 2$. The result for this column is 2_{16} .

$$\begin{array}{r} 4F. \\ B \\ - 29. \\ \hline A \\ 26. \\ 1 \end{array}$$

Step 4: The final step is to assemble the results from each column to form the complete answer. By combining the integer and fractional parts, we obtain the final result of the subtraction. The answer is $(26.1)_{16}$.

💡 Quick Tip

When performing hexadecimal arithmetic, remember the decimal equivalents: A=10, B=11, C=12, D=13, E=14, F=15. For subtraction, if you need to borrow, you borrow 16 from the column to the left. In this case, no borrowing was needed.

6. If the load resistance decreases in a zener regulator, the series current:

- (1) decreases.
- (2) stays the same.
- (3) increases.
- (4) equals the source voltage divided by the series resistance

Correct Answer: (2) stays the same.

Solution: Step 1: Let's first visualize and understand the operation of an ideal Zener diode voltage regulator circuit. This circuit typically features a series resistor, denoted as R_S , which is connected between a higher, unregulated input voltage source (V_{in}) and the rest of the circuit. A Zener diode is placed in parallel with the load resistor (R_L), and it is reverse-biased. The primary function of the Zener diode, when operating in its breakdown region, is to establish and maintain a nearly constant voltage, the Zener voltage (V_Z), across itself and therefore across the parallel-connected load.

Step 2: Now, we will determine the expression for the series current, which is the total current flowing from the source, denoted as I_S . This current must pass through the series resistor R_S . The voltage drop across this resistor, V_{RS} , is determined by the difference between the constant input voltage and the constant Zener voltage, so $V_{RS} = V_{in} - V_Z$. By applying Ohm's law to this series resistor, we can express the series current as:

$$I_S = \frac{V_{RS}}{R_S} = \frac{V_{in} - V_Z}{R_S}$$

Step 3: Let's consider what happens when the load resistance, R_L , is changed. The problem states that R_L decreases. In a properly functioning Zener regulator, the Zener diode ensures that the voltage across the load, V_Z , remains constant. The input voltage V_{in} is also assumed to be constant, and the series resistor R_S has a fixed value. Observing the equation for the series current, I_S , we see that it depends only on V_{in} , V_Z , and R_S . Since all three of these quantities are constant, the series current I_S must also remain constant, regardless of any changes in the load resistance.

Step 4: To understand the internal dynamics of the circuit, let's look at how the currents are distributed. According to Kirchhoff's current law, the constant series current I_S splits at the junction, dividing between the Zener diode (current I_Z) and the load resistor (current I_L). This gives the relationship $I_S = I_Z + I_L$. When the load resistance R_L decreases, the current drawn by the load, calculated as $I_L = V_Z/R_L$, will increase. Because the total series current I_S must stay constant, the circuit automatically compensates by reducing the current flowing through the Zener diode, I_Z . The regulator continues to function correctly as long as the Zener current I_Z does not drop to zero.

💡 Quick Tip

In a working Zener regulator, think of the series resistor and the constant Zener voltage as setting a constant total current supply (I_S). The Zener diode then "absorbs" whatever current is not drawn by the load to keep the voltage constant.

7. In a controlled current source with OP-Amp the circuit acts as:

- (1) voltage amplifier.
- (2) current-to-voltage converter.
- (3) voltage-to-current converter.
- (4) current amplifier.

Correct Answer: (3) voltage-to-current converter.

Solution: Step 1: First, we need to grasp the fundamental purpose of a circuit referred to as a "controlled current source." This is a type of electronic circuit specifically designed to produce and maintain a constant, predictable current that flows through a load. Crucially, the magnitude of this output current is not arbitrary; it is precisely determined or "controlled" by an independent input signal.

Step 2: Now, let's examine how an Operational Amplifier (Op-Amp) is employed to achieve this functionality. In a standard Op-Amp based design for a controlled current source, the control signal is an input voltage, V_{in} , which is applied to one of the Op-Amp's input terminals. The Op-Amp, leveraging its defining characteristics of extremely high open-loop gain and a carefully configured negative feedback network, continually adjusts its own output voltage. This adjustment is done in such a way as to force the current passing through the load, I_{out} , to be directly and linearly proportional to the input voltage, V_{in} . This creates the relationship $I_{out} = k \cdot V_{in}$, where k is a constant of proportionality determined by the circuit's resistors.

Step 3: Finally, we classify the circuit's function by considering the nature of its input and output signals. The input to the circuit is a voltage (V_{in}), and the primary output is a current (I_{out}). A circuit that accepts a voltage as its input and produces a proportional current as its output is, by definition, a voltage-to-current converter. This type of circuit is also known by the more formal name of a transconductance amplifier, as it converts a voltage signal into a current signal.

💡 Quick Tip

Remember the four basic types of amplifiers based on their input and output signals: - Voltage In, Voltage Out → Voltage Amplifier - Current In, Voltage Out → Transresistance Amplifier (Current-to-Voltage Converter) - Voltage In, Current Out → Transconductance Amplifier (Voltage-to-Current Converter) - Current In, Current Out → Current Amplifier

8. Match the LIST-I with LIST-II

LIST-I (Logic Gates)		LIST-II (Expressions)	
A.	EX-OR	I.	$AB + \bar{A}\bar{B}$
B.	NAND	II.	$A + B$
C.	OR	III.	AB
D.	EX-NOR	IV.	$\bar{A}\bar{B} + AB$

Choose the correct answer from the options given below:

- (1) A - I, B - II, C - III, D - IV
- (2) A - I, B - III, C - II, D - IV
- (3) A - I, B - II, C - IV, D - III
- (4) A - III, B - IV, C - I, D - II

Correct Answer: (2) A - I, B - III, C - II, D - IV

Solution: Step 1: The task is to correctly pair each logic gate from LIST-I with its defining Boolean algebraic expression from LIST-II. We will proceed by examining each gate individually. - **A. EX-OR:** The Exclusive-OR (EX-OR) gate is defined by its output being true (logic 1) if and only if its inputs are different from each other. The Boolean expression that represents this condition is $A \oplus B = AB + \bar{A}\bar{B}$. This expression perfectly matches expression **I**. Therefore, the correct pairing is **A → I**. - **C. OR:** The OR gate produces a true output if one or more of its inputs are true. Its standard Boolean expression is the logical sum of its inputs, written as $A + B$. This corresponds exactly to expression **II**. Thus, the correct pairing is **C → II**. - **D. EX-NOR:** The Exclusive-NOR (EX-NOR) gate, also known as the equivalence gate, outputs true if and only if its inputs are the same (both true or both false). Its Boolean expression is the complement of the EX-OR gate, $\overline{A \oplus B}$, which simplifies to $AB + \bar{A}\bar{B}$. This expression is identical to expression **IV**. Hence, the correct pairing is **D → IV**. - **B. NAND:** The NAND gate provides an output that is the negation of an AND gate's output. Its correct Boolean expression is $\overline{AB}$. Looking at the options, expression **III** is given as AB , which is the expression for a simple AND gate, not a NAND gate.

Step 2: Now we will use our confirmed pairings to determine the correct option from the choices provided. We have definitively established the following correct matches: A→I, C→II, and D→IV. Let's inspect the options: - Option (1) suggests C→III, which is incorrect. - Option (2) suggests A→I, C→II, and D→IV. These three matches are correct based on our analysis. This option pairs B (NAND) with III (AND). This indicates a probable typographical error in the question itself, where either the gate B should have been listed as AND, or expression III should have been $\overline{AB}$. Nevertheless, because the other three pairs in this option are perfectly correct, this option stands out as the most likely intended answer. - Option (3) suggests C→IV, which is incorrect. - Option (4) suggests A→III, which is incorrect.

Based on this logical deduction, option (2) is the only plausible choice, despite the apparent error regarding the NAND gate.

💡 Quick Tip

When facing matching questions with a potential error, first identify all the pairs you are certain about. Then, use the process of elimination to find the option that correctly matches all the unambiguous pairs. The remaining pair in that option is likely the intended, albeit flawed, answer.

9. Arrange the following numbers in ascending order:

A. $(10110.011)_2$

B. $(32)_{10}$

C. $(5F.8)_{16}$

D. F_{16}

Choose the Correct answer from the options given below:

(1) A, B, C, D

(2) D, A, B, C

(3) B, A, D, C

(4) C, B, D, A

Correct Answer: (2) D, A, B, C

Solution: Step 1: In order to compare these numbers, which are given in different number systems (binary, decimal, and hexadecimal), we must first convert them all into a single, common base. The decimal (base-10) system is the most convenient choice for this purpose. -

A. Convert $(10110.011)_2$ to decimal: We expand the binary number by multiplying each digit by the corresponding power of 2.

$$\begin{aligned} & (1 \cdot 2^4) + (0 \cdot 2^3) + (1 \cdot 2^2) + (1 \cdot 2^1) + (0 \cdot 2^0) + (0 \cdot 2^{-1}) + (1 \cdot 2^{-2}) + (1 \cdot 2^{-3}) \\ &= (1 \cdot 16) + (0 \cdot 8) + (1 \cdot 4) + (1 \cdot 2) + (0 \cdot 1) + (0 \cdot 0.5) + (1 \cdot 0.25) + (1 \cdot 0.125) \\ &= 16 + 0 + 4 + 2 + 0 + 0 + 0.25 + 0.125 = 22.375_{10} \end{aligned}$$

- **B. Convert $(32)_{10}$ to decimal:** This number is already in the decimal system, so no conversion is necessary. Its value is 32_{10} . - **C. Convert $(5F.8)_{16}$ to decimal:** We expand the hexadecimal number, remembering that F is equivalent to 15 in decimal.

$$\begin{aligned} & (5 \cdot 16^1) + (F \cdot 16^0) + (8 \cdot 16^{-1}) \\ &= (5 \cdot 16) + (15 \cdot 1) + (8/16) = 80 + 15 + 0.5 = 95.5_{10} \end{aligned}$$

- **D. Convert F_{16} to decimal:** The hexadecimal digit F directly corresponds to the decimal number 15.

$$F_{16} = 15_{10}$$

Step 2: Now that all the numbers have been converted to their decimal equivalents, we can directly compare their magnitudes. The decimal values are as follows: - A = 22.375 - B = 32 - C = 95.5 - D = 15

Step 3: The final step is to arrange these decimal values in ascending order, which means from the smallest value to the largest. Comparing the numbers, we see that 15 is the smallest, followed by 22.375, then 32, and the largest is 95.5. Therefore, the correct ascending order of the original labels is D, A, B, C.

💡 Quick Tip

To compare numbers in different bases (binary, decimal, hexadecimal), the most reliable method is to convert all of them to a single base, usually decimal. Remember the positional values are powers of the base (2, 10, or 16).

10. Match the LIST-I with LIST-II

LIST-I (Configuration of Bipolar Transistors)		LIST-II (Characteristics)	
A.	Common Base	I.	Current Gain but no Voltage Gain
B.	Common Emitter	II.	Voltage Gain but no Current Gain
C.	Common Collector	III.	Both Current and Voltage Gain

Choose the correct answer from the options given below:

- (1) A - I, B - II, C - III
- (2) A - II, B - III, C - I
- (3) A - I, B - III, C - II
- (4) A - III, B - II, C - I

Correct Answer: (2) A - II, B - III, C - I

Solution: Step 1: We need to systematically analyze the fundamental properties of each of the three Bipolar Junction Transistor (BJT) amplifier configurations to match them with their correct characteristics. - **A. Common Base (CB):** In this configuration, the input signal is applied to the emitter and the output is taken from the collector, with the base being common to both. It is characterized by a very low input impedance and a very high output impedance. Its current gain, denoted by α , is inherently slightly less than unity (typically 0.95 to 0.99). However, because of the large difference between its high output impedance and low input impedance, it can provide a substantial voltage gain. Therefore, its key characteristic is providing "Voltage Gain but no Current Gain" (since the current gain is not greater than 1). This description corresponds to characteristic **II**. - **B. Common Emitter (CE):** This is the most prevalent amplifier configuration. The input is at the base, the output is at the collector, and the emitter is the common terminal. It exhibits moderate input and output impedances. The CE configuration is unique in that it provides significant amplification for both the current (with a gain of β) and the voltage. This makes it a versatile, general-purpose amplifier. Its characteristic is accurately described as "Both Current and Voltage Gain". This matches characteristic **III**. - **C. Common Collector (CC):** This configuration is also widely known as an emitter follower. The input is at the base, the output is taken from the emitter, and the collector is common. It is defined by its very high input impedance and very low output impedance, making it an excellent buffer. Its voltage gain is always slightly less than, but approximately equal to, 1 (unity). In contrast, it offers a high current gain, which is equal to $\beta + 1$. Its primary feature is thus providing "Current Gain but no Voltage Gain" (as the voltage gain is approximately one). This matches characteristic **I**. **Step 2:** Based on our detailed analysis, we can now establish the correct pairings and select the corresponding option. - Common Base (A) $\rightarrow$ Voltage Gain but no Current Gain (II) - Common Emitter (B) $\rightarrow$ Both Current and Voltage Gain (III) - Common Collector (C) $\rightarrow$ Current Gain but no Voltage Gain (I) This sequence of pairings, A - II, B - III, C - I, is precisely what is listed in option (2).

💡 Quick Tip

A simple way to remember BJT configurations: - **Common Emitter (CE)**: The all-rounder. Good voltage and current gain. Inverts the signal. - **Common Collector (CC)**: The "buffer". High current gain, unity voltage gain. Used for impedance matching. - **Common Base (CB)**: The "current follower". Unity current gain, high voltage gain. Used for high-frequency applications.

11. The first maxima for Bragg's diffraction pattern by a crystal is observed at 30° when X-rays wavelength of 0.32 nm are used. The distance between the atomic planes is:

- (1) 0.32 nm
- (2) 0.48 nm
- (3) 0.84 Å
- (4) 0.48 Å

Correct Answer: (1) 0.32 nm

Solution: Step 1: The first step in solving this problem is to invoke the fundamental principle governing the diffraction of X-rays by a crystal lattice, which is known as Bragg's Law. This law provides the condition for constructive interference of the X-rays scattered by parallel atomic planes. The mathematical expression for Bragg's Law is:

$$n\lambda = 2d \sin \theta$$

In this equation, n is an integer representing the order of the diffraction maximum, λ is the wavelength of the incident X-rays, d is the perpendicular distance between the parallel atomic planes (the interplanar spacing), and θ is the glancing angle of incidence of the X-rays relative to these planes.

Step 2: Next, we must carefully extract the specific values for the variables in Bragg's Law from the text of the problem. - The problem states it is the "first maxima," which implies that we are dealing with the first-order diffraction. Therefore, we set $n = 1$. - The wavelength of the X-rays used is explicitly given as $\lambda = 0.32$ nm. - The angle at which this diffraction maximum is observed is given as $\theta = 30^\circ$.

Step 3: With the formula and all the necessary values identified, we can now substitute these values into the Bragg's Law equation to solve for the unknown interplanar distance, d .

$$(1)(0.32 \text{ nm}) = 2 \cdot d \cdot \sin(30^\circ)$$

We recall the standard trigonometric value for the sine of 30 degrees, which is $\sin(30^\circ) = 0.5$. Substituting this value into the equation gives:

$$0.32 \text{ nm} = 2 \cdot d \cdot (0.5)$$

Multiplying the terms on the right side of the equation simplifies it to:

$$0.32 \text{ nm} = d$$

This result shows that the distance between the atomic planes in the crystal is exactly 0.32 nm.

 Quick Tip

Bragg's Law is a fundamental equation in solid-state physics. Remember that n must be an integer (1, 2, 3, ...) representing the order of the reflection. For "first maxima" or "first-order diffraction," always use $n = 1$.

12. The stopping potential for a fast moving photo-electron is independent of:

- (1) the frequency of incident photon.
- (2) the intensity of incident photon.
- (3) the wavelength of the incident photon.
- (4) type of metals

Correct Answer: (2) the intensity of incident photon.

Solution: Step 1: To determine the dependencies of the stopping potential, we must start with the foundational equation of the photoelectric effect. Einstein's photoelectric equation provides a relationship between the maximum kinetic energy ($K_{\max}$) of an emitted electron (photoelectron), the frequency of the incoming light (f), and a material property called the work function (ϕ). The equation is:

$$K_{\max} = hf - \phi$$

Here, h is Planck's constant. This equation states that the maximum kinetic energy of a photoelectron is the energy of the incident photon (hf) minus the minimum energy required to liberate the electron from the metal surface (ϕ).

Step 2: We now need to connect this maximum kinetic energy to the concept of stopping potential. The stopping potential, denoted V_s , is defined as the minimum reverse voltage that must be applied to completely halt the emission of even the most energetic photoelectrons. The work done by this potential on an electron (with charge e) must equal the electron's initial maximum kinetic energy. This relationship is given by:

$$eV_s = K_{\max}$$

By substituting the first equation into the second, we can derive a direct expression for the stopping potential:

$$V_s = \frac{K_{\max}}{e} = \frac{hf - \phi}{e} = \left(\frac{h}{e}\right) f - \frac{\phi}{e}$$

Step 3: By examining this final equation for V_s , we can systematically analyze its dependencies. - It is directly proportional to the frequency (f) of the incident photon. - Since frequency and wavelength (λ) are related by $f = c/\lambda$, the stopping potential also depends on the wavelength of the incident photon. - The stopping potential depends on the work function (ϕ), which is a characteristic property of the material being illuminated. Therefore, it depends on the type of metal. The intensity of the incident light corresponds to the number of photons arriving per second. A higher intensity means more photons, which will result in more photoelectrons being emitted (a higher photoelectric current). However, the intensity does not change the energy of each individual photon (hf). Since the stopping potential is

determined by the energy of the *most energetic* photoelectron, which depends solely on the photon's energy, it is completely independent of the light's intensity.

💡 Quick Tip

For the photoelectric effect, remember this key distinction: - **Frequency/Wavelength** determines the **Energy** of photoelectrons ($K_{\max}$, V_s). - **Intensity** determines the **Number** of photoelectrons (photocurrent).

13. Ravi and Swati are twins and they are being separated at a rate of $0.80c$. Ravi and Swati each send out a radio signal once a year while Ravi is away. How many signals does Ravi receive for a trip of 15 years?

- (1) 3 signals
- (2) 5 signals
- (3) 9 signals
- (4) No signal

Correct Answer: (2) 5 signals

Solution: Step 1: The core physical principle governing this scenario is the relativistic Doppler effect. Because the source of the radio signals (Swati) and the observer (Ravi) are moving apart from each other at a significant fraction of the speed of light, the time interval between the reception of consecutive signals will be different from the time interval between their transmission. This effect must be accounted for.

Step 2: We need to use the specific formula from special relativity that describes the Doppler effect for the time interval (or period) of a signal when the source and observer are receding from each other. The time interval measured by the receiver (T_{received}) is related to the time interval at the source (T_{sent}) by the following equation:

$$T_{\text{received}} = T_{\text{sent}} \sqrt{\frac{1 + v/c}{1 - v/c}}$$

where v is the relative speed of separation and c is the speed of light.

Step 3: Now, we will insert the numerical values given in the problem into this formula. - The time interval between the signals being sent by Swati is given as one per year, so $T_{\text{sent}} = 1$ year. - The relative velocity of separation is given as $v = 0.80c$, which means $v/c = 0.80$. Substituting these values, we get:

$$T_{\text{received}} = (1 \text{ year}) \sqrt{\frac{1 + 0.80}{1 - 0.80}} = (1 \text{ year}) \sqrt{\frac{1.8}{0.2}} = (1 \text{ year}) \sqrt{9} = 3 \text{ years}$$

This calculation tells us that although Swati sends a signal every year, Ravi, who is moving away, only receives one of these signals every three years from his perspective.

Step 4: The final step is to determine the total number of signals Ravi receives over the entire duration of his trip. The trip's duration in Ravi's frame of reference is 15 years. Since he receives one signal every 3 years, we can calculate the total number of signals by dividing

the total time by the time interval between receptions:

$$\text{Number of signals} = \frac{\text{Total trip time}}{T_{\text{received}}} = \frac{15 \text{ years}}{3 \text{ years/signal}} = 5 \text{ signals}$$

💡 Quick Tip

This is a classic twin paradox-style problem. The key is recognizing that the rate of receiving signals is altered by the relative motion. For objects moving apart, the time between received signals is longer than the time between sent signals.

14. The Schrodinger wave equation is:

- (1) non-linear differential equation.
- (2) linear differential equation.
- (3) second order equation in time.
- (4) first order equation in space.

Correct Answer: (2) linear differential equation.

Solution: Step 1: To analyze the properties of the Schrodinger wave equation, let's first write down its general time-dependent form, which describes how a quantum state evolves over time.

$$i\hbar \frac{\partial \Psi(\mathbf{r}, t)}{\partial t} = \left[-\frac{\hbar^2}{2m} \nabla^2 + V(\mathbf{r}, t) \right] \Psi(\mathbf{r}, t)$$

In this equation, $\Psi(\mathbf{r}, t)$ is the wave function, which contains all the information about the quantum system.

Step 2: Now, we will examine the mathematical characteristics of this equation based on its structure. - **Linearity:** An equation is defined as linear if the dependent variable (in this case, the wave function Ψ) and all of its derivatives appear only to the first power. We can inspect the Schrodinger equation and see that there are no terms involving Ψ^2 , $(\frac{\partial \Psi}{\partial t})^2$, or products like $\Psi \frac{\partial \Psi}{\partial x}$. Every term is proportional to either Ψ or one of its derivatives. This property is crucial because it leads to the principle of superposition, which states that if Ψ_1 and Ψ_2 are two valid solutions, then any linear combination of them, such as $c_1 \Psi_1 + c_2 \Psi_2$, is also a valid solution. - **Order in time:** The order of a differential equation with respect to a variable is the highest order of derivative with respect to that variable. The Schrodinger equation contains the term $\frac{\partial \Psi}{\partial t}$, which is a first derivative with respect to time. There are no higher time derivatives. Therefore, the equation is first-order in time. - **Order in space:** The spatial derivatives are contained within the Laplacian operator, ∇^2 , which is shorthand for $\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2}$. Since this term involves second derivatives with respect to the spatial coordinates, the equation is second-order in space.

Step 3: With this analysis complete, we can now evaluate the given options. 1. non-linear differential equation: This is incorrect. As established, the equation is linear. 2. linear differential equation: This is correct. The wave function and its derivatives appear only to the first power. 3. second order equation in time: This is incorrect. The equation is first-order in time. 4. first order equation in space: This is incorrect. The equation is second-order in space.

 Quick Tip

The linearity of the Schrodinger equation is one of its most important features. It is the mathematical basis for the principle of superposition in quantum mechanics, which allows for phenomena like interference of wave functions.

15. In Compton scattering, Compton shift equals Compton wavelength if angle of scattering is:

- (1) 0
- (2) $\pi/4$
- (3) $\pi/2$
- (4) π

Correct Answer: (3) $\pi/2$

Solution: Step 1: The first step is to recall the fundamental formula that describes the phenomenon of Compton scattering. The Compton shift, symbolized as $\Delta\lambda$, represents the increase in the wavelength of a photon after it has scattered off a charged particle, typically an electron at rest. This shift is a function of the scattering angle θ and is given by the equation:

$$\Delta\lambda = \frac{h}{m_e c} (1 - \cos \theta)$$

Here, h is Planck's constant, m_e is the rest mass of the electron, and c is the speed of light.

Step 2: Next, we need to understand the definition of the Compton wavelength. The Compton wavelength, denoted by λ_c , is a quantum mechanical property of a particle. For an electron, it is defined by the constant group of terms that appears in the Compton shift formula:

$$\lambda_c = \frac{h}{m_e c}$$

It represents the wavelength shift that would occur for a 90-degree scattering event.

Step 3: The problem asks us to find the scattering angle θ for the specific case where the Compton shift is equal to the Compton wavelength. We can express this condition mathematically as $\Delta\lambda = \lambda_c$. We now substitute the full expressions for these two quantities into this equality:

$$\frac{h}{m_e c} (1 - \cos \theta) = \frac{h}{m_e c}$$

Since the term $\frac{h}{m_e c}$ is a non-zero constant present on both sides of the equation, we can divide both sides by it, which simplifies the equation significantly:

$$1 - \cos \theta = 1$$

Subtracting 1 from both sides gives:

$$-\cos \theta = 0$$

This implies that:

$$\cos \theta = 0$$

We now need to find the angle θ (within the physically possible range of 0 to π radians) for which the cosine is zero. This occurs precisely when $\theta = \pi/2$ radians (or 90 degrees).

 Quick Tip

Remember the physical meaning of the limits for Compton scattering: - $\theta = 0$: No scattering, $\Delta\lambda = 0$. - $\theta = \pi/2$ (90 degrees): Shift equals the Compton wavelength, $\Delta\lambda = \lambda_c$. - $\theta = \pi$ (180 degrees, backscattering): Maximum shift, $\Delta\lambda = 2\lambda_c$.

16. Wavelength of X-rays having the largest penetrating power is:

- (1) 1.2 Å
- (2) 6 Å
- (3) 9 Å
- (4) 12 Å

Correct Answer: (1) 1.2 Å

Solution: Step 1: First, we must establish the connection between the penetrating power of a photon and its energy. The ability of electromagnetic radiation, such as X-rays, to pass through material is known as its penetrating power. This property is directly correlated with the energy of the individual photons; photons with higher energy are able to penetrate matter more effectively.

Step 2: Next, we need to recall the fundamental relationship in quantum physics that links a photon's energy (E) to its wavelength (λ). This relationship is described by the Planck-Einstein relation, which states that energy is inversely proportional to wavelength:

$$E = hf = \frac{hc}{\lambda}$$

Here, h is Planck's constant, and c is the speed of light. This equation makes it clear that as wavelength decreases, the photon's energy increases.

Step 3: Now, we can combine the insights from the first two steps. To achieve the largest penetrating power, the X-ray photons must possess the highest possible energy. Based on the energy-wavelength formula, the highest energy is associated with the shortest, or smallest, wavelength.

Step 4: The final step is to examine the provided options and identify the shortest wavelength. The choices given are 1.2 Å, 6 Å, 9 Å, and 12 Å. By simple comparison, the smallest numerical value is 1.2 Å. Therefore, X-rays with this wavelength will have the highest energy and, consequently, the greatest penetrating power.

 Quick Tip

This is a fundamental concept in electromagnetic radiation: Short Wavelength $\leftrightarrow$ High Frequency $\leftrightarrow$ High Energy $\leftrightarrow$ High Penetrating Power. This applies to the entire EM spectrum, from radio waves to gamma rays.

17. Match the LIST-I with LIST-II

LIST-I (Type of decay in Radioactivity)		LIST-II (Reason for stability)	
A.	Alpha decay	I.	Nucleus has excess energy in an excited state
B.	Beta negative decay	II.	Nucleus has too many protons relative to neutrons
C.	Gamma decay	III.	Nucleus is mostly heavier than Pb (Z=82)
D.	Positron Emission	IV.	Nucleus has too many neutrons relative to protons

Choose the correct answer from the options given below:

- (1) A - I, B - II, C - III, D - IV
- (2) A - I, B - III, C - II, D - IV
- (3) A - I, B - II, C - IV, D - III
- (4) A - III, B - IV, C - I, D - II

Correct Answer: (4) A - III, B - IV, C - I, D - II

Solution: Step 1: We will systematically examine each mode of radioactive decay from List-I and determine the underlying nuclear instability from List-II that causes it. - **A.**

Alpha decay: This process involves the emission of an alpha particle, which is a helium nucleus (${}^4_2\text{He}$). This decay reduces the parent nucleus's mass number by 4 and its atomic number by 2. It is a mechanism for very large, heavy nuclei to reduce their overall size and move towards a more stable configuration. This decay mode is predominantly observed in nuclei that are significantly heavier than lead (Z=82). This description matches reason **III**. -

B. Beta negative decay: In this decay, a neutron within the nucleus transforms into a proton, while an electron (e^-) and an electron antineutrino are emitted. The transformation is $n \rightarrow p + e^- + \bar{\nu}_e$. The result is that the atomic number increases by one, and the neutron number decreases by one. This process occurs in nuclei that are "neutron-rich," meaning they have an excess of neutrons relative to protons for their given mass. This corresponds to reason **IV**. -

C. Gamma decay: This is not a transmutation but an energy-releasing process. It occurs when a nucleus is in a metastable, excited energy state. To return to its ground state, it releases the surplus energy by emitting a high-energy photon, known as a gamma ray. The numbers of protons and neutrons remain unchanged. This decay is a consequence of the nucleus having excess energy. This matches reason **I**. -

D. Positron Emission (β^+ decay): In this type of decay, a proton within the nucleus converts into a neutron, and a positron (e^+ , the antiparticle of the electron) and an electron neutrino are emitted. The transformation is $p \rightarrow n + e^+ + \nu_e$. This causes the atomic number to decrease by one and the neutron number to increase by one. This process happens in "proton-rich" nuclei, which have too many protons compared to neutrons for stability. This description matches reason **II**.

Step 2: Now, we assemble the correct pairings to form the complete sequence. - A (Alpha decay) $\rightarrow$ III (Nucleus is too heavy) - B (Beta negative decay) $\rightarrow$ IV (Too many neutrons) - C (Gamma decay) $\rightarrow$ I (Excess energy) - D (Positron Emission) $\rightarrow$ II (Too many protons) This sequence, A - III, B - IV, C - I, D - II, directly corresponds to the combination presented in option (4).

💡 Quick Tip

Think of nuclear decay as a nucleus's way of adjusting its proton-to-neutron ratio to reach the "valley of stability". - Too heavy? $\rightarrow$ Alpha decay. - Too many neutrons? $\rightarrow$ Beta-minus decay. - Too many protons? $\rightarrow$ Positron emission or electron capture. - Too much energy? $\rightarrow$ Gamma decay.

18. The de-Broglie wavelength of an electron moving with a velocity of 10^7 m/s is:

- (1) 7.3×10^{-11} m
- (2) 1.3×10^{-11} m
- (3) 7.3×10^{-7} m
- (4) 3.1×10^{-7} m

Correct Answer: (1) 7.3×10^{-11} m

Solution: Step 1: The first step is to recall the de-Broglie hypothesis, which postulates that all matter exhibits wave-like properties. The wavelength associated with a particle, known as the de-Broglie wavelength (λ), is given by the formula:

$$\lambda = \frac{h}{p} = \frac{h}{mv}$$

where h represents Planck's constant, p is the momentum of the particle, m is its mass, and v is its velocity.

Step 2: Next, we need to gather all the necessary physical constants and the values provided in the problem statement. - Planck's constant is a fundamental constant of nature, $h \approx 6.626 \times 10^{-34}$ J · s. - The mass of an electron is also a fundamental constant, $m_e \approx 9.11 \times 10^{-31}$ kg. - The velocity of the electron is given as $v = 10^7$ m/s.

Step 3: Now we can substitute these values into the de-Broglie wavelength formula and perform the calculation.

$$\lambda = \frac{6.626 \times 10^{-34} \text{ J} \cdot \text{s}}{(9.11 \times 10^{-31} \text{ kg}) \times (10^7 \text{ m/s})}$$

To simplify the calculation, let's handle the numerical part and the powers of ten separately.

$$\lambda = \left(\frac{6.626}{9.11} \right) \times \left(\frac{10^{-34}}{10^{-31} \times 10^7} \right) \text{ m}$$

$$\lambda \approx 0.727 \times 10^{-34 - (-31) - 7} \text{ m}$$

$$\lambda \approx 0.727 \times 10^{-34 + 31 - 7} \text{ m}$$

$$\lambda \approx 0.727 \times 10^{-10} \text{ m}$$

To express this in standard scientific notation, we adjust the decimal point:

$$\lambda \approx 7.27 \times 10^{-11} \text{ m}$$

This calculated value is closest to the option 7.3×10^{-11} m.

Quick Tip

For calculations involving fundamental constants, it's often useful to know approximate ratios. For instance, $h/m_e \approx 7.27 \times 10^{-4}$. This can sometimes speed up calculations. Also, always double-check the powers of ten.

19. Consider the following statements about light:
- A. photoelectric effect exhibits wave nature of light.
 - B. Compton effect exhibits wave nature of light.
 - C. photoelectric effect exhibits particle nature of light.
 - D. Compton effect exhibits particle nature of light.

Choose the CORRECT answer from the options given below:

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

Correct Answer: (4) C and D only

Solution: Step 1: Let us first analyze the photoelectric effect and its implications for the nature of light. The photoelectric effect is the phenomenon where electrons are ejected from a material's surface when it is illuminated by light. Critical experimental observations—such as the fact that electron emission only occurs if the light's frequency is above a certain threshold, and that this emission is nearly instantaneous—could not be reconciled with the classical wave theory of light. Albert Einstein provided the explanation by proposing that light energy is quantized into discrete packets, or particles, called photons. The energy of a photon is proportional to its frequency. This model perfectly explained the experimental data, thereby providing compelling evidence for the particle nature of light. Consequently, statement A is incorrect, and statement C is correct.

Step 2: Now, let's examine the Compton effect. This effect involves the scattering of high-frequency photons (like X-rays or gamma rays) by charged particles, usually electrons. It is observed that the scattered photons have a longer wavelength (lower energy) than the incident photons, with the change in wavelength depending on the scattering angle. This phenomenon is explained by treating the interaction as an elastic collision between two particles: a photon and an electron. In this model, both kinetic energy and momentum are conserved. The "billiard-ball" nature of this interaction is a powerful demonstration of light behaving as a particle with momentum. Thus, the Compton effect supports the particle nature of light. This means statement B is incorrect, and statement D is correct.

Step 3: Based on the analysis of both phenomena, we can conclude which of the given statements are factually correct. We have determined that the photoelectric effect (statement C) and the Compton effect (statement D) are two key experiments that demonstrate the particle-like characteristics of light. Therefore, the correct option must include only statements C and D.

 Quick Tip

Remember the key experiments for wave-particle duality: - **Wave Nature:** Interference, Diffraction, Polarization. - **Particle Nature:** Photoelectric Effect, Compton Scattering, Blackbody Radiation.

20. Match the LIST-I with LIST-II

LIST-I (Energy of a particle box of length L)		LIST-II (Degeneracy of the states)	
A.	$14h^2/(8mL^2)$	I.	1
B.	$11h^2/(8mL^2)$	II.	3
C.	$3h^2/(8mL^2)$	III.	6

Choose the correct answer from the options given below:

- (1) A - I, B - II, C - III
- (2) A - I, B - III, C - II
- (3) A - III, B - II, C - I
- (4) A - III, B - I, C - II

Correct Answer: (3) A - III, B - II, C - I

Solution: Step 1: To begin, we must recall the formula for the quantized energy levels of a particle of mass m confined within a three-dimensional cubic potential well (a box) with side length L . The energy, E , is determined by a set of three quantum numbers, (n_x, n_y, n_z) , which must be positive integers (1, 2, 3, ...). The formula is:

$$E = \frac{h^2}{8mL^2}(n_x^2 + n_y^2 + n_z^2)$$

The degeneracy of an energy level is defined as the number of distinct quantum states (n_x, n_y, n_z) that correspond to the same energy value.

Step 2: We will now determine the degeneracy for each energy value given in List-I. - **A.**

Energy $14h^2/(8mL^2)$: By comparing this to the general formula, we need to find the number of unique sets of positive integers (n_x, n_y, n_z) that satisfy the condition $n_x^2 + n_y^2 + n_z^2 = 14$.

Let's test combinations of squares of small integers (1, 4, 9, 16,...). We find that

$1^2 + 2^2 + 3^2 = 1 + 4 + 9 = 14$. The quantum numbers are (1, 2, 3). Since all three numbers are different, any permutation of them results in a distinct quantum state. The number of permutations of three distinct items is $3! = 3 \times 2 \times 1 = 6$. These states are (1,2,3), (1,3,2), (2,1,3), (2,3,1), (3,1,2), and (3,2,1). Thus, the degeneracy is 6. This means **A** $\rightarrow$ **III**.

- **B. Energy $11h^2/(8mL^2)$:** Here, we require $n_x^2 + n_y^2 + n_z^2 = 11$. Testing combinations, we find $1^2 + 1^2 + 3^2 = 1 + 1 + 9 = 11$. The set of quantum numbers is (1, 1, 3). To find the number of distinct states, we need to find the number of unique permutations of these numbers. The formula for permutations with repetitions is $\frac{n!}{n_1!n_2!\dots}$. Here, we have 3 numbers with a repetition of '1' (twice), so the number of permutations is $\frac{3!}{2!} = 3$. The distinct states are (1,1,3), (1,3,1), and (3,1,1). The degeneracy is 3. This means **B** $\rightarrow$ **II**.

- **C. Energy $3h^2/(8mL^2)$:** For this energy, we need $n_x^2 + n_y^2 + n_z^2 = 3$. The only possible way to achieve this sum using squares of positive integers is $1^2 + 1^2 + 1^2 = 1 + 1 + 1 = 3$. The only quantum state is (1,1,1), which represents the ground state. Since there is only one combination of quantum numbers for this energy, the state is non-degenerate. The degeneracy is 1. This means **C** $\rightarrow$ **I**.

Step 3: Finally, we assemble the correct pairings to find the right option. The correct matches we have determined are: A $\rightarrow$ III, B $\rightarrow$ II, and C $\rightarrow$ I. This sequence corresponds exactly to option (3).

 Quick Tip

To find the degeneracy for a 3D particle in a box, you are looking for the number of ways you can sum the squares of three positive integers to get a specific number. Systematically test combinations of small integers (1, 2, 3, 4, ...) to find the sets that work.

21. A particle of mass m is in an infinite square potential of length L . The wave function is superimposed state of first two energy eigen states, given by

$\Psi(x) = \sqrt{\frac{1}{3}}\Psi_{n=1}(x) + \sqrt{\frac{2}{3}}\Psi_{n=2}(x)$. Identify the correct statements:

- A. $\langle p \rangle = 0$
- B. $\Delta p = \sqrt{3}h/2L$
- C. $\langle E \rangle = 3h^2/8mL^2$
- D. $\Delta x = 0$

Choose the correct answer from the options given below:

- (1) A, B and D only
- (2) A, B and C only
- (3) A, B, C and D
- (4) B, C and D only

Correct Answer: (2) A, B and C only

Solution: Step 1: First, we must analyze the given quantum state. The wave function is a superposition of the first two energy eigenstates, $\Psi = c_1\Psi_1 + c_2\Psi_2$, where the coefficients are $c_1 = \sqrt{1/3}$ and $c_2 = \sqrt{2/3}$. We can verify that the state is properly normalized, as the sum of the probabilities is $|c_1|^2 + |c_2|^2 = (\sqrt{1/3})^2 + (\sqrt{2/3})^2 = 1/3 + 2/3 = 1$. The energy eigenvalues for an infinite square well are given by $E_n = n^2h^2/(8mL^2)$.

Step 2: Now, let's evaluate each statement's validity. - **A. Expectation value of momentum $\langle p \rangle$:** For any stationary state (an energy eigenstate) in a symmetric potential like the infinite square well, the probability density is symmetric, leading to an expectation value of momentum of zero. Since the given state is a superposition of such states, its expectation value of momentum is also zero. Therefore, statement A is correct. - **C.**

Expectation value of energy $\langle E \rangle$: The expectation value of energy in a superposition state is the weighted average of the energy eigenvalues, where the weights are the probabilities of being in each state.

$$\begin{aligned}\langle E \rangle &= |c_1|^2 E_1 + |c_2|^2 E_2 = \left(\frac{1}{3}\right) \left(\frac{1^2 h^2}{8mL^2}\right) + \left(\frac{2}{3}\right) \left(\frac{2^2 h^2}{8mL^2}\right) \\ \langle E \rangle &= \frac{h^2}{8mL^2} \left(\frac{1}{3} + \frac{2 \cdot 4}{3}\right) = \frac{h^2}{8mL^2} \left(\frac{1+8}{3}\right) = \frac{h^2}{8mL^2} \left(\frac{9}{3}\right) = \frac{3h^2}{8mL^2}\end{aligned}$$

Therefore, statement C is correct. - **B. Uncertainty in momentum Δp :** The uncertainty is defined by $(\Delta p)^2 = \langle p^2 \rangle - \langle p \rangle^2$. Since we have already established that $\langle p \rangle = 0$, this simplifies to $(\Delta p)^2 = \langle p^2 \rangle$. For an infinite square well, the energy operator is $\hat{H} = \hat{p}^2/(2m)$, which allows us to find the expectation value of p^2 from the expectation value of energy:

$$\langle p^2 \rangle = 2m \langle E \rangle.$$

$$\langle p^2 \rangle = 2m \left(\frac{3h^2}{8mL^2} \right) = \frac{6mh^2}{8mL^2} = \frac{3h^2}{4L^2}$$

The uncertainty in momentum is the square root of this value:

$$\Delta p = \sqrt{\langle p^2 \rangle} = \sqrt{\frac{3h^2}{4L^2}} = \frac{\sqrt{3}h}{2L}$$

Therefore, statement B is correct. - **D. Uncertainty in position Δx :** The uncertainty in position, Δx , represents the standard deviation of the particle's position. For a particle confined to a box of length L, its position is not precisely known. An uncertainty of $\Delta x = 0$ would imply that the particle's position is known with absolute certainty, which would, according to the Heisenberg Uncertainty Principle ($\Delta x \Delta p \geq \hbar/2$), require an infinite uncertainty in momentum. This is not the case. Therefore, statement D is incorrect.

Step 3: Finally, we identify the collection of all correct statements. From our analysis, statements A, B, and C have been shown to be correct, while D is incorrect. This combination corresponds to option (2).

💡 Quick Tip

For a superposition state $\Psi = \sum c_n \Psi_n$, the expectation value of an operator $\hat{A}$ with eigenvalues a_n is $\langle A \rangle = \sum |c_n|^2 a_n$, provided the Ψ_n are eigenstates of $\hat{A}$. This works for energy, but for operators like momentum, you need to be more careful. The relation $\langle p^2 \rangle = 2m \langle E \rangle$ is a useful shortcut for the infinite square well.

22. If any two rows (or columns) of a determinant are identical then the value of the determinant is:

- (1) 1
- (2) 0
- (3) ∞
- (4) unchanged

Correct Answer: (2) 0

Solution: This question refers to a fundamental property of determinants in linear algebra. Let's consider a matrix A. One of the key properties of determinants states that if we swap any two rows (or columns) of A to create a new matrix A', the determinant of the new matrix is the negative of the original determinant (i.e., $\det(A') = -\det(A)$).

Now, let's apply this property to the specific case where two rows of matrix A are identical. If we swap these two identical rows, the matrix remains completely unchanged. So, in this case, the new matrix A' is exactly the same as the original matrix A.

This leads to the mathematical statement:

$$\det(A) = \det(A')$$

But from the row-swapping property, we also know:

$$\det(A') = -\det(A)$$

Combining these two equations, we get:

$$\det(A) = -\det(A)$$

The only number that is equal to its own negative is zero. Therefore, it must be that the value of the determinant is 0.

 Quick Tip

Memorizing the basic properties of determinants is essential for linear algebra. The key properties include: $\det(A^T) = \det(A)$, $\det(AB) = \det(A)\det(B)$, and the effect of row/column operations (swapping, scaling, adding a multiple of another row/column).

23. The eigen values of matrix A are 1, -2, 3. The eigen values of $3I - 2A + A^2$ are:

- A. 2
- B. 6
- C. 8
- D. 11

Choose the correct answer from the options given below:

- (1) A, B and D only
- (2) A, B and C only
- (3) A, B, C and D
- (4) B, C and D only

Correct Answer: (1) A, B and D only

Solution: Step 1: The key to solving this problem lies in a powerful property of eigenvalues. This property states that if a matrix A has an eigenvalue λ , then any polynomial function of that matrix, let's call it $P(A)$, will have an eigenvalue equal to the same polynomial function evaluated at λ , i.e., $P(\lambda)$.

Step 2: First, we must identify the polynomial function and list the known eigenvalues. The new matrix is given by the expression $3I - 2A + A^2$. This corresponds to a matrix polynomial $P(A) = 3I - 2A + A^2$. The equivalent scalar polynomial, which we will apply to the eigenvalues, is $P(\lambda) = 3 - 2\lambda + \lambda^2$. The given eigenvalues of the original matrix A are $\lambda_1 = 1$, $\lambda_2 = -2$, and $\lambda_3 = 3$.

Step 3: We will now compute the new eigenvalues by substituting each of A 's eigenvalues into our scalar polynomial $P(\lambda)$. - For the first eigenvalue, $\lambda_1 = 1$: The new eigenvalue is $P(1) = 3 - 2(1) + (1)^2 = 3 - 2 + 1 = 2$. - For the second eigenvalue, $\lambda_2 = -2$: The new eigenvalue is $P(-2) = 3 - 2(-2) + (-2)^2 = 3 + 4 + 4 = 11$. - For the third eigenvalue, $\lambda_3 = 3$: The new eigenvalue is $P(3) = 3 - 2(3) + (3)^2 = 3 - 6 + 9 = 6$.

Step 4: Finally, we compare our calculated eigenvalues with the values provided in statements A, B, C, and D. Our calculations show that the eigenvalues of the new matrix are 2, 11, and 6. These values correspond precisely to statements A, D, and B, respectively. The value 8, given in statement C, is not one of the calculated eigenvalues. Therefore, the correct statements are A, B, and D.

 Quick Tip

This property is very powerful. It means you don't need to find the matrix A itself. You can find the eigenvalues of any function of A (like A^2 , A^{-1} , or e^A) directly from the eigenvalues of A .

24. For an even function, the Fourier coefficients are:

- A. $a_0 \neq 0$
- B. $a_n \neq 0$
- C. $a_n = 0$
- D. $b_n = 0$

Choose the correct answer from the options given below:

- (1) A, C and D only
- (2) A, B and D only
- (3) C and D only
- (4) B and D only

Correct Answer: (2) A, B and D only

Solution: Step 1: Let's start by defining the key terms. An even function is a function $f(x)$ that satisfies the symmetry property $f(-x) = f(x)$. The Fourier series expansion of a periodic function is given by $f(x) = \frac{a_0}{2} + \sum_{n=1}^{\infty} (a_n \cos(nx) + b_n \sin(nx))$. The coefficients are calculated using the following integrals over a symmetric interval, typically $[-\pi, \pi]$: - $a_0 = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) dx$ - $a_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \cos(nx) dx$ - $b_n = \frac{1}{\pi} \int_{-\pi}^{\pi} f(x) \sin(nx) dx$

Step 2: Now, we analyze how these coefficients behave when $f(x)$ is an even function, using properties of integrals of even and odd functions. - **Analysis of a_0 and a_n :** The function $\cos(nx)$ is itself an even function. The product of two even functions, $f(x) \times \cos(nx)$, results in another even function. A key property of definite integrals is that the integral of a non-zero even function over a symmetric interval (like $[-\pi, \pi]$) is generally non-zero. (It is equal to twice the integral over half the interval). Therefore, for a typical even function, we expect the coefficients a_0 (which involves the integral of just $f(x)$) and a_n to be non-zero. This means statements A and B are correct, and statement C is incorrect. - **Analysis of b_n :** The function $\sin(nx)$ is an odd function. The product of an even function, $f(x)$, and an odd function, $\sin(nx)$, results in an odd function. Another key property of definite integrals is that the integral of any odd function over a symmetric interval is always exactly zero. Therefore, for any even function, the coefficient b_n must be zero. This means statement D is correct.

Step 3: We can now conclude which statements accurately describe the Fourier coefficients for an even function. Based on our analysis, the true statements are that a_0 is generally non-zero, a_n is generally non-zero, and b_n is always zero. This set of correct statements corresponds to A, B, and D.

💡 Quick Tip

Remember the symmetry properties for Fourier series: - ****Even function:**** Contains only cosine terms ($b_n = 0$). - ****Odd function:**** Contains only sine terms ($a_0 = 0$ and $a_n = 0$). This saves a lot of calculation time.

25. Match the LIST-I with LIST-II

LIST-I (Expressions)		LIST-II (Values)	
A.	i^{49}	I.	1
B.	i^{38}	II.	$-i$
C.	i^{103}	III.	i
D.	i^{92}	IV.	-1

Choose the correct answer from the options given below:

- (1) A - I, B - II, C - III, D - IV
- (2) A - I, B - III, C - II, D - IV
- (3) A - III, B - IV, C - II, D - I
- (4) A - III, B - IV, C - I, D - II

Correct Answer: (3) A - III, B - IV, C - II, D - I

Solution: Step 1: The first step is to recognize the cyclical nature of the powers of the imaginary unit, i . The pattern repeats every four powers: $i^1 = i$, $i^2 = -1$, $i^3 = -i$, and $i^4 = 1$. This cycle means that to evaluate i^n , we only need to find the remainder when the exponent n is divided by 4. The value of i^n will be the same as i^k , where k is this remainder. If the remainder is 0, it corresponds to i^4 .

Step 2: Now, we will apply this rule to evaluate each expression from List-I. - **A. i^{49} :** We divide the exponent 49 by 4. $49 \div 4 = 12$ with a remainder of 1. Therefore, i^{49} is equivalent to i^1 , which is i . This matches value **III.** - **B. i^{38} :** We divide the exponent 38 by 4. $38 \div 4 = 9$ with a remainder of 2. Therefore, i^{38} is equivalent to i^2 , which is -1 . This matches value **IV.** - **C. i^{103} :** We divide the exponent 103 by 4. $103 \div 4 = 25$ with a remainder of 3. Therefore, i^{103} is equivalent to i^3 , which is $-i$. This matches value **II.** - **D. i^{92} :** We divide the exponent 92 by 4. $92 \div 4 = 23$ with a remainder of 0. A remainder of 0 corresponds to the end of a cycle, which is i^4 . Therefore, i^{92} is equivalent to i^4 , which is 1. This matches value **I.**

Step 3: With all expressions evaluated, we can now formulate the correct matching sequence. The correct pairings are: A→III, B→IV, C→II, and D→I. This sequence corresponds to option (3).

💡 Quick Tip

To quickly find the remainder when dividing a number by 4, you only need to look at its last two digits. For example, for i^{103} , you just need the remainder of 03 divided by 4, which is 3.

26. The real and imaginary parts of $\log(x + iy)$ are:

- (1) Real part = $\log(x^2 + y^2)$ and Imaginary part = $\tan^{-1}\left(\frac{y}{x}\right)$
 (2) Real part = $\log(x^2 + y^2)$ and Imaginary part = $\tan^{-1}\left(\frac{x}{y}\right)$
 (3) Real part = $\log \sqrt{x^2 + y^2}$ and Imaginary part = $\tan^{-1}\left(\frac{x}{y}\right)$
 (4) Real part = $\log \sqrt{x^2 + y^2}$ and Imaginary part = $\tan^{-1}\left(\frac{y}{x}\right)$

Correct Answer: (4) Real part = $\log \sqrt{x^2 + y^2}$ and Imaginary part = $\tan^{-1}\left(\frac{y}{x}\right)$

Solution: Step 1: The key to finding the logarithm of a complex number is to first convert it from its Cartesian form ($z = x + iy$) to its polar form. The polar form expresses the complex number in terms of its magnitude (modulus) and angle (argument). The polar form is $z = re^{i\theta}$, where the modulus r is the distance from the origin, calculated as $r = |z| = \sqrt{x^2 + y^2}$, and the argument θ is the angle with the positive real axis, calculated as $\theta = \arg(z) = \tan^{-1}(y/x)$.

Step 2: Now that the complex number is in polar form, we can apply the natural logarithm function.

$$\log(z) = \log(re^{i\theta})$$

Using the standard property of logarithms that $\log(ab) = \log(a) + \log(b)$, we can separate the modulus and the exponential term:

$$\log(z) = \log(r) + \log(e^{i\theta})$$

Next, using the property that the logarithm is the inverse of the exponential function, $\log(e^w) = w$, we simplify the second term:

$$\log(z) = \log(r) + i\theta$$

Step 3: The expression is now in the form of a complex number (Real Part + i · Imaginary Part). We can identify these parts by substituting back the expressions for r and θ . - The real part is the term without i : $\text{Re}(\log(z)) = \log(r) = \log(\sqrt{x^2 + y^2})$. - The imaginary part is the coefficient of i : $\text{Im}(\log(z)) = \theta = \tan^{-1}(y/x)$.

Comparing this result with the given options, we find that it perfectly matches the expressions provided in option (4). It is also worth noting that due to logarithm properties, the real part can also be written as $\frac{1}{2} \log(x^2 + y^2)$.

💡 Quick Tip

To find the logarithm of a complex number, the first step is always to convert it from Cartesian form ($x + iy$) to polar form ($re^{i\theta}$). The logarithm then separates neatly into its real and imaginary parts.

27. If $x = r \cos \theta, y = r \sin \theta$ then Match the LIST-I with LIST-II

LIST-I		LIST-II	
A.	$\frac{\partial r}{\partial x}$	I.	$\frac{1}{r}$
B.	$\frac{\partial r}{\partial y}$	II.	$\frac{y}{r}$
C.	$\frac{\partial(x,y)}{\partial(r,\theta)}$	III.	$\frac{x}{r}$
D.	$\frac{\partial(r,\theta)}{\partial(x,y)}$	IV.	r

(Note: Typo in question, should be $y = r \sin \theta$)

- (1) A - III, B - II, C - I, D - IV
- (2) A - III, B - II, C - IV, D - I
- (3) A - II, B - III, C - IV, D - I
- (4) A - II, B - III, C - I, D - IV

Correct Answer: (2) A - III, B - II, C - IV, D - I

Solution: Step 1: First, we must establish the fundamental relationships that connect Cartesian coordinates (x, y) and polar coordinates (r, θ) . The problem gives the transformation from polar to Cartesian as $x = r \cos \theta$ and $y = r \sin \theta$ (correcting the typo). From these, we can derive the inverse relationships: squaring and adding gives $x^2 + y^2 = r^2 \cos^2 \theta + r^2 \sin^2 \theta = r^2$, so $r = \sqrt{x^2 + y^2}$. Dividing gives $y/x = \tan \theta$, so $\theta = \tan^{-1}(y/x)$.

Step 2: We now proceed to calculate the required partial derivatives using implicit differentiation on the relation $r^2 = x^2 + y^2$. - **A.** $\frac{\partial r}{\partial x}$: We differentiate the equation $r^2 = x^2 + y^2$ with respect to x , treating r as a function of x and y .
 $\frac{\partial}{\partial x}(r^2) = \frac{\partial}{\partial x}(x^2 + y^2) \implies 2r \frac{\partial r}{\partial x} = 2x$. Solving for the derivative gives $\frac{\partial r}{\partial x} = \frac{x}{r}$. This corresponds to item **III**. - **B.** $\frac{\partial r}{\partial y}$: Similarly, we differentiate $r^2 = x^2 + y^2$ with respect to y .
 $\frac{\partial}{\partial y}(r^2) = \frac{\partial}{\partial y}(x^2 + y^2) \implies 2r \frac{\partial r}{\partial y} = 2y$. Solving for the derivative gives $\frac{\partial r}{\partial y} = \frac{y}{r}$. This corresponds to item **II**.

Step 3: Next, we compute the Jacobians for the coordinate transformations. - **C.** $\frac{\partial(x, y)}{\partial(r, \theta)}$: This notation represents the Jacobian determinant of the transformation from polar to Cartesian coordinates. It is calculated as follows:

$$J = \begin{vmatrix} \frac{\partial x}{\partial r} & \frac{\partial x}{\partial \theta} \\ \frac{\partial y}{\partial r} & \frac{\partial y}{\partial \theta} \end{vmatrix} = \begin{vmatrix} \cos \theta & -r \sin \theta \\ \sin \theta & r \cos \theta \end{vmatrix}$$
$$= (\cos \theta)(r \cos \theta) - (-r \sin \theta)(\sin \theta) = r \cos^2 \theta + r \sin^2 \theta = r(\cos^2 \theta + \sin^2 \theta) = r$$

This result matches item **IV**. - **D.** $\frac{\partial(r, \theta)}{\partial(x, y)}$: This is the Jacobian of the inverse transformation. A property of Jacobians states that the Jacobian of the inverse transformation is the reciprocal of the original Jacobian.

$$\frac{\partial(r, \theta)}{\partial(x, y)} = \left(\frac{\partial(x, y)}{\partial(r, \theta)} \right)^{-1} = \frac{1}{r}$$

This result matches item **I**.

Step 4: Finally, we assemble the correct sequence of matches. Based on our calculations, the correct pairings are: A→III, B→II, C→IV, and D→I. This sequence corresponds to option (2).

💡 Quick Tip

The Jacobian $\frac{\partial(x, y)}{\partial(r, \theta)} = r$ is a crucial result used for changing variables in double integrals from Cartesian to polar coordinates: $dx dy = r dr d\theta$.

28. For the differential equation $\left(1 + \frac{d^2y}{dx^2}\right)^{3/2} = y \frac{d^2y}{dx^2}$ the order, degree and linearity respectively are:

- (1) 3, 2 and non-linear
- (2) 2, 3 and non-linear
- (3) 2, 3 and linear
- (4) 3, 2 and linear

Correct Answer: (2) 2, 3 and non-linear

Solution: Step 1: First, we must determine the order of the ordinary differential equation (ODE). The order is defined as the order of the highest derivative that appears in the equation. In the given equation, the highest derivative term is $\frac{d^2y}{dx^2}$, which is a second derivative. Therefore, the **order is 2**.

Step 2: Next, we determine the degree of the ODE. The degree is defined as the highest power (exponent) of the highest-order derivative, but only after the equation has been rationalized to remove any radicals or fractional powers of the derivatives. To eliminate the fractional power of $3/2$, we must square both sides of the equation:

$$\left[\left(1 + \frac{d^2y}{dx^2} \right)^{3/2} \right]^2 = \left(y \frac{d^2y}{dx^2} \right)^2$$

This simplifies to:

$$\left(1 + \frac{d^2y}{dx^2} \right)^3 = y^2 \left(\frac{d^2y}{dx^2} \right)^2$$

If we were to expand the left side using the binomial theorem $(a + b)^3 = a^3 + 3a^2b + 3ab^2 + b^3$, we would get:

$$1 + 3 \left(\frac{d^2y}{dx^2} \right) + 3 \left(\frac{d^2y}{dx^2} \right)^2 + \left(\frac{d^2y}{dx^2} \right)^3 = y^2 \left(\frac{d^2y}{dx^2} \right)^2$$

Now that the equation is in polynomial form with respect to its derivatives, we can identify the highest power of the highest-order derivative $\left(\frac{d^2y}{dx^2}\right)$. The highest power is 3, appearing in the term $\left(\frac{d^2y}{dx^2}\right)^3$. Thus, the **degree is 3**.

Step 3: Finally, we assess the linearity of the equation. An ODE is linear if the dependent variable, y , and all of its derivatives appear only to the first power and are not part of any product with each other. In our rationalized equation, we have terms like $\left(\frac{d^2y}{dx^2}\right)^3$ and $y^2 \left(\frac{d^2y}{dx^2}\right)^2$. The presence of derivatives raised to powers higher than one, as well as the product of y^2 with a derivative, violates the conditions for linearity. Therefore, the equation is **non-linear**.

Combining our findings, the equation has an order of 2, a degree of 3, and is non-linear.

Quick Tip

To find the degree, always make sure the equation is a polynomial in its derivatives first. This means eliminating all fractional powers and denominators involving derivatives.

29. For particular Integral, Match the LIST-I with LIST-II

LIST-I		LIST-II	
A.	$\frac{1}{(D-1)}x^2$	I.	xe^x
B.	$\frac{1}{D^2+D+1}\cos x$	II.	$\sin x$
C.	$\frac{1}{(D-1)^2}e^x$	III.	$\frac{x^2e^x}{2}$
D.	$\frac{1}{D^3-3D^2+4D-2}e^x$	IV.	$-(x^2+2x+2)$

(Note: List-I Item A assumed to be $\frac{1}{D-1}x^2$ based on options)

- (1) A - I, B - II, C - III, D - IV
 (2) A - I, B - III, C - II, D - IV
 (3) A - IV, B - II, C - III, D - I
 (4) A - IV, B - II, C - I, D - III

Correct Answer: (3) A - IV, B - II, C - III, D - I

Solution: Step 1: We must evaluate the particular integral (P.I.) for each of the given expressions using the appropriate methods for the operator $D = d/dx$. - **B.** $\frac{1}{D^2+D+1}\cos x$: The standard rule for a function $\cos(ax)$ is to replace every instance of D^2 with $-a^2$. Here, $a = 1$, so we substitute $D^2 \rightarrow -1^2 = -1$.

$$P.I. = \frac{1}{(-1) + D + 1} \cos x = \frac{1}{D} \cos x$$

The operator $1/D$ represents integration. Therefore, $P.I. = \int \cos x dx = \sin x$. This matches **II**.

- **C.** $\frac{1}{(D-1)^2}e^x$: For an exponential function e^{ax} , we normally substitute $D \rightarrow a$. Here, $a = 1$, and substituting $D = 1$ results in a zero denominator, indicating a case of failure. We use the shift theorem: $\frac{1}{f(D)}e^{ax}V(x) = e^{ax}\frac{1}{f(D+a)}V(x)$. Here $V(x) = 1$.

$$P.I. = e^x \frac{1}{((D+1)-1)^2}(1) = e^x \frac{1}{D^2}(1) = e^x \left(\frac{1}{D} \int 1 dx \right) = e^x \left(\int x dx \right) = e^x \frac{x^2}{2}$$

This matches **III**.

- **D.** $\frac{1}{D^3-3D^2+4D-2}e^x$: We again try substituting $D = 1$. The denominator becomes $1^3 - 3(1)^2 + 4(1) - 2 = 1 - 3 + 4 - 2 = 0$. This is another case of failure. When the substitution $D = a$ makes the denominator $f(D)$ zero, but $f'(a)$ is non-zero, the rule is $P.I. = x \frac{1}{f'(a)}e^{ax}$. Let $f(D) = D^3 - 3D^2 + 4D - 2$. The derivative is $f'(D) = 3D^2 - 6D + 4$. Evaluating at $a = 1$: $f'(1) = 3(1)^2 - 6(1) + 4 = 3 - 6 + 4 = 1$.

$$P.I. = x \frac{1}{1}e^x = xe^x$$

This matches **I**.

- **A.** $\frac{1}{(D-1)}x^2$: For a polynomial function, we use a binomial expansion of the operator.

$$\frac{1}{D-1} = \frac{1}{-(1-D)} = -(1-D)^{-1} = -(1+D+D^2+D^3+\dots)$$

We apply this expanded operator to x^2 , noting that derivatives of x^2 beyond the second are zero.

$$P.I. = -(1+D+D^2)(x^2) = -(x^2 + D(x^2) + D^2(x^2)) = -(x^2 + 2x + 2)$$

This matches **IV**.

Step 2: Now we assemble the correct sequence of matches based on our calculations. The correct pairings are: A→IV, B→II, C→III, and D→I. This sequence corresponds to option (3).

 Quick Tip

Master the different methods for finding particular integrals based on the form of the function on the right side: exponential (e^{ax}), trigonometric ($\sin(ax)$, $\cos(ax)$), polynomial (x^n), and their products. Special attention is needed for "cases of failure" where the simple substitution method fails.

30. If $\vec{A} = \vec{\nabla}\phi$ and $\phi = xy + yz + zx$, then the true statements are:

- A. $\vec{\nabla} \cdot \vec{A} = 0$
- B. $\vec{\nabla} \cdot \vec{A} \neq 0$
- C. $\vec{\nabla} \times \vec{A} = \vec{0}$
- D. $\vec{\nabla} \times \vec{A} \neq \vec{0}$

Choose the correct answer from the option given below:

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

Correct Answer: (1) A and C only

Solution: Step 1: This problem can be solved most efficiently by applying two fundamental identities of vector calculus, which bypass the need for extensive calculation. These identities are: 1. **Curl of a Gradient:** The curl of the gradient of any scalar field ϕ is always identically zero: $\vec{\nabla} \times (\vec{\nabla}\phi) = \vec{0}$. A vector field that can be expressed as the gradient of a scalar is known as a conservative or irrotational field. 2. **Divergence of a Gradient:** The divergence of the gradient of a scalar field ϕ is equal to the Laplacian of that field: $\vec{\nabla} \cdot (\vec{\nabla}\phi) = \nabla^2\phi$.

Step 2: Now we apply these identities to the given vector field $\vec{A}$ and scalar potential ϕ .

Analysis of the Curl of $\vec{A}$: The problem states that $\vec{A} = \vec{\nabla}\phi$. Using the first identity, we can immediately conclude that the curl of $\vec{A}$ must be zero, without needing to know the specific form of ϕ .

$$\vec{\nabla} \times \vec{A} = \vec{\nabla} \times (\vec{\nabla}\phi) = \vec{0}$$

This confirms that statement C is true and statement D is false.

Analysis of the Divergence of $\vec{A}$: Using the second identity, the divergence of $\vec{A}$ is equal to the Laplacian of ϕ . We must calculate this.

$$\vec{\nabla} \cdot \vec{A} = \nabla^2\phi = \frac{\partial^2\phi}{\partial x^2} + \frac{\partial^2\phi}{\partial y^2} + \frac{\partial^2\phi}{\partial z^2}$$

First, we find the first partial derivatives of $\phi = xy + yz + zx$:

$$\frac{\partial\phi}{\partial x} = y + z, \quad \frac{\partial\phi}{\partial y} = x + z, \quad \frac{\partial\phi}{\partial z} = y + x$$

Next, we find the second partial derivatives:

$$\frac{\partial^2 \phi}{\partial x^2} = \frac{\partial}{\partial x}(y + z) = 0$$

$$\frac{\partial^2 \phi}{\partial y^2} = \frac{\partial}{\partial y}(x + z) = 0$$

$$\frac{\partial^2 \phi}{\partial z^2} = \frac{\partial}{\partial z}(y + x) = 0$$

Summing these gives the Laplacian: $\vec{\nabla} \cdot \vec{A} = 0 + 0 + 0 = 0$. This confirms that statement A is true and statement B is false. A scalar potential whose Laplacian is zero is called a harmonic function.

Step 3: Finally, we identify the correct set of true statements. Our analysis has shown that both statement A ($\vec{\nabla} \cdot \vec{A} = 0$) and statement C ($\vec{\nabla} \times \vec{A} = \vec{0}$) are true.

💡 Quick Tip

Knowing the identities $\nabla \times (\nabla \phi) = 0$ and $\nabla \cdot (\nabla \times \vec{F}) = 0$ can save you a lot of time. If a vector field is given as a gradient of a scalar ($\vec{A} = \vec{\nabla} \phi$), its curl must be zero. If it's given as the curl of another vector ($\vec{B} = \vec{\nabla} \times \vec{F}$), its divergence must be zero.

31. The projection of vector $\vec{A} = \hat{i} - 2\hat{j} + \hat{k}$ on vector $\vec{B} = 4\hat{i} - 4\hat{j} + 7\hat{k}$ is:

- (1) $\frac{17}{9}$
- (2) $\frac{17}{7}$
- (3) $\frac{19}{7}$
- (4) $\frac{19}{9}$

Correct Answer: (4) $\frac{19}{9}$

Solution: Step 1: We must begin by recalling the correct formula for the scalar projection of one vector onto another. The scalar projection of a vector $\vec{A}$ onto a vector $\vec{B}$ represents the component of $\vec{A}$ that lies in the direction of $\vec{B}$. It is calculated by taking the dot product of the two vectors and dividing by the magnitude of the vector being projected onto. The formula is:

$$\text{proj}_{\vec{B}} \vec{A} = \frac{\vec{A} \cdot \vec{B}}{|\vec{B}|}$$

Step 2: The next step is to compute the dot product, $\vec{A} \cdot \vec{B}$, of the given vectors. This is done by multiplying their corresponding components and summing the results.

$$\vec{A} \cdot \vec{B} = (1)(4) + (-2)(-4) + (1)(7)$$

$$\vec{A} \cdot \vec{B} = 4 + 8 + 7 = 19$$

Step 3: Now, we need to calculate the magnitude (or length) of the vector $\vec{B}$. The magnitude is found by taking the square root of the sum of the squares of its components.

$$|\vec{B}| = \sqrt{4^2 + (-4)^2 + 7^2}$$

$$|\vec{B}| = \sqrt{16 + 16 + 49} = \sqrt{81} = 9$$

Step 4: Finally, we can compute the scalar projection by substituting the values of the dot product and the magnitude into the formula from Step 1.

$$\text{proj}_{\vec{B}} \vec{A} = \frac{\text{Dot Product}}{\text{Magnitude of } \vec{B}} = \frac{19}{9}$$

 Quick Tip

Be careful to distinguish between scalar projection (which is a number, as asked for here) and vector projection (which would be a vector: $\frac{\vec{A} \cdot \vec{B}}{|\vec{B}|^2} \vec{B}$). The formula divides by the magnitude $|\vec{B}|$, not the magnitude squared.

32. If $\vec{A} = \vec{\nabla} \times \vec{F}$, then $\oiint_S \vec{A} \cdot \hat{n} dS$ (for any closed surface S) is:

- (1) 0
- (2) 4S
- (3) 3S
- (4) 4V

Correct Answer: (1) 0

Solution: Step 1: The problem presents a surface integral of a vector field over a closed surface S. The first logical step is to apply the Divergence Theorem, also known as Gauss's Theorem. This theorem provides a powerful connection between the total flux of a vector field through a closed surface and the behavior of the field within the volume V enclosed by that surface. The theorem states:

$$\oiint_S \vec{A} \cdot \hat{n} dS = \iiint_V (\vec{\nabla} \cdot \vec{A}) dV$$

This transforms the surface integral into a volume integral of the divergence of the vector field.

Step 2: Now, we must use the information given in the problem about the vector field $\vec{A}$. We are told that $\vec{A}$ is the curl of another vector field $\vec{F}$, i.e., $\vec{A} = \vec{\nabla} \times \vec{F}$. We substitute this expression for $\vec{A}$ into the volume integral from the Divergence Theorem:

$$\iiint_V \vec{\nabla} \cdot (\vec{\nabla} \times \vec{F}) dV$$

Step 3: At this point, we employ a fundamental identity from vector calculus. This identity states that the divergence of the curl of any sufficiently smooth vector field is always identically zero.

$$\vec{\nabla} \cdot (\vec{\nabla} \times \vec{F}) = 0$$

This is a mathematical truism, reflecting the fact that a field which is a curl has no sources or sinks.

Step 4: With the integrand known to be zero, we can now evaluate the integral. The integral of zero over any volume is simply zero.

$$\iiint_V (0) dV = 0$$

Since the volume integral is zero, it follows from the Divergence Theorem that the original surface integral must also be zero.

 Quick Tip

The identity $\text{div}(\text{curl } \vec{F}) = 0$ is extremely useful. It implies that a vector field which is the curl of another field (like the magnetic field $\vec{B} = \vec{\nabla} \times \vec{A}$) must be solenoidal (divergence-free). This means it has no sources or sinks, and its field lines must form closed loops.

33. If $|\vec{A} + \vec{B}| = |\vec{A} - \vec{B}|$ then the angle between vectors $\vec{A}$ and $\vec{B}$ is:

- (1) 0
- (2) $\pi/4$
- (3) $\pi/2$
- (4) $3\pi/4$

Correct Answer: (3) $\pi/2$

Solution: Step 1: The initial step is to manipulate the given equation into a more useful form. Since the magnitude of a vector is always a non-negative scalar, we are free to square both sides of the equation without changing the underlying relationship.

$$|\vec{A} + \vec{B}|^2 = |\vec{A} - \vec{B}|^2$$

Step 2: We now utilize a key property of the dot product, which relates the square of a vector's magnitude to the dot product of the vector with itself. For any vector $\vec{V}$, the relationship is $|\vec{V}|^2 = \vec{V} \cdot \vec{V}$. Applying this to our equation gives:

$$(\vec{A} + \vec{B}) \cdot (\vec{A} + \vec{B}) = (\vec{A} - \vec{B}) \cdot (\vec{A} - \vec{B})$$

Step 3: The next step is to expand the dot products on both sides of the equation, much like expanding a binomial product in regular algebra.

$$\vec{A} \cdot \vec{A} + \vec{A} \cdot \vec{B} + \vec{B} \cdot \vec{A} + \vec{B} \cdot \vec{B} = \vec{A} \cdot \vec{A} - \vec{A} \cdot \vec{B} - \vec{B} \cdot \vec{A} + \vec{B} \cdot \vec{B}$$

Recalling that $\vec{A} \cdot \vec{A} = |\vec{A}|^2$ and that the dot product is commutative ($\vec{A} \cdot \vec{B} = \vec{B} \cdot \vec{A}$), we can simplify the expression:

$$|\vec{A}|^2 + 2(\vec{A} \cdot \vec{B}) + |\vec{B}|^2 = |\vec{A}|^2 - 2(\vec{A} \cdot \vec{B}) + |\vec{B}|^2$$

Step 4: We can now simplify this equation by canceling the identical terms ($|\vec{A}|^2$ and $|\vec{B}|^2$) that appear on both sides.

$$2(\vec{A} \cdot \vec{B}) = -2(\vec{A} \cdot \vec{B})$$

Moving all terms to one side gives:

$$4(\vec{A} \cdot \vec{B}) = 0$$

Which simplifies to:

$$\vec{A} \cdot \vec{B} = 0$$

Step 5: The final step is to interpret this result. The dot product of two non-zero vectors is zero if, and only if, the vectors are perpendicular (orthogonal) to each other. An angle of perpendicularity is 90 degrees, which in radians is $\pi/2$.

💡 Quick Tip

Geometrically, the vectors $\vec{A} + \vec{B}$ and $\vec{A} - \vec{B}$ represent the diagonals of a parallelogram formed by vectors $\vec{A}$ and $\vec{B}$. The condition that the diagonals have equal length means the parallelogram must be a rectangle, which implies that $\vec{A}$ and $\vec{B}$ are perpendicular.

34. The place at which plane of vibration of Foucault's pendulum does not rotate at all, is:

- (1) Pole
- (2) Equator
- (3) Tropic of Cancer
- (4) Tropic of Capricorn

Correct Answer: (2) Equator

Solution: Step 1: To solve this, we must first recall the physics behind the Foucault pendulum's apparent rotation. The effect is due to the Earth's rotation, and the rate at which the pendulum's swing plane precesses (rotates) depends on its geographical latitude. The standard formula for the angular speed of this precession, ω_P , is:

$$\omega_P = \Omega \sin \phi$$

where Ω is the angular speed of the Earth's rotation (one full circle, or 360° , per sidereal day) and ϕ is the latitude of the pendulum's location.

Step 2: The question asks for the location where the plane of vibration "does not rotate at all." This corresponds to a situation where the rate of precession, ω_P , is equal to zero.

$$\omega_P = \Omega \sin \phi = 0$$

Step 3: Now we must solve this equation for the latitude, ϕ . We know that the Earth is rotating, so its angular speed Ω is a non-zero constant. Therefore, for the product to be zero, the other term must be zero:

$$\sin \phi = 0$$

The sine function is zero when its argument is zero. This condition is met when the latitude $\phi = 0^\circ$.

Step 4: The final step is to identify the geographical location on Earth that corresponds to a latitude of 0° . This is, by definition, the Earth's Equator. At other locations, like the North

Pole ($\phi = +90^\circ$) or South Pole ($\phi = -90^\circ$), the value of $\sin \phi$ is ± 1 , and the rate of rotation is at its maximum.

 Quick Tip

Visualize the Earth's rotation. At the poles, an observer is simply spinning in place, so the pendulum's plane appears to rotate a full circle in 24 hours. At the equator, an observer is carried along without any local twisting motion, so the pendulum's plane remains fixed relative to the ground.

35. A 1500 kg car traveling east with a speed of 25 m/s collides at an intersection with a 2500 kg van traveling north at a speed of 20 m/s. The direction of wreckage after collision, assuming that the vehicles undergo a perfectly inelastic collision is:

- (1) 58.2°
- (2) 47.2°
- (3) 53.1°
- (4) 50.6°

Correct Answer: (3) 53.1°

Solution: Step 1: The first step is to establish a coordinate system and calculate the initial momentum of each vehicle as a vector. Let's define the eastward direction as the positive x-axis and the northward direction as the positive y-axis. Since momentum is a vector quantity ($\vec{p} = m\vec{v}$), we calculate the components of the total initial momentum. - The car travels east, so its momentum is entirely in the x-direction:

$p_x = m_{car}v_{car} = (1500 \text{ kg})(25 \text{ m/s}) = 37500 \text{ kg} \cdot \text{m/s}$. - The van travels north, so its momentum is entirely in the y-direction: $p_y = m_{van}v_{van} = (2500 \text{ kg})(20 \text{ m/s}) = 50000 \text{ kg} \cdot \text{m/s}$.

The total initial momentum of the system is the vector sum of these components:

$$\vec{P}_{initial} = 37500\hat{i} + 50000\hat{j}.$$

Step 2: Next, we apply the law of conservation of linear momentum. The problem states that the collision is "perfectly inelastic," which means the two vehicles stick together and move as a single mass after the collision. In any closed system, the total momentum before the collision must equal the total momentum after the collision.

$$\vec{P}_{final} = \vec{P}_{initial} = 37500\hat{i} + 50000\hat{j}$$

Step 3: The direction of the wreckage after the collision will be the same as the direction of the final total momentum vector. We can find the angle, θ , that this vector makes with our defined x-axis (East) using trigonometry. The components of the final momentum are $P_x = 37500$ and $P_y = 50000$.

$$\tan \theta = \frac{\text{Opposite}}{\text{Adjacent}} = \frac{P_y}{P_x} = \frac{50000}{37500}$$

Step 4: Now, we perform the calculation to find the angle θ .

$$\tan \theta = \frac{500}{375} = \frac{4 \times 125}{3 \times 125} = \frac{4}{3} \approx 1.333$$

To find the angle, we take the arctangent (or inverse tangent) of this ratio:

$$\theta = \arctan\left(\frac{4}{3}\right) \approx 53.13^\circ$$

This angle represents the direction of the wreckage, measured north of the eastward direction. The closest answer is 53.1° .

 Quick Tip

In 2D collision problems, always break the momentum into x and y components. Momentum is conserved independently in each direction. For a perfectly inelastic collision, the final velocity vector points in the same direction as the total initial momentum vector.

36. Moment of inertia of a solid cone about its vertical axis is:

- (1) $MR^2/10$
- (2) $3MR^2/10$
- (3) $5MR^2/10$
- (4) $7MR^2/10$

Correct Answer: (2) $3MR^2/10$

Solution: This question asks for a standard result for the moment of inertia of a common geometrical shape. The moment of inertia, I , is a measure of an object's resistance to angular acceleration. For a solid cone with a total mass M and a circular base of radius R , the moment of inertia about its central axis of symmetry (the vertical axis that passes through the apex and the center of the base) is given by the well-known formula:

$$I = \frac{3}{10}MR^2$$

This formula is typically derived using integral calculus. The derivation involves conceptually slicing the cone into an infinite number of infinitesimally thin circular disks stacked along the vertical axis. The moment of inertia of each small disk (with mass dm and radius r) is $dI = \frac{1}{2}r^2dm$. To find the total moment of inertia of the cone, one must integrate this expression over the entire height of the cone, taking into account how the radius r and the differential mass dm change with height. This integration process yields the final result.

 Quick Tip

It's helpful to memorize the moments of inertia for common shapes: - Thin Rod (center): $ML^2/12$ - Hoop (central axis): MR^2 - Solid Cylinder/Disk (central axis): $MR^2/2$ - Solid Sphere (center): $2MR^2/5$ - Solid Cone (central axis): $3MR^2/10$

37. If the torque remains constant while the angle changes, the work done is equal to:

- (1) ratio of torque and angular displacement
- (2) ratio of the angular displacement and square root of torque
- (3) product of torque and angular displacement
- (4) product of torque and square root of angular displacement

Correct Answer: (3) product of torque and angular displacement

Solution: Step 1: The most direct way to understand work in a rotational context is to draw an analogy with linear (translational) motion. In linear mechanics, the work done (W) by a constant force (F) that moves an object through a displacement (d) is defined as the product of the force and the displacement: $W = F \cdot d$.

Step 2: We can apply this same conceptual framework to rotational motion by replacing the linear quantities with their rotational counterparts. - In rotation, the quantity that causes an angular acceleration (the rotational equivalent of force) is **torque**, denoted by τ . - The rotational equivalent of linear displacement (d) is **angular displacement**, denoted by θ .

Step 3: By directly substituting the rotational analogues into the linear work formula, we can formulate the expression for rotational work. The work done (W) by a **constant torque** (τ) that rotates an object through an angular displacement (θ) is the product of the torque and the angular displacement.

$$W = \tau\theta$$

This expression directly matches the description given in option (3). It is important to note that this simple product form is only valid for a constant torque. If the torque varies with the angle, the work done must be calculated by integrating the torque over the angular displacement: $W = \int \tau d\theta$.

💡 Quick Tip

Many concepts in rotational mechanics have direct analogues in linear mechanics. - Position $x \leftrightarrow$ Angle θ - Velocity $v \leftrightarrow$ Angular Velocity ω - Mass $m \leftrightarrow$ Moment of Inertia I - Force $F \leftrightarrow$ Torque τ - Momentum $p = mv \leftrightarrow$ Angular Momentum $L = I\omega$ - Work $W = Fd \leftrightarrow$ Work $W = \tau\theta$

38. For a force $\mathbf{F}$ to be conservative, the relations to be satisfied are:

- A. $\frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y} = 0$
- B. $\frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z} = 0$
- C. $\frac{\partial F_x}{\partial z} - \frac{\partial F_z}{\partial x} = 0$
- D. $\frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y} = \frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z} = \frac{\partial F_x}{\partial z} - \frac{\partial F_z}{\partial x} \neq 0$

Choose the correct answer from the options given below:

- (1) A and B only
- (2) A, B and C only
- (3) B, C and D only
- (4) A, B, C and D only

Correct Answer: (2) A, B and C only

Solution: Step 1: First, we must recall the fundamental mathematical condition that a force field must satisfy to be classified as conservative. A force field $\vec{F}$ is conservative if and only if its curl is equal to the zero vector. This is a complete and sufficient condition.

$$\vec{\nabla} \times \vec{F} = \vec{0}$$

Step 2: Next, we need to write out the full expression for the curl of a vector field $\vec{F} = F_x\hat{i} + F_y\hat{j} + F_z\hat{k}$ in Cartesian coordinates. This is typically calculated using a symbolic determinant:

$$\vec{\nabla} \times \vec{F} = \begin{vmatrix} \hat{i} & \hat{j} & \hat{k} \\ \frac{\partial}{\partial x} & \frac{\partial}{\partial y} & \frac{\partial}{\partial z} \\ F_x & F_y & F_z \end{vmatrix}$$

Expanding this determinant yields the three components of the curl vector:

$$\vec{\nabla} \times \vec{F} = \left(\frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z} \right) \hat{i} - \left(\frac{\partial F_z}{\partial x} - \frac{\partial F_x}{\partial z} \right) \hat{j} + \left(\frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y} \right) \hat{k}$$

Step 3: For the curl vector to be the zero vector ($\vec{0} = 0\hat{i} + 0\hat{j} + 0\hat{k}$), each of its individual components must be equal to zero. This gives us a set of three required conditions. - The $\hat{i}$ component must be zero: $\frac{\partial F_z}{\partial y} - \frac{\partial F_y}{\partial z} = 0$. This is precisely statement B. - The $\hat{j}$ component must be zero: $-\left(\frac{\partial F_z}{\partial x} - \frac{\partial F_x}{\partial z} \right) = 0$, which simplifies to $\frac{\partial F_x}{\partial z} - \frac{\partial F_z}{\partial x} = 0$. This is precisely statement C. - The $\hat{k}$ component must be zero: $\frac{\partial F_y}{\partial x} - \frac{\partial F_x}{\partial y} = 0$. This is precisely statement A. Thus, for a force to be conservative, all three of the relations given in statements A, B, and C must hold true. Statement D describes a condition where the curl is non-zero, which defines a non-conservative force.

💡 Quick Tip

A force is conservative if it can be written as the gradient of a scalar potential, $\vec{F} = -\vec{\nabla}V$. The condition $\vec{\nabla} \times \vec{F} = 0$ is equivalent to this, due to the identity $\vec{\nabla} \times (\vec{\nabla}V) = 0$.

39. As water flows from a faucet, stream of water becomes narrower as it descends. The guiding principle for this observation is:

- (1) Bernoulli's equation in fluid dynamics
- (2) Pascal's law
- (3) Continuity equation in fluid dynamics
- (4) Archimedes's principle

Correct Answer: (3) Continuity equation in fluid dynamics

Solution: Step 1: First, let us analyze the physical situation described. When a stream of water exits a faucet, it is in freefall. Under the influence of gravity, it accelerates, meaning its downward speed continuously increases as it descends.

Step 2: Now, we must apply the appropriate physical principle to relate this change in speed to the shape of the stream. The guiding principle here is the conservation of mass as it applies to fluid flow, which is encapsulated in the continuity equation. For a steady flow of an

incompressible fluid (a very good approximation for water), the volume flow rate must be constant at all points along the stream. This is expressed mathematically as:

$$A \times v = \text{Constant}$$

or

$$A_1 v_1 = A_2 v_2$$

where A is the cross-sectional area of the fluid stream and v is the speed of the fluid at that cross-section.

Step 3: We can now connect the observation from Step 1 with the principle from Step 2. Let's consider a point 1 just as the water leaves the faucet and a point 2 at some distance below it. Due to gravitational acceleration, we know that the speed at point 2 is greater than the speed at point 1 ($v_2 > v_1$). According to the continuity equation, for the product Av to remain constant, an increase in speed (v) must be accompanied by a decrease in the cross-sectional area (A). A smaller cross-sectional area means that the stream of water must become narrower. This directly explains the observed phenomenon. While Bernoulli's equation is also relevant to the overall flow, it is the continuity equation that specifically and directly explains the change in the stream's width.

💡 Quick Tip

The continuity equation $Av = \text{constant}$ is a very intuitive principle of fluid flow. It simply means that "what goes in must come out." If you squeeze a hose to make the opening smaller (decrease A), the water must speed up (increase v).

40. When in a small pond a person in rowboat, throws an anchor overboard, what happens to the water level?

- (1) Goes down
- (2) Goes up
- (3) First goes up and then goes down
- (4) Remains same

Correct Answer: (1) Goes down

Solution: Step 1: Analyze the initial state (anchor inside the boat). In the beginning, the boat and the anchor form a single floating system. According to Archimedes' principle, any object that floats displaces a volume of fluid whose weight is equal to the total weight of the floating object. Let W_{boat} be the weight of the boat and W_{anchor} be the weight of the anchor. The total weight of the system is $W_{total} = W_{boat} + W_{anchor}$. The weight of the water displaced is therefore $W_{disp,1} = W_{total}$. We can express the volume of displaced water as $V_{disp,1} = \frac{W_{disp,1}}{\rho_{water}g} = \frac{W_{boat} + W_{anchor}}{\rho_{water}g}$.

Step 2: Analyze the final state (anchor at the bottom of the pond). After the anchor is thrown overboard, it sinks to the bottom, and the boat is left floating by itself. We must now consider the water displaced by each object separately. - The boat, now lighter, floats and displaces a volume of water corresponding only to its own weight:

$V_{boat,disp} = \frac{W_{boat}}{\rho_{water}g}$. - The anchor, being fully submerged, displaces a volume of water equal to

its own physical volume. We can express the anchor's volume in terms of its weight and density: $V_{\text{anchor}} = \frac{W_{\text{anchor}}}{\rho_{\text{anchor}}g}$. The total volume of water displaced in this final state is the sum of the volume displaced by the boat and the volume displaced by the anchor:

$$V_{\text{disp},2} = V_{\text{boat,disp}} + V_{\text{anchor}} = \frac{W_{\text{boat}}}{\rho_{\text{water}}g} + \frac{W_{\text{anchor}}}{\rho_{\text{anchor}}g}$$

Step 3: Compare the volume of displaced water in the two states. Let's compare $V_{\text{disp},1}$ with $V_{\text{disp},2}$.

$$V_{\text{disp},1} = \frac{W_{\text{boat}}}{\rho_{\text{water}}g} + \frac{W_{\text{anchor}}}{\rho_{\text{water}}g}$$

The only difference between the total displaced volume in the two states is the term related to the anchor. In the initial state, the anchor's contribution to displacement is $\frac{W_{\text{anchor}}}{\rho_{\text{water}}g}$. In the final state, its contribution is $\frac{W_{\text{anchor}}}{\rho_{\text{anchor}}g}$. An anchor is made of a dense material like iron, so its density, ρ_{anchor} , is significantly greater than the density of water, ρ_{water} . Because $\rho_{\text{anchor}} > \rho_{\text{water}}$, it follows that $\frac{1}{\rho_{\text{water}}} > \frac{1}{\rho_{\text{anchor}}}$. This directly implies that the volume of water displaced by the anchor's weight (initial state) is greater than the volume of water displaced by the anchor's volume (final state): $\frac{W_{\text{anchor}}}{\rho_{\text{water}}g} > \frac{W_{\text{anchor}}}{\rho_{\text{anchor}}g}$. Therefore, the total initial displaced volume is greater than the total final displaced volume: $V_{\text{disp},1} > V_{\text{disp},2}$.

Step 4: Determine the effect on the water level. Since the total volume of water being displaced by the boat-anchor system is less in the final state than in the initial state, the overall water level in the pond must fall.

💡 Quick Tip

The key insight is that when the anchor is in the boat, it contributes its full *weight* to water displacement. When it's at the bottom, it only displaces its own *volume*. Since the anchor is denser than water, the volume of water equivalent to its weight is much larger than its actual volume.

41. A skater is using very low-friction rollerblades. A friend throws a Frisbee straight at her. In which case does the Frisbee impart the greatest impulse to the skater:

- (1) she catches the Frisbee and holds it
- (2) she catches it momentarily but drops it
- (3) she catches it and at once throws it back to her friend
- (4) she can't catch it at all

Correct Answer: (3) she catches it and at once throws it back to her friend

Solution: Step 1: First, let's define the key physical quantity, impulse. Impulse (J) is defined as the change in an object's momentum (Δp). According to the impulse-momentum theorem, the impulse delivered to an object equals the change in its momentum. By Newton's third law of action-reaction, the impulse the Frisbee imparts to the skater is equal in magnitude and opposite in direction to the impulse the skater imparts to the Frisbee. Therefore, to find the greatest impulse on the skater, we should find the scenario where the

Frisbee's momentum changes by the largest amount. Let the initial momentum of the Frisbee be $\vec{p}_i = m\vec{v}$, where the direction is toward the skater.

Step 2: Now, let's analyze the change in the Frisbee's momentum for each scenario. - **Case 1 (Catches and holds):** The skater brings the Frisbee to a stop relative to herself. The final momentum of the Frisbee is now part of the skater's momentum, but its momentum relative to the skater is zero. The change in the Frisbee's momentum is

$\Delta\vec{p} = \vec{p}_f - \vec{p}_i = 0 - m\vec{v} = -m\vec{v}$. The magnitude of the impulse is $|mv|$. - **Case 2 (Catches and drops):** To catch the Frisbee, even momentarily, the skater must first absorb all of its initial momentum. The process of dropping it vertically afterwards does not affect the horizontal impulse. The change in horizontal momentum is still $-m\vec{v}$, and the impulse magnitude is $|mv|$. - **Case 3 (Catches and throws back):** This is a two-part process.

First, the skater absorbs the initial momentum $m\vec{v}$ to catch it (an impulse of magnitude $|mv|$). Then, she applies an additional impulse to throw it back, giving it a new momentum in the opposite direction, $\vec{p}_f = -m\vec{v}'$. The total change in the Frisbee's momentum is

$\Delta\vec{p} = \vec{p}_f - \vec{p}_i = (-m\vec{v}') - (m\vec{v}) = -m(v + v')$. The magnitude of this change is $m(v + v')$.

Since v and v' are both positive speeds, this magnitude is clearly greater than $|mv|$. - **Case 4 (Can't catch):** If there is no interaction, there is no force exerted, and therefore no change in momentum. The impulse is zero.

Step 3: By comparing the magnitudes of the impulse in each case, it is clear that the largest change in momentum occurs in Case 3, where the Frisbee's direction of motion is completely reversed. This reversal imparts the greatest possible impulse to the skater.

💡 Quick Tip

Impulse is maximized when there is a reversal of momentum. Think of a bouncy ball versus a lump of clay hitting a wall; the bouncy ball imparts a greater impulse because its momentum changes from $+p$ to $-p$, a total change of $2p$.

42. The engine of a rocket in outer space, far from any planet is turned on. The rocket ejects burnt fuel at constant rate. In the first second of firing, it ejects 1/100 of its initial mass at relative speed of 2000 m/s. The initial acceleration of the rocket is:

- (1) 5 m/s^2
- (2) -10 m/s^2
- (3) $+20 \text{ m/s}^2$
- (4) -30 m/s^2

Correct Answer: (3) $+20 \text{ m/s}^2$

Solution: Step 1: The first step is to recall the fundamental equation that describes the thrust force generated by a rocket engine. This force is a direct result of the momentum change of the ejected fuel. The formula for thrust, F , is:

$$F = v_{rel} \left| \frac{dm}{dt} \right|$$

Here, v_{rel} is the speed of the ejected fuel relative to the rocket, and $\left| \frac{dm}{dt} \right|$ is the magnitude of the rate at which mass is ejected (the mass flow rate).

Step 2: Next, we connect this thrust force to the rocket's acceleration using Newton's second law of motion, $F = Ma$. In this context, F is the thrust, M is the instantaneous mass of the rocket, and a is its acceleration.

$$Ma = v_{rel} \left| \frac{dm}{dt} \right|$$

We can rearrange this to solve for the acceleration:

$$a = \frac{v_{rel}}{M} \left| \frac{dm}{dt} \right|$$

Step 3: We are asked to find the *initial* acceleration. Therefore, we must use the values of the variables at the very beginning of the firing ($t = 0$). - The mass of the rocket at this moment is its initial mass, so we use $M = M_0$. - The relative speed of the exhaust is given as $v_{rel} = 2000 \text{ m/s}$. - The problem states that the rocket ejects $1/100$ of its initial mass ($M_0/100$) in the first second. Assuming a constant rate, this gives us the mass flow rate:

$$\left| \frac{dm}{dt} \right| = \frac{\text{mass ejected}}{\text{time taken}} = \frac{M_0/100}{1 \text{ s}} = \frac{M_0}{100} \text{ kg/s}.$$

Step 4: Finally, we substitute these initial values into our acceleration equation to find the initial acceleration.

$$a_{initial} = \frac{2000 \text{ m/s}}{M_0} \left(\frac{M_0}{100} \text{ kg/s} \right)$$

The M_0 terms cancel out, leaving:

$$a_{initial} = \frac{2000}{100} \text{ m/s}^2 = 20 \text{ m/s}^2$$

The acceleration is positive, signifying that it is in the direction of the rocket's motion, opposite to the direction of the ejected fuel.

💡 Quick Tip

The rocket thrust equation is a direct application of the conservation of momentum.
The force on the rocket is the reaction force to the force required to eject the fuel mass.

43. Match List-I with List-II on basis of two simple harmonic signals of same frequency and various phase difference interacts with each other:

LIST-I (Lissajous Figure)		LIST-II (Phase Difference)	
A.	Right handed elliptically polarised vibrations	I.	Phase difference = $\pi/4$
B.	Left handed elliptically polarised vibrations	II.	Phase difference = $3\pi/4$
C.	Circularly polarized vibrations	III.	No phase difference
D.	Linearly polarized vibrations	IV.	Phase difference = $\pi/2$

Choose the correct answer from the options given below:

- (1) A - I, B - II, C - III, D - IV
- (2) A - I, B - III, C - II, D - IV
- (3) A - I, B - II, C - IV, D - III
- (4) A - III, B - IV, C - I, D - II

Correct Answer: (3) A - I, B - II, C - IV, D - III

Solution: Step 1: The core of this problem is understanding how the superposition of two perpendicular simple harmonic motions (SHMs) of the same frequency creates different Lissajous figures based on their phase difference, δ . - **D. Linearly polarized vibrations:** This occurs when the resulting motion of the particle is confined to a straight line. This happens when the two SHMs are perfectly in phase ($\delta = 0$) or perfectly out of phase ($\delta = \pi$). In this case, the y-displacement is always directly proportional to the x-displacement. This case corresponds to **III (No phase difference)**. - **C. Circularly polarized vibrations:** A circular path is a special case that occurs under two conditions: the amplitudes of the two SHMs must be equal, and the phase difference must be exactly $\delta = \pi/2$ (90 degrees) or $\delta = 3\pi/2$. This puts the motions in quadrature. This description matches **IV**. - **A. B. Elliptically polarized vibrations:** For all other phase differences between 0 and π , the resulting path is an ellipse. The orientation and direction of tracing of the ellipse depend on the exact phase difference. - A phase difference between 0 and $\pi/2$, such as $\delta = \pi/4$, results in an ellipse that is, by convention, described as right-handed. Therefore, **A matches I**. - A phase difference between $\pi/2$ and π , such as $\delta = 3\pi/4$, results in another ellipse, which can be defined as left-handed based on its tracing direction. Therefore, **B matches II**.
Step 2: With the individual pairings established, we can now assemble the complete matching sequence. - A (Right handed ellipse) $\rightarrow$ I ($\pi/4$) - B (Left handed ellipse) $\rightarrow$ II ($3\pi/4$) - C (Circle) $\rightarrow$ IV ($\pi/2$) - D (Line) $\rightarrow$ III (0) This sequence, A - I, B - II, C - IV, D - III, corresponds to the arrangement in option (3).

 Quick Tip

For Lissajous figures with two SHMs of the same frequency: - Phase diff 0 or $\pi \rightarrow$ Straight Line - Phase diff $\pi/2$ or $3\pi/2$ (and equal amplitudes) $\rightarrow$ Circle - Any other phase diff $\rightarrow$ Ellipse

44. Displacement of a particle at any instant of time t is $y = 5 \sin(100\pi t + \phi)$. The frequency of oscillation of the particle is:

- (1) 100 Hz
- (2) 25 Hz
- (3) 200 Hz
- (4) 50 Hz

Correct Answer: (4) 50 Hz

Solution: Step 1: We begin by identifying the general mathematical form for the displacement in simple harmonic motion (SHM). The displacement, y , as a function of time, t , is typically expressed as $y = A \sin(\omega t + \phi)$. In this equation, A represents the amplitude, ω is the angular frequency (measured in radians per second), and ϕ is the initial phase angle.

Step 2: The next step is to compare the specific equation given in the problem, $y = 5 \sin(100\pi t + \phi)$, with the standard form. By direct comparison, we can extract the value of the angular frequency, ω . The term multiplying the time variable t inside the sine function is the angular frequency. Thus, we have $\omega = 100\pi$ rad/s.

Step 3: We need to find the linear frequency (f), which is measured in Hertz (Hz), not the angular frequency (ω). The relationship between these two quantities is fundamental to oscillatory motion:

$$\omega = 2\pi f$$

We can rearrange this formula to solve for the linear frequency, f :

$$f = \frac{\omega}{2\pi}$$

Step 4: Finally, we substitute the value of ω that we extracted in Step 2 into this relationship to calculate the frequency.

$$f = \frac{100\pi}{2\pi} = 50 \text{ Hz}$$

Therefore, the frequency of the particle's oscillation is 50 Hz.

💡 Quick Tip

Always be careful to distinguish between angular frequency ω (in rad/s) and frequency f (in Hz). The term multiplying t inside the sine/cosine function is always ω .

45. Which of the following conditions will lead to Anomalous dispersion?

- (1) Group velocity $\neq$ Phase Velocity
- (2) Group velocity $\neq$ Phase Velocity
- (3) Group velocity = Phase Velocity
- (4) Doesn't depend on relation of group and phase velocity

Correct Answer: (1) Group velocity $\neq$ Phase Velocity

Solution: Step 1: First, we must define what dispersion means in the context of wave propagation. Dispersion is the phenomenon where the speed at which a wave travels through a medium depends on its frequency or wavelength. To analyze this, we use two different measures of velocity: the phase velocity, $v_p = \frac{\omega}{k}$ (the speed of a single frequency component), and the group velocity, $v_g = \frac{d\omega}{dk}$ (the speed of the overall wave packet or envelope).

Step 2: The relationship between group velocity and phase velocity allows us to classify the type of dispersion occurring in a medium. - ****No Dispersion:**** In a non-dispersive medium (like a vacuum for light), all frequencies travel at the same speed. This leads to the condition that the group velocity is equal to the phase velocity, $v_g = v_p$. - ****Normal Dispersion:**** This is the more common type of dispersion, observed, for example, when white light passes through a glass prism. In this case, higher frequencies (like blue light) travel slower than lower frequencies (like red light), meaning the phase velocity decreases as frequency increases. This corresponds to a situation where the group velocity is less than the phase velocity, $v_g < v_p$. - ****Anomalous Dispersion:**** This is a less common phenomenon that occurs in specific frequency ranges for certain materials, usually near a resonant absorption frequency of the medium. In this regime, the phase velocity actually increases as the frequency increases. This counter-intuitive behavior corresponds to a situation where the group velocity is greater than the phase velocity, $v_g > v_p$.

Based on this classification, the condition that leads to anomalous dispersion is when the group velocity exceeds the phase velocity.

💡 Quick Tip

A simple way to remember is: - Normal: $v_g < v_p$ (The "normal" situation for light through a prism). - Anomalous: $v_g > v_p$ (The "anomalous" or unusual case). - Non-dispersive: $v_g = v_p$ (e.g., light in a vacuum).

46. A tuning fork of unknown frequency sounded with a tuning fork of frequency 256 Hz produces 4 beats per second. If a small quantity of wax is fixed on first fork so that it produces 3 beats per second with tuning fork, what will be the frequency of first fork (in Hz)?

- (1) 260
- (2) 252
- (3) 256
- (4) 280

Correct Answer: (1) 260

Solution: Step 1: The initial information allows us to determine the possible frequencies of the unknown tuning fork. Let the unknown frequency be f_1 and the reference frequency be $f_2 = 256$ Hz. The beat frequency is the absolute difference between the two frequencies, $f_{beat} = |f_1 - f_2|$. We are given that $f_{beat} = 4$ Hz. This gives us two possible values for f_1 : - Possibility A: $f_1 - 256 = 4 \implies f_1 = 260$ Hz. - Possibility B: $256 - f_1 = 4 \implies f_1 = 252$ Hz.

Step 2: Next, we must understand the physical effect of adding wax to a tuning fork. Adding wax increases the mass of the prongs. An increase in mass causes the fork to vibrate more slowly, thereby decreasing its natural frequency of oscillation. Let's call the new, lower frequency of the first fork f'_1 , where we know that $f'_1 < f_1$.

Step 3: We can now use the second piece of information (the new beat frequency of 3 Hz) to determine which of our two initial possibilities is correct. We will test each case. - **Testing**

Case A: Assume the initial frequency was $f_1 = 260$ Hz. When wax is added, the frequency f'_1 will become slightly lower than 260 Hz (e.g., 259 Hz, 258 Hz, etc.). The new beat frequency will be $|f'_1 - 256|$. As f'_1 decreases from 260 and moves closer to 256, the difference between the frequencies will decrease. A new beat frequency of 3 Hz (which would occur at $f'_1 = 259$ Hz) is entirely consistent with this scenario. - **Testing Case B: Assume the initial frequency was $f_1 = 252$ Hz.** When wax is added, the frequency f'_1 will become slightly lower than 252 Hz (e.g., 251 Hz). The new beat frequency will be $|f'_1 - 256| = 256 - f'_1$. As f'_1 decreases from 252, it moves further away from 256. This means the difference between the frequencies, and thus the beat frequency, must increase to a value greater than 4 Hz. This directly contradicts the observation that the beat frequency decreased to 3 Hz.

Step 4: Based on our logical test, only Case A is consistent with all the information provided in the problem. Therefore, the original frequency of the unknown tuning fork must have been 260 Hz.

 Quick Tip

Remember the effects of modifying a tuning fork: - ****Adding mass (waxing):**** Decreases frequency. - ****Removing mass (filing):**** Increases frequency.

47. When light ray refracts on entering from one medium to another medium of different refractive indices, it follows:

- (1) Minimum distance
- (2) Minimum time
- (3) Maximum time
- (4) Maximum distance

Correct Answer: (2) Minimum time

Solution: The fundamental principle that governs the path taken by light rays through different media is known as ****Fermat's Principle of Least Time****. This principle is more general than the laws of reflection and refraction and can be used to derive them. It states that out of all possible paths a light ray might take to get from one point to another, it will always follow the path that takes the minimum amount of time.

When light travels within a single, uniform medium, its speed is constant. In this case, the path of minimum time is also the path of minimum distance—a straight line. However, when a light ray crosses the boundary from one medium to another (refraction), its speed changes (since $v = c/n$). The straight-line path is no longer the quickest path. To minimize its total travel time, the light ray will bend at the interface, spending a different amount of time in each medium than it would have along a straight line. This bent path, governed by Snell's Law, is precisely the path of least time.

 Quick Tip

Fermat's Principle is a powerful concept in optics. Think of it like a lifeguard saving a swimmer: the lifeguard doesn't run in a straight line to the swimmer, but runs a longer distance on the sand (where they are fast) and a shorter distance in the water (where they are slow) to minimize the total time.

48. Two thin convex lenses of focal lengths 2 cm and 6 cm are separated by a distance of 4 cm in air. Arrange the following cardinal points in ascending order on basis of their distance from second lens:

- A. First Principal Point
- B. First Focal Point
- C. Second Focal Point
- D. Second Nodal Point

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

(4) C, B, D, A

Correct Answer: (2) A, C, B, D

Solution: Step 1: First, we establish a coordinate system and define the parameters. Let the first lens (L_1) be at the origin ($x = 0$) and the second lens (L_2) be at $x = 4$ cm. The given values are: focal length of the first lens, $f_1 = 2$ cm; focal length of the second lens, $f_2 = 6$ cm; and the separation distance, $d = 4$ cm.

Step 2: We calculate the equivalent focal length, F , of the lens combination using the formula:

$$\frac{1}{F} = \frac{1}{f_1} + \frac{1}{f_2} - \frac{d}{f_1 f_2} = \frac{1}{2} + \frac{1}{6} - \frac{4}{(2)(6)} = \frac{6 + 2 - 4}{12} = \frac{4}{12} = \frac{1}{3}$$

This gives an equivalent focal length of $F = 3$ cm.

Step 3: We calculate the positions of the two principal points, P_1 and P_2 . - The position of the first principal point (P_1) is measured from the first lens (L_1). Its position is

$\alpha_1 = \frac{dF}{f_2} = \frac{4 \times 3}{6} = 2$ cm. So, the absolute position of P_1 is at $x = 2$. - The position of the

second principal point (P_2) is measured from the second lens (L_2). Its position is

$\alpha_2 = -\frac{dF}{f_1} = -\frac{4 \times 3}{2} = -6$ cm. This means P_2 is 6 cm to the left of L_2 , so its absolute position is at $x = 4 - 6 = -2$.

Step 4: We calculate the positions of the two focal points, F_1 and F_2 , relative to their respective principal points. - The first focal point (F_1) is located at a distance $-F$ from P_1 .

Its absolute position is $x = (\text{pos of } P_1) - F = 2 - 3 = -1$. - The second focal point (F_2) is located at a distance $+F$ from P_2 . Its absolute position is $x = (\text{pos of } P_2) + F = -2 + 3 = 1$.

Step 5: We determine the positions of the nodal points, N_1 and N_2 . Because the optical system is in a uniform medium (air) on both sides, the nodal points coincide exactly with the principal points. Thus, N_1 is at the same position as P_1 ($x = 2$), and N_2 is at the same position as P_2 ($x = -2$).

Step 6: Finally, we calculate the distance of each specified cardinal point from the second lens (L_2 , which is at $x = 4$) and arrange them in ascending order. - A (First Principal Point, P_1 at $x = 2$): Distance = $|2 - 4| = 2$ cm. - C (Second Focal Point, F_2 at $x = 1$): Distance = $|1 - 4| = 3$ cm. - B (First Focal Point, F_1 at $x = -1$): Distance = $|-1 - 4| = 5$ cm. - D (Second Nodal Point, N_2 at $x = -2$): Distance = $|-2 - 4| = 6$ cm.

Arranging these distances from smallest to largest gives the sequence: A (2 cm) ; C (3 cm) ; B (5 cm) ; D (6 cm). The correct order of the labels is therefore A, C, B, D.

Quick Tip

For a two-lens system, always establish a coordinate system first (e.g., first lens at the origin). Calculate the positions of the principal points relative to the physical lenses, then use these principal points as the reference for locating the focal points.

49. Match the LIST-I with LIST-II

LIST-I		LIST-II	
A.	Compton Effect	I.	Diffraction
B.	Colors in thin film	II.	Interference
C.	Double Refraction	III.	Polarization
D.	Bragg's Equation	IV.	Scattering

Choose the correct answer from the options given below:

- (1) A - IV, B - II, C - III, D - I
- (2) A - I, B - III, C - II, D - IV
- (3) A - I, B - II, C - IV, D - III
- (4) A - III, B - IV, C - I, D - II

Correct Answer: (1) A - IV, B - II, C - III, D - I

Solution: Step 1: We must systematically analyze each physical phenomenon in List-I and identify the fundamental principle from List-II that best describes it. - **A. Compton Effect:** This effect describes the interaction where a high-energy photon (like an X-ray) collides with an electron, resulting in a change in the photon's wavelength. This is a classic example of an inelastic collision between particles and is fundamentally a **Scattering** process. This matches **IV**. - **B. Colors in thin film:** The shimmering, vibrant colors observed on soap bubbles, oil slicks, or coatings on lenses are produced by the superposition of light waves. Specifically, it is the constructive and destructive **Interference** between light rays reflecting from the top and bottom surfaces of the thin film that causes certain colors to be enhanced and others to be cancelled. This matches **II**. - **C. Double Refraction:** Also known as birefringence, this is an optical property of certain anisotropic materials (like calcite crystals) where an incoming ray of unpolarized light is split into two rays that travel at different speeds and are polarized perpendicularly to each other. This phenomenon is a direct manifestation of the **Polarization** of light. This matches **III**. - **D. Bragg's Equation:** The equation $n\lambda = 2d \sin \theta$ provides the condition for constructive interference of waves that are scattered by the periodic atomic planes within a crystal. This selective reflection is the principle behind X-ray crystallography and is a form of **Diffraction**. This matches **I**.

Step 2: Now we assemble the correct pairings into a single sequence. - A (Compton Effect) → IV (Scattering) - B (Thin film colors) → II (Interference) - C (Double Refraction) → III (Polarization) - D (Bragg's Equation) → I (Diffraction) This sequence corresponds to the arrangement in option (1).

💡 Quick Tip

Associate these key pairs: - Thin Films ↔ Interference - Crystals/Slits ↔ Diffraction - Calcite/Polaroids ↔ Polarization - Photon-Electron Collision ↔ Scattering

50. The resolving power of a grating:

- A. increases with increase in total number of lines ruled on grating.
- B. increases with increase in total width of grating.
- C. increases with increasing the order of spectrum as in Echelon grating.
- D. increases with decreasing the order of spectrum as in Echelon grating.

Choose the **CORRECT** answer from the options given below:

- (1) A, B and D only
- (2) A, B and C only
- (3) B and D only
- (4) A and D only

Correct Answer: (2) A, B and C only

Solution: Step 1: We must begin by recalling the definition and the standard formula for the resolving power of a diffraction grating. The resolving power, R , of a grating is a measure of its ability to distinguish between two closely spaced spectral lines. It is defined as $R = \frac{\lambda}{\Delta\lambda}$, where $\Delta\lambda$ is the minimum wavelength separation that can be resolved at a wavelength λ . The formula derived for a grating's resolving power is:

$$R = mN$$

where m is the spectral order (a positive integer like 1, 2, 3...) and N is the total number of lines or rulings illuminated on the grating.

Step 2: We will now analyze each statement based on this formula and related principles. -

A. increases with increase in total number of lines ruled on grating. Looking at the formula $R = mN$, it is evident that the resolving power R is directly proportional to the total number of lines N . Therefore, a grating with more lines will have a higher resolving power. This statement is correct. -

B. increases with increase in total width of grating. The total width of the illuminated part of the grating, W , is related to the number of lines N and the spacing between them, d , by the equation $W = Nd$. We can rewrite the resolving power formula using the grating equation ($d \sin \theta = m\lambda$). From this, $m = \frac{d \sin \theta}{\lambda}$. Substituting this into $R = mN$ gives $R = \left(\frac{d \sin \theta}{\lambda}\right) N = \frac{(Nd) \sin \theta}{\lambda} = \frac{W \sin \theta}{\lambda}$. This shows that R is directly proportional to the total width W . This statement is correct. -

C. increases with increasing the order of spectrum as in Echelon grating. Again, from the primary formula $R = mN$, the resolving power R is directly proportional to the spectral order m . Observing the spectrum in a higher order (e.g., second or third order) will yield better resolution. The Echelon grating is a specific type of grating designed to operate at very high orders, precisely to achieve extremely high resolving power. This statement is correct. -

D. increases with decreasing the order of spectrum as in Echelon grating. This statement is the direct opposite of statement C and contradicts the formula $R = mN$. Therefore, it is incorrect.

Step 3: We conclude by identifying the set of correct statements. Based on our analysis, statements A, B, and C all correctly describe how the resolving power of a grating can be increased.

Quick Tip

To get high resolution with a grating, you want to use a grating with many total lines (N) and observe the spectrum in a high order (m). Both factors directly improve the ability to separate close spectral lines.

51. A 20g of cane sugar is dissolved in water to make 50 cc of solution. A 20 cm length of tube filled with this solution causes $+53^{\circ}30'$ optical rotation. What will be the specific rotation?

- (1) 66.9 degree (decimeter) $^{-1}$ (g/cc) $^{-1}$
- (2) 7.09 degree (decimeter) $^{-1}$ (g/cc) $^{-1}$
- (3) 76.9 degree (decimeter) $^{-1}$ (g/cc) $^{-1}$
- (4) 6.69 degree (decimeter) $^{-1}$ (g/cc) $^{-1}$

Correct Answer: (1) 66.9 degree (decimeter) $^{-1}$ (g/cc) $^{-1}$

Solution: Step 1: The first step is to recall the defining formula for specific rotation. Specific rotation, denoted by $[\alpha]$, is an intrinsic property of a chiral substance. It relates the observed optical rotation to the concentration of the solution and the path length of the light through it. The formula is:

$$[\alpha] = \frac{\theta}{l \times c}$$

where θ is the observed optical rotation in degrees, l is the path length of the sample tube in **decimeters (dm)**, and c is the concentration of the solution in **g/cc** (grams per cubic centimeter, which is equivalent to g/mL).

Step 2: Next, we must process the given values and convert them into the standard units required by the formula. - **Observed rotation (θ):** The value is given as $53^{\circ}30'$. We need to convert the minutes to decimal degrees. Since there are 60 minutes in a degree, $30' = 30/60 = 0.5^{\circ}$. So, $\theta = 53.5^{\circ}$. - **Path length (l):** The length is given as 20 cm. The formula requires decimeters. Since 1 dm = 10 cm, we have $l = 20 \text{ cm}/10 = 2 \text{ dm}$. - **Concentration (c):** The solution was made by dissolving a mass of 20 g in a final volume of 50 cc.

$$c = \frac{\text{mass of solute}}{\text{volume of solution}} = \frac{20 \text{ g}}{50 \text{ cc}} = 0.4 \text{ g/cc}$$

Step 3: Now we can substitute these correctly-formatted values into the specific rotation formula.

$$[\alpha] = \frac{53.5}{2 \times 0.4} = \frac{53.5}{0.8}$$

Performing the division gives:

$$[\alpha] = 66.875$$

Step 4: Finally, we compare our calculated result to the given options. The value 66.875 is approximately 66.9. The units are degree·dm $^{-1}$ ·(g/cc) $^{-1}$, matching option (1).

Quick Tip

Pay close attention to units when calculating specific rotation. The path length must be in decimeters and the concentration in g/cc or g/mL. A common mistake is forgetting to convert from centimeters to decimeters.

52. True conditions for sustained interference of light waves are:

- A. Two interfering sources must be coherent.
- B. Two interfering waves must be propagated along the same line.

C. Two interfering waves must have equal amplitude.

D. If the interfering waves are polarized, they must be in the same state of polarization.

Choose the correct answer from the options given below:

(1) A, B and D only

(2) A, B and C only

(3) A, B, C and D

(4) B, C and D only

Correct Answer: (3) A, B, C and D

Solution: Step 1: We need to evaluate each statement to see if it represents a necessary or ideal condition for producing a "sustained" interference pattern, which means a pattern that is stable, observable, and has good contrast. - **A. Two interfering sources must be coherent.** This is the most fundamental requirement. Coherence means that the phase difference between the two sources remains constant over time. If the phase relationship were random, the positions of maxima and minima would fluctuate rapidly, and the interference pattern would be washed out and unobservable. Thus, statement A is true. - **B. Two interfering waves must be propagated along the same line.** For the waves to interfere, they must overlap in the same region of space. While they don't need to be perfectly collinear, for a clear and large-scale interference pattern to form (like the fringes in Young's double-slit experiment), the waves must be propagating in very nearly the same direction. Thus, statement B is considered a true condition. - **C. Two interfering waves must have equal amplitude.** While interference can occur between waves of unequal amplitude, the visibility or contrast of the interference pattern is maximized when the amplitudes are equal. With equal amplitudes, the intensity at the minima (destructive interference) becomes zero, leading to the clearest distinction between bright and dark fringes. For ideal sustained interference, this is a required condition. Thus, statement C is true. - **D. If the interfering waves are polarized, they must be in the same state of polarization.** Interference is the superposition of electric field vectors. Two light waves that are polarized perpendicularly to each other (e.g., one vertically polarized and one horizontally polarized) cannot interfere to produce intensity variations. Their electric field vectors are orthogonal and add up in a way that doesn't produce the characteristic interference pattern. Therefore, the waves must have at least partial components of their polarization in the same direction. Thus, statement D is true.

Step 2: Having evaluated all four statements, we can conclude which set of conditions is correct. All four statements—A, B, C, and D—describe the ideal conditions necessary to observe a clear, stable, high-contrast, and sustained interference pattern with light waves.

 Quick Tip

The three "Cs" of interference are Coherence, Collinearity (or near-collinearity), and Comparable amplitudes. A fourth condition, related to polarization, is also crucial for light waves.

53. A long, straight wire carries a current of 10 A. The magnitude of the magnetic field at a distance of 5 cm from the wire is:

- (1) $4 \times 10^{-5} \text{ T}$
- (2) $8 \times 10^{-5} \text{ T}$
- (3) $12 \times 10^{-5} \text{ T}$
- (4) $16 \times 10^{-5} \text{ T}$

Correct Answer: (1) $4 \times 10^{-5} \text{ T}$

Solution: Step 1: The first step is to recall the formula that describes the magnetic field generated by a long, straight, current-carrying wire. This result is derived from Ampere's Law and states that the magnetic field magnitude, B , at a perpendicular distance r from the wire is:

$$B = \frac{\mu_0 I}{2\pi r}$$

In this formula, I is the current in the wire, and μ_0 is the permeability of free space, a fundamental constant with the value $\mu_0 = 4\pi \times 10^{-7} \text{ T}\cdot\text{m/A}$.

Step 2: Next, we identify the values given in the problem and ensure they are in standard SI units before we use them in the formula. - The current is given as $I = 10 \text{ A}$. - The distance is given as $r = 5 \text{ cm}$. We must convert this to meters: $r = 5 \times 10^{-2} \text{ m}$ or 0.05 m .

Step 3: Now we substitute the known values into the formula and calculate the magnitude of the magnetic field, B .

$$B = \frac{(4\pi \times 10^{-7} \text{ T}\cdot\text{m/A}) \times (10 \text{ A})}{2\pi \times (0.05 \text{ m})}$$

We can simplify the calculation by canceling 2π from the numerator and denominator:

$$B = \frac{(2 \times 10^{-7}) \times (10)}{0.05} \text{ T} = \frac{2 \times 10^{-6}}{5 \times 10^{-2}} \text{ T}$$

$$B = \left(\frac{2}{5}\right) \times 10^{-6-(-2)} \text{ T} = 0.4 \times 10^{-4} \text{ T}$$

To express this in standard scientific notation, we adjust the decimal point:

$$B = 4 \times 10^{-5} \text{ T}$$

 Quick Tip

The expression $\frac{\mu_0}{2\pi}$ can be simplified to $2 \times 10^{-7} \text{ T}\cdot\text{m/A}$, which makes calculations for the magnetic field of a straight wire quicker.

54. A point charge $+Q$ is placed at the origin. The electric potential at point (3a,4a,0) in terms of k , Q and a is:

- (1) $kQ/5a$
- (2) $kQ/3a$
- (3) $kQ/7a$
- (4) $kQ/2a$

Correct Answer: (1) $kQ/5a$

Solution: Step 1: We begin by recalling the fundamental formula for the electric potential created by a single point charge. The electric potential, V , which is a scalar quantity, at a specific location is given by:

$$V = \frac{kQ}{r}$$

where Q is the magnitude of the point charge, k is Coulomb's constant ($k = \frac{1}{4\pi\epsilon_0}$), and r is the straight-line distance from the charge to the point where the potential is being calculated.

Step 2: Next, we must determine the distance, r , between the source of the potential (the charge $+Q$ at the origin) and the point of interest. The charge is located at coordinate $(0,0,0)$, and the point is at $(3a,4a,0)$. We use the three-dimensional distance formula to find r :

$$r = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2 + (z_2 - z_1)^2}$$

$$r = \sqrt{(3a - 0)^2 + (4a - 0)^2 + (0 - 0)^2}$$

$$r = \sqrt{(3a)^2 + (4a)^2} = \sqrt{9a^2 + 16a^2} = \sqrt{25a^2}$$

Assuming a is a positive length, the distance is $r = 5a$.

Step 3: Finally, we substitute this calculated distance into the electric potential formula from Step 1.

$$V = \frac{kQ}{r} = \frac{kQ}{5a}$$

This is the electric potential at the specified point.

 Quick Tip

Electric potential is a scalar quantity, so you only need the magnitude of the distance between the charge and the point of interest. Remember the 3-4-5 right triangle, as it frequently appears in physics problems to simplify distance calculations.

55. A conducting sphere of radius R carries a total charge Q . The electric field at a distance $r > R$ from the center is:

- (1) kQ/r^2
- (2) $kQ/2R^2$
- (3) kQ/R
- (4) 0

Correct Answer: (1) kQ/r^2

Solution: Step 1: The most effective way to determine the electric field outside a spherically symmetric charge distribution is to apply Gauss's Law. A key result from Gauss's Law, often called the Shell Theorem, states that for any point outside a spherically symmetric distribution of charge (such as a uniformly charged shell or a conducting sphere), the electric field is exactly the same as if all of the charge were concentrated into a single point charge located at the center of the sphere.

Step 2: With this principle in mind, we can simply use the well-known formula for the electric field produced by a point charge. The magnitude of the electric field, E , at a distance r from a point charge Q is given by Coulomb's Law:

$$E = \frac{kQ}{r^2}$$

where k is Coulomb's constant, defined as $k = \frac{1}{4\pi\epsilon_0}$.

Step 3: Now, we apply this point-charge formula to our conducting sphere. Since the problem specifies that we are interested in a point at a distance $r > R$ (i.e., outside the sphere), we can directly use the result from the Shell Theorem. The electric field at this external point is:

$$E = \frac{kQ}{r^2}$$

This means the sphere acts electrically as if it were a point charge Q at the origin for all external points.

💡 Quick Tip

This is a key result from Gauss's Law. For a spherical shell or solid conducting sphere:
 - **Outside ($r \geq R$):** $E = kQ/r^2$ (acts like a point charge). - **Inside ($r < R$):** $E = 0$ (for a conductor or hollow shell).

56. A siren on a tall pole radiates sound waves uniformly in all directions. The sound intensity at a distance of 15 m from the siren, is 0.250 W/m^2 . The intensity of sound at distance 75 m from siren is:

- (1) 0.250 W/m^2
- (2) 0.010 W/m^2
- (3) 0.100 W/m^2
- (4) 6.250 W/m^2

Correct Answer: (2) 0.010 W/m^2

Solution: Step 1: We must first recall the physical principle that governs how intensity changes with distance from a source that radiates uniformly. When a source radiates power isotropically (uniformly in all directions), the power spreads out over the surface of an expanding sphere. The surface area of a sphere is $4\pi r^2$. Since the total power is constant, the intensity I (Power per unit area) must decrease in proportion to the surface area it covers. This leads to the inverse square law:

$$I \propto \frac{1}{r^2}$$

From this relationship, we can establish a ratio for two different points: $I_1 r_1^2 = I_2 r_2^2$, as the product $I \cdot r^2$ must be constant.

Step 2: Next, we identify the initial and final conditions from the problem statement. - The initial intensity is $I_1 = 0.250 \text{ W/m}^2$. - The initial distance is $r_1 = 15 \text{ m}$. - The final distance is $r_2 = 75 \text{ m}$. - We need to find the final intensity, I_2 .

Step 3: We can now use the ratio derived from the inverse square law to solve for the unknown intensity, I_2 .

$$I_2 = I_1 \left(\frac{r_1}{r_2} \right)^2$$

Substituting the given values:

$$I_2 = (0.250 \text{ W/m}^2) \times \left(\frac{15 \text{ m}}{75 \text{ m}} \right)^2$$

Simplify the ratio of the distances:

$$I_2 = 0.250 \times \left(\frac{1}{5} \right)^2 = 0.250 \times \frac{1}{25}$$

Performing the final calculation:

$$I_2 = \frac{0.250}{25} = 0.010 \text{ W/m}^2$$

 Quick Tip

The inverse square law is fundamental for any quantity that spreads out uniformly from a point source, including sound intensity, light intensity, and gravitational and electrostatic fields. If the distance increases by a factor of n , the intensity decreases by a factor of n^2 .

57. Inside a uniformly charged spherical shell, the value of the electric field at distance r from the center is:

- (1) 0
- (2) kQ/r
- (3) kQ/r^2
- (4) Constant

Correct Answer: (1) 0

Solution: Step 1: The most direct method to determine the electric field inside a charged spherical shell is to apply Gauss's Law. To do this, we imagine a theoretical, closed spherical surface (a "Gaussian surface") with a radius r that is smaller than the radius of the actual shell, and concentric with it. Gauss's Law relates the net electric flux ($\oint \vec{E} \cdot d\vec{A}$) through this surface to the total electric charge enclosed within it (Q_{enclosed}):

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

Step 2: The next crucial step is to determine the amount of charge, Q_{enclosed} , that is inside our imaginary Gaussian surface. In the case of a spherical shell, all of the electric charge resides exclusively on the surface of the shell itself. Since our Gaussian surface has a radius r

that is smaller than the shell's radius, it is located entirely within the hollow region and does not enclose any of the charge. Therefore, the enclosed charge is zero:

$$Q_{\text{enclosed}} = 0$$

Step 3: Now we can solve for the electric field, $\vec{E}$. Substituting $Q_{\text{enclosed}} = 0$ into Gauss's Law gives:

$$\oint \vec{E} \cdot d\vec{A} = 0$$

This means the total electric flux through our Gaussian surface is zero. Due to the spherical symmetry of the problem, if the electric field were not zero, it would have to be constant in magnitude and directed radially at every point on our Gaussian surface. The integral would then simplify to $E \times (4\pi r^2)$. For this product to equal zero, and since the area $4\pi r^2$ is not zero, the magnitude of the electric field, E , must itself be zero.

$$E = 0$$

This holds true for any point inside the charged spherical shell.

💡 Quick Tip

This is a classic result of Gauss's Law. It also applies to the inside of any hollow conductor in electrostatic equilibrium, a principle used in Faraday cages for electrostatic shielding.

58. If an electromagnetic wave is totally reflected, the radiation pressure in terms of average Poynting vector S_{av} is:

- (1) $\frac{S_{av}}{2c}$
- (2) $\frac{2S_{av}}{c}$
- (3) $\frac{S_{av}}{c}$
- (4) $\frac{S_{av}}{c^2}$

Correct Answer: (2) $\frac{2S_{av}}{c}$

Solution: Step 1: First, we must understand the relationship between the energy and momentum carried by an electromagnetic wave. The average Poynting vector, S_{av} , represents the average rate of energy flow per unit area. The momentum carried by the wave per unit area per unit time, also known as the momentum flux, is related to the energy flux by the speed of light, c . This momentum flux, $\frac{S_{av}}{c}$, is precisely equal to the radiation pressure that the wave would exert if it were completely absorbed by a surface.

Step 2: Now, we must consider the specific case of total reflection. When an electromagnetic wave is perfectly reflected from a surface, its momentum vector is reversed. Let's analyze the change in momentum. - The momentum arriving at the surface per unit area per unit time is given by the incident momentum flux, which has a magnitude of $\frac{S_{av}}{c}$. - After reflection, the wave is traveling in the opposite direction, so the momentum leaving the surface per unit area per unit time has a magnitude of $\frac{S_{av}}{c}$ in the reverse direction. The total change in the wave's

momentum per unit area per unit time is the difference between the final and initial momentum fluxes:

$$\Delta(\text{momentum flux}) = (\text{final momentum flux}) - (\text{initial momentum flux}) = \left(-\frac{S_{av}}{c}\right) - \left(\frac{S_{av}}{c}\right) = -\frac{2S_{av}}{c}$$

Step 3: The radiation pressure, P_{rad} , is the force exerted on the surface per unit area. By the principle of conservation of momentum, the momentum gained by the surface is equal in magnitude and opposite in direction to the momentum lost by the wave. Therefore, the pressure is the magnitude of the change in the wave's momentum flux.

$$P_{rad} = |\Delta(\text{momentum flux})| = \left|-\frac{2S_{av}}{c}\right| = \frac{2S_{av}}{c}$$

💡 Quick Tip

Remember the two cases for radiation pressure: - **Perfect Absorption:** $P_{rad} = S_{av}/c$
- **Perfect Reflection:** $P_{rad} = 2S_{av}/c$ This is analogous to the impulse imparted by a particle: a particle that bounces back imparts twice the impulse of a particle that sticks.

59. The electric potential inside a charged conducting sphere is constant. The charge distribution inside the sphere will be:

- (1) Uniform
- (2) Non uniform
- (3) Zero charge
- (4) Radially increasing

Correct Answer: (3) Zero charge

Solution: Step 1: The first step is to connect the concepts of electric potential (V) and electric field ($\vec{E}$). These two quantities are not independent; the electric field is defined as the negative gradient of the electric potential. Mathematically, this is written as $\vec{E} = -\vec{\nabla}V$. In a situation with spherical symmetry, this relationship simplifies to $E = -dV/dr$, meaning the electric field is the negative rate of change of potential with distance.

Step 2: We are given the crucial piece of information that the electric potential, V , is constant everywhere inside the conducting sphere. If a quantity is constant, its derivative (or its rate of change with respect to position) must be zero. Therefore:

$$\frac{dV}{dr} = 0$$

From the relationship in Step 1, this directly implies that the electric field inside the sphere must be zero:

$$E = -0 = 0$$

Step 3: Now we must relate the electric field to the charge distribution. The differential form of Gauss's Law, also known as Poisson's equation, provides this link: $\vec{\nabla} \cdot \vec{E} = \rho/\epsilon_0$, where ρ is

the volume charge density (the amount of charge per unit volume). Since we have established that the electric field $\vec{E}$ is zero everywhere inside the sphere, its divergence, $\vec{\nabla} \cdot \vec{E}$, must also be zero.

$$\vec{\nabla} \cdot (0) = 0$$

Substituting this into Gauss's Law gives:

$$0 = \frac{\rho}{\epsilon_0} \implies \rho = 0$$

A volume charge density of zero means that there is no net electric charge at any point within the volume of the sphere. This is a fundamental property of conductors in electrostatic equilibrium: any net charge they possess must reside entirely on their outer surface.

💡 Quick Tip

For a conductor in electrostatic equilibrium: 1. The electric field inside is zero. 2. The electric potential inside is constant and equal to the potential on the surface. 3. Any net charge resides entirely on the surface.

60. In the context of conductors in electrostatic equilibrium, the relationship between electric field and the conductor's surface is:

- (1) Electric field is parallel to the surface
- (2) Electric field is perpendicular to the surface
- (3) Electric field is tangential to the surface
- (4) Electric field is zero on the surface

Correct Answer: (2) Electric field is perpendicular to the surface

Solution: Step 1: We must first recall a critical property of any conductor that is in a state of electrostatic equilibrium. In this state, there is no net flow of charge, which implies that the electric potential (V) is constant everywhere throughout the conductor, including its entire surface. By definition, a surface of constant potential is known as an equipotential surface.

Step 2: Next, we consider the fundamental geometric relationship between electric field lines and equipotential surfaces. Electric field lines always intersect equipotential surfaces at a right angle (i.e., they are perpendicular). We can understand this by considering what would happen if this were not the case.

Step 3: Let's apply this to our conductor. The term "electrostatic equilibrium" means that all the free charges within the conductor have settled into a stable configuration and are no longer moving. If the electric field had a component parallel (or tangential) to the conductor's surface, this component would exert a force on the free charges located on the surface. This force would cause the charges to move along the surface, which would constitute a current. This contradicts the very definition of electrostatic equilibrium.

Step 4: Therefore, for the charges to remain stationary, the net force on them along the surface must be zero. This requires that the tangential component of the electric field at the surface must be zero. If the tangential component is zero, the electric field vector can only have a perpendicular component. Thus, the electric field at the surface of a conductor in

electrostatic equilibrium must be perpendicular to the surface. The field is generally not zero on the surface itself; it is only zero if the conductor has no net charge and is in a region of zero external field.

💡 Quick Tip

Remember that electric field lines point from higher potential to lower potential and are always perpendicular to equipotential lines/surfaces. Since a conductor's surface is an equipotential, the E-field lines must emerge from it at a right angle.

61. The electric field just outside a charged conductor is E . The electric field just inside the conductor is:

- (1) E
- (2) $E/2$
- (3) $2E$
- (4) 0

Correct Answer: (4) 0

Solution: One of the most fundamental and defining properties of an electrical conductor in a state of electrostatic equilibrium is that the net electric field within the bulk of the material must be zero.

The reasoning is as follows: Conductors, by definition, contain mobile charge carriers (typically electrons). If an electric field were to exist inside the conductor, these free charges would experience an electric force ($\vec{F} = q\vec{E}$). This force would cause the charges to accelerate and move, creating a current. However, the condition of "electrostatic equilibrium" explicitly means that there is no net motion of charge.

For the charges to be in equilibrium, the net force on them must be zero. This can only be true if the net electric field inside the conductor is zero. The free charges within the conductor will rapidly redistribute themselves (accumulating on the surface) to create an internal electric field that perfectly cancels out any externally applied field. Therefore, regardless of the strength of the electric field E just outside the conductor, the electric field just inside the conductor is always zero.

💡 Quick Tip

This principle is the basis for electrostatic shielding. A conducting box, known as a Faraday cage, allows the electric field inside to remain zero even when the box is placed in a strong external electric field.

62. A parallel-plate capacitor has a dielectric slab of thickness d and dielectric constant K inserted between the plates. The capacitance change compared to the vacuum case (when no slab is inserted) is:

- (1) Increases by a factor of K

- (2) Decreases by a factor of K
- (3) Increases by a factor of K^2
- (4) Decreases by a factor of K^2

Correct Answer: (1) Increases by a factor of K

Solution: Step 1: Let's start by recalling the formula for the capacitance of a standard parallel-plate capacitor when the space between the plates is a vacuum. The capacitance, which we will call C_0 , is determined by the geometry of the plates (Area A and separation distance d) and is given by:

$$C_0 = \frac{\epsilon_0 A}{d}$$

Here, ϵ_0 is the permittivity of free space, a fundamental constant.

Step 2: Now, let's consider the effect of introducing a dielectric material. When a dielectric with a dielectric constant K is inserted to completely fill the gap between the capacitor plates, it alters the electrical properties of the space. The permittivity of the space is no longer ϵ_0 but becomes $\epsilon = K\epsilon_0$.

Step 3: We can now write the new formula for the capacitance, C , with the dielectric slab in place. We simply replace ϵ_0 in the original formula with the new permittivity, ϵ .

$$C = \frac{\epsilon A}{d} = \frac{K\epsilon_0 A}{d}$$

Step 4: The final step is to compare the new capacitance, C , with the original vacuum capacitance, C_0 . By rearranging the new formula, we can see the relationship clearly:

$$C = K \left(\frac{\epsilon_0 A}{d} \right)$$

Since the term in the parentheses is exactly the original capacitance C_0 , we have:

$$C = KC_0$$

This shows that the capacitance has increased by a multiplicative factor of K . Since the dielectric constant K for any material is always greater than 1 (it is 1 for a vacuum), the insertion of a dielectric always increases the capacitance.

💡 Quick Tip

A dielectric material reduces the electric field between the capacitor plates for a given charge. This allows more charge to be stored at the same potential difference, thereby increasing the capacitance ($C = Q/V$).

63. The gravitational field at a point in space is:

- (1) Force per unit mass
- (2) Force per unit charge
- (3) Mass per unit volume
- (4) Mass per unit charge

Correct Answer: (1) Force per unit mass

Solution: The concept of a gravitational field is used to describe the influence that a massive body has on the space around it. The gravitational field, denoted by the vector $\vec{g}$, at any specific point in space is defined by the gravitational force, $\vec{F}_g$, that would be exerted on a hypothetical small "test mass," m , if it were placed at that point.

The relationship is formally defined as the ratio of the gravitational force to the test mass:

$$\vec{g} = \frac{\vec{F}_g}{m}$$

From this definition, it is clear that the gravitational field has units of force (Newtons) divided by units of mass (kilograms). Therefore, the gravitational field is precisely the **force per unit mass**.

This field value is independent of the test mass used to measure it; it is a property of the source mass and the location in space. The units of N/kg are dimensionally equivalent to m/s^2 , which is why the gravitational field strength is often referred to as the acceleration due to gravity.

 Quick Tip

This definition is directly analogous to the definition of the electric field, which is the electric force per unit charge ($\vec{E} = \vec{F}_e/q$).

64. For a system of particles, if the external net force acting on the system is zero, the system's center of mass is:

- (1) at rest
- (2) moving at a constant velocity
- (3) accelerating
- (4) rotating

Correct Answer: (2) moving at a constant velocity

Solution: Step 1: The motion of a system of particles can be concisely described by considering the motion of its center of mass. Newton's second law, when applied to a system of particles, states that the net external force acting on the system, $\vec{F}_{net,ext}$, is equal to the product of the total mass of the system, M , and the acceleration of its center of mass, $\vec{a}_{CM}$.

$$\vec{F}_{net,ext} = M\vec{a}_{CM}$$

Step 2: The problem provides a specific condition: the net external force acting on the system is zero. We apply this condition to the equation from Step 1.

$$\vec{F}_{net,ext} = 0$$

This gives us the equation:

$$0 = M\vec{a}_{CM}$$

Step 3: We can now solve this equation for the acceleration of the center of mass, $\vec{a}_{CM}$. Since the total mass of the system, M , is a positive, non-zero value, the only way for the product $M\vec{a}_{CM}$ to be zero is if the acceleration vector itself is zero.

$$\vec{a}_{CM} = 0$$

Step 4: The final step is to interpret what zero acceleration means for the motion of the center of mass. Acceleration is the rate of change of velocity. If the acceleration is zero, it means that the velocity of the center of mass, $\vec{v}_{CM}$, is not changing; it must be constant. Therefore, the center of mass is moving at a constant velocity. It is important to note that being "at rest" (option 1) is simply a special case of this, where the constant velocity happens to be zero. Option (2) is the more general and complete answer.

💡 Quick Tip

This is a statement of the conservation of momentum for a system. If the net external force is zero, the total momentum of the system ($M\vec{v}_{CM}$) is conserved.

65. Which of the following statements are correct:

- A. Specific heat of saturated water vapour at 100°C is negative.**
- B. There is only one triple point of a substance.**
- C. Boiling point of every liquid rises with increase in pressure.**
- D. Latent heat can not become zero.**

Choose the CORRECT answer from the options given below:

- (1) A, B and D only
- (2) A, B and C only
- (3) A, B, C and D
- (4) B, C and D only

Correct Answer: (2) A, B and C only

Solution: Step 1: We must carefully analyze the thermodynamic validity of each statement.

- **A. Specific heat of saturated water vapour at 100°C is negative:** Saturated vapor is vapor in equilibrium with its liquid phase. If you add a small amount of heat (dQ) to it, it tends to superheat. To keep it saturated at a new, higher temperature (dT), its pressure must also increase. To achieve this pressure increase, the vapor must be compressed. The work done during this compression can raise the internal energy and temperature by a large amount. To prevent the temperature from overshooting the new saturation point, it is often necessary to simultaneously remove heat. Thus, for a positive temperature change ($dT > 0$), the net heat added (dQ) can be negative. Since specific heat is $c = dQ/dT$, a negative dQ with a positive dT results in a negative specific heat. This statement is correct. - **B. There is only one triple point of a substance:** The triple point is defined as the specific, unique combination of temperature and pressure at which the solid, liquid, and gaseous phases of a pure substance can coexist in thermodynamic equilibrium. On a phase diagram, this is a single, invariant point. Therefore, this statement is correct. - **C. Boiling point of every liquid rises with increase in pressure:** The boiling point is the temperature where a

liquid's vapor pressure equals the surrounding environmental pressure. For any liquid, its vapor pressure is a strongly increasing function of temperature. Consequently, if the external pressure is increased, the liquid must be heated to a higher temperature to raise its vapor pressure to match the new, higher external pressure. Thus, the boiling point always increases with pressure. This statement is correct. - **D. Latent heat can not become zero:** The latent heat of vaporization is the energy needed for the liquid-to-gas phase transition. As one moves up the liquid-vapor coexistence curve on a phase diagram (to higher temperatures and pressures), the distinction between the liquid and gas phases diminishes. This curve terminates at a specific point called the critical point. At the critical point, the liquid and gas phases become identical, and the phase transition ceases to exist. Consequently, the latent heat of vaporization becomes zero at the critical point. Therefore, this statement is incorrect.

Step 2: Based on our analysis, we can identify the set of correct statements. Statements A, B, and C are all correct principles of thermodynamics. Statement D is incorrect.

💡 Quick Tip

Phase diagrams are essential for understanding these concepts. The triple point is a single point, the boiling point is a line (the liquid-vapor coexistence curve), and this line terminates at the critical point, where the latent heat of vaporization vanishes.

66. In the steady state of temperature, the flow of heat across the body depends upon its:

- A. thermal capacity
- B. thermal conductivity
- C. temperature difference across its opposite faces
- D. thermal resistivity

Choose the CORRECT answer from the options given below:

- (1) A,B and D only
- (2) A,B and C only
- (3) A, B, C and D
- (4) B, C and D only

Correct Answer: (4) B, C and D only

Solution: Step 1: To determine the factors affecting heat flow in a steady state, we must refer to Fourier's Law of Heat Conduction. "Steady state" implies that the temperature at any given point within the body is constant over time. The rate of heat flow, H (or $\frac{dQ}{dt}$), through a material is given by:

$$H = kA \frac{\Delta T}{L}$$

where k is the thermal conductivity, A is the cross-sectional area, L is the thickness, and ΔT is the temperature difference across the thickness L .

Step 2: Now, we analyze each of the given options in the context of this law. - ****B. thermal conductivity (k):**** The rate of heat flow H is directly proportional to k . A material with high thermal conductivity will allow heat to flow more readily. Thus, statement B is correct. - ****C. temperature difference (ΔT):**** The rate of heat flow H is directly proportional to the

temperature difference ΔT . A larger temperature difference drives a higher rate of heat flow. Thus, statement C is correct. - **D. thermal resistivity (ρ_T):** Thermal resistivity is defined as the reciprocal of thermal conductivity, $\rho_T = 1/k$. Since the heat flow depends on k , it must also depend on its reciprocal, ρ_T . We can write Fourier's law as $H = \frac{A\Delta T}{\rho_T L}$. Thus, statement D is correct. - **A. thermal capacity:** Thermal capacity (or heat capacity) is the amount of heat required to raise the temperature of the body by one degree. It relates to the storage of thermal energy. In a steady state, by definition, the temperatures within the body are not changing, so no additional heat is being stored. Therefore, thermal capacity does not influence the rate of heat flow in the steady state. Thus, statement A is incorrect.

Step 3: We can now conclude which factors are relevant. The steady-state flow of heat is dependent on the material's thermal conductivity (B), the temperature difference across it (C), and its thermal resistivity (D).

💡 Quick Tip

Think of the analogy with electrical circuits (Ohm's Law, $I = V/R$): - Heat Flow $H \leftrightarrow$ Current I - Temperature Difference $\Delta T \leftrightarrow$ Voltage V - Thermal Resistance $R_T = L/(kA) \leftrightarrow$ Electrical Resistance R Thermal capacity is analogous to electrical capacitance, which is irrelevant for steady DC current.

67. Degree of degeneracy will be large when:

- (1) temperature is high, particle density is large
- (2) temperature is high, particle density is small
- (3) temperature is low, particle density is small
- (4) temperature is low, particle density is large

Correct Answer: (4) temperature is low, particle density is large

Solution: In the context of statistical mechanics, a system of particles (a quantum gas) is described as "degenerate" when its behavior deviates significantly from classical predictions and must be described by quantum statistics (either Fermi-Dirac or Bose-Einstein). This quantum behavior becomes prominent when the wave-like nature of the particles becomes significant.

The condition for degeneracy is met when the thermal de Broglie wavelength of the particles, λ_{th} , becomes comparable to or larger than the average distance between the particles.

Let's analyze these two factors: 1. **Thermal de Broglie Wavelength:** This wavelength is associated with a particle's thermal motion and is given by $\lambda_{th} = \frac{h}{\sqrt{2\pi m k_B T}}$. From this formula, we can see that λ_{th} is inversely proportional to the square root of the temperature, T . Therefore, the de Broglie wavelength becomes large when the **temperature is low**.

2. **Inter-particle Separation:** The average distance between particles is determined by the particle density, $n = N/V$ (number of particles per unit volume). A **large particle density** means the particles are crowded closely together, resulting in a small average inter-particle separation.

For the de Broglie wavelength to be comparable to or larger than the inter-particle separation, we need a large λ_{th} and a small separation distance. This occurs when the

****temperature is low**** and the ****particle density is large****. Under these conditions, the wave functions of the particles overlap, and quantum effects like the Pauli exclusion principle (for fermions) or Bose-Einstein condensation (for bosons) become dominant.

 Quick Tip

Degeneracy in this context means "quantum-ness." Quantum effects dominate when particles are cold (slow-moving, large de Broglie wavelength) and crowded (small inter-particle spacing).

68. Quantum statistics changes into classical statistics if: (Symbols have their usual meaning)

- (1) $\frac{g_i}{n_i} = 1$
- (2) $\frac{g_i}{n_i} \gg 1$
- (3) $\frac{g_i}{n_i} \ll 1$
- (4) $\frac{g_i}{n_i} = 0$

Correct Answer: (2) $\frac{g_i}{n_i} \gg 1$

Solution: Step 1: First, let's understand the meaning of the symbols in the ratio. - n_i represents the number of particles occupying a particular energy level i . - g_i represents the degeneracy of that energy level, which is the number of distinct quantum states that have the same energy E_i . - The ratio $\frac{n_i}{g_i}$ is therefore the average number of particles per available quantum state at that energy level. This is often called the "occupation index."

Step 2: Now, let's distinguish between the domains of quantum and classical statistics. - ****Quantum Statistics (Fermi-Dirac for fermions, Bose-Einstein for bosons)**** are required when the occupation index is not negligible. In this regime, the particles are "crowded" into the available quantum states, and the quantum rules (like the Pauli exclusion principle for fermions, which limits the occupation index to ≤ 1) become critically important. - ****Classical Statistics (Maxwell-Boltzmann)**** provides a valid approximation when the particles are very sparsely distributed among the available states. In this "non-crowded" or dilute limit, the probability of two particles attempting to occupy the same state is so low that the quantum rules become irrelevant.

Step 3: From this understanding, we can formulate the condition for the classical limit. The classical approximation is valid when the average number of particles per state is much, much less than one.

$$\frac{n_i}{g_i} \ll 1$$

This condition means that the number of available states, g_i , is vastly larger than the number of particles, n_i , that need to be placed in them. We can rearrange this inequality by dividing both sides by n_i and multiplying by g_i , which gives:

$$g_i \gg n_i$$

Or, expressed in the form given in the options:

$$\frac{g_i}{n_i} \gg 1$$

This condition is typically met in gases at high temperatures and low densities.

 Quick Tip

The classical limit is the "non-crowded" limit. If there are many more available quantum "seats" (g_i) than there are particles (n_i) to sit in them, the particles are unlikely to interact in a way that requires quantum rules, so classical statistics work fine.

69. Fermions have spin value equal to:

- (1) zero
- (2) $\frac{1}{2}$
- (3) 1
- (4) 2

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

Solution: In the framework of quantum mechanics and particle physics, all fundamental and composite particles are classified into one of two categories based on their intrinsic angular momentum, known as spin. Spin is a quantized property, measured in units of the reduced Planck constant, $\hbar$.

- **Fermions:** These are particles characterized by having **half-integer spin**. This means their spin quantum number can be $\frac{1}{2}, \frac{3}{2}, \frac{5}{2}$, and so on. A defining characteristic of fermions is that they obey the Pauli exclusion principle, which forbids two identical fermions from occupying the same quantum state. All fundamental matter particles, such as electrons, protons, neutrons, and quarks, are fermions (all having a spin of $1/2$).

- **Bosons:** These are particles characterized by having **integer spin**. Their spin quantum number can be 0, 1, 2, and so on. Bosons do not obey the Pauli exclusion principle, and multiple identical bosons can occupy the same quantum state. Force-carrying particles, such as photons (spin 1), are bosons.

The question asks for the spin value of fermions. From the options provided, the only half-integer value is $1/2$.

 Quick Tip

A simple mnemonic: **F**ermions have **F**ractional spin (half-integer), while **B**osons have... not fractional spin (integer).

70. According to the Dulong and Petit's law, the atomic heat of an element at constant volume:

- (1) increases with increase of temperature
- (2) decreases with increase of temperature
- (3) becomes zero at absolute zero
- (4) is constant

Correct Answer: (4) is constant

Solution: The Law of Dulong and Petit is a historical and classical law in thermodynamics that makes a prediction about the molar specific heat capacity of solid elements. The term "atomic heat" is an older term for molar specific heat capacity. The law is derived from the principles of classical statistical mechanics, specifically the equipartition theorem.

This theorem assigns an average energy of $\frac{1}{2}k_B T$ to each quadratic degree of freedom of a system. For atoms in a solid crystal lattice, each atom can oscillate in three dimensions, and each oscillation has both kinetic and potential energy components. This gives a total of 6 degrees of freedom per atom. The total thermal energy per mole is then

$$U = N_A \times 6 \times \left(\frac{1}{2}k_B T\right) = 3N_A k_B T = 3RT.$$

The molar specific heat at constant volume, C_V , is defined as the rate of change of internal energy with respect to temperature: $C_V = \left(\frac{\partial U}{\partial T}\right)_V$. Differentiating the energy expression gives:

$$C_V = \frac{d}{dT}(3RT) = 3R$$

where R is the universal gas constant. According to the Dulong and Petit law, the molar specific heat is approximately equal to the constant value of $3R$ (about $25 \text{ J}/(\text{mol}\cdot\text{K})$). The law itself does not include any temperature dependence; it predicts that the atomic heat **is constant**.

It is important to note that this classical law is only accurate at high temperatures.

Experiments show that at low temperatures, the specific heat decreases and approaches zero as the temperature approaches absolute zero, a phenomenon that can only be explained by quantum mechanics (e.g., the Debye model). However, the question specifically asks what the Dulong and Petit law states.

Quick Tip

Dulong-Petit is a classical high-temperature limit. It fails at low temperatures where quantum effects become important. Einstein's and Debye's models were developed to explain this low-temperature deviation.

71. The S.I. unit of compressibility is:

- (1) N/m^2
- (2) $\text{N}/\text{m}^2\text{s}$
- (3) m^2/N
- (4) Ns/m^2

Correct Answer: (3) m^2/N

Solution: Step 1: The first step is to correctly define compressibility. Compressibility, often symbolized by κ (kappa) or β , is a measure of how much a substance's volume changes in response to a change in pressure. It is formally defined as the reciprocal of the bulk modulus, B .

$$\kappa = \frac{1}{B}$$

Step 2: To find the unit of compressibility, we must first determine the S.I. unit of the bulk modulus. The bulk modulus is a measure of a substance's resistance to uniform compression (its "stiffness"). It is defined as the ratio of the change in pressure (ΔP) to the resulting fractional change in volume (the volumetric strain, $\Delta V/V_0$).

$$B = -\frac{\Delta P}{\Delta V/V_0}$$

The volumetric strain term, $\Delta V/V_0$, is a ratio of two volumes, making it a dimensionless quantity. Therefore, the S.I. unit of the bulk modulus is the same as the S.I. unit of pressure. The S.I. unit for pressure is the Pascal (Pa), which is defined as one Newton of force per square meter of area (N/m^2).

Step 3: Now we can find the S.I. unit of compressibility. Since compressibility is the reciprocal of the bulk modulus, its unit must be the reciprocal of the unit of pressure.

$$\text{Unit of } \kappa = \frac{1}{\text{Unit of } B} = \frac{1}{\text{Pa}} = \frac{1}{\text{N/m}^2}$$

Inverting the fraction gives the final unit:

$$\text{Unit of } \kappa = \frac{\text{m}^2}{\text{N}}$$

Quick Tip

Remembering that compressibility is the inverse of stiffness (bulk modulus) is key. Stiff materials have a high bulk modulus and low compressibility, while easily compressed materials have a low bulk modulus and high compressibility.

72. Specific heat of saturated water vapour at 100°C is.

- (1) zero
- (2) positive
- (3) negative
- (4) sometimes positive, sometimes negative

Correct Answer: (3) negative

Solution: The concept of the specific heat of a saturated vapor is a classic and somewhat counter-intuitive topic in thermodynamics. A "saturated" vapor is one that is in equilibrium with its liquid phase, meaning it is at its boiling point for the given pressure.

Let's consider what happens when we try to raise the temperature of saturated water vapor while keeping it in a saturated state. If we simply add heat ($dQ > 0$) to the vapor, it will superheat and no longer be saturated. To maintain saturation at the new, higher temperature ($dT > 0$), we must also increase the pressure to the corresponding new saturation pressure.

To increase the pressure of the vapor, we must do work on it by compressing it.

The key insight is that the work done on the vapor during this compression adds a significant amount of energy to it, which also raises its temperature. It turns out that for water vapor,

the temperature increase from the work of compression is greater than the desired temperature increase (dT). Therefore, to prevent the vapor from overheating and to keep it exactly on the saturation curve, we must simultaneously remove heat from the system. This leads to a situation where the net heat added to the system is negative ($dQ < 0$) in order to achieve a positive change in temperature ($dT > 0$). Since specific heat is defined as $c = \frac{dQ}{mdT}$, having a negative dQ and a positive dT results in a **negative** specific heat capacity.

💡 Quick Tip

This is a famous paradox in thermodynamics. Adding heat to saturated steam can make it condense if the pressure isn't increased enough, and expanding it adiabatically causes condensation. This behavior leads to the negative specific heat capacity along the saturation curve.

73. The gas constant is:

- (1) ratio of Boltzmann constant and Avogadro's number
- (2) product of Boltzmann constant and Avogadro's number
- (3) ratio of Avogadro's number and Boltzmann constant
- (4) product of square of Boltzmann constant and Avogadro's number

Correct Answer: (2) product of Boltzmann constant and Avogadro's number

Solution: Step 1: We can establish the relationship between the universal gas constant (R) and the Boltzmann constant (k_B) by examining the two common forms of the ideal gas law. - The first form, used in chemistry and macroscopic thermodynamics, is the molar form: $PV = nRT$. Here, n is the number of moles of the gas. - The second form, used in physics and statistical mechanics, is the molecular form: $PV = Nk_B T$. Here, N is the total number of individual molecules in the gas.

Step 2: The link between these two forms is Avogadro's number, N_A . Avogadro's number is defined as the number of constituent particles (molecules, in this case) per mole of a substance. Therefore, the total number of molecules, N , is equal to the number of moles, n , multiplied by Avogadro's number.

$$N = n \times N_A$$

Step 3: We can now equate the two expressions for PV from the ideal gas law and substitute our expression for N .

$$nRT = Nk_B T$$

Substitute for N :

$$nRT = (nN_A)k_B T$$

The terms n (number of moles) and T (temperature) appear on both sides of the equation and can be canceled out.

$$R = N_A k_B$$

This result shows that the universal gas constant, R , is the **product of Avogadro's number and the Boltzmann constant**.

💡 Quick Tip

Think of the constants this way: The universal gas constant R is for macroscopic amounts (per mole), while the Boltzmann constant k_B is for microscopic amounts (per molecule). Avogadro's number is the conversion factor between them.

74. Match List-I with List-II for the index of refraction for yellow light of sodium (589 nm).

LIST-I (Materials)		LIST-II (Refractive Indices)	
A.	Ice	I.	1.309
B.	Rock salt (NaCl)	II.	1.460
C.	CCl ₄	III.	1.544
D.	Diamond	IV.	2.417

Choose the correct answer from the options given below:

- (1) A - II, B - III, C - I, D - IV
- (2) A - I, B - III, C - II, D - IV
- (3) A - IV, B - III, C - II, D - I
- (4) A - I, B - IV, C - III, D - II

Correct Answer: (2) A - I, B - III, C - II, D - IV

Solution: This question requires matching common materials with their known refractive indices. We can solve this by using general knowledge about the optical density of these substances, from least dense to most dense. - **D. Diamond:** Diamond is famously known for its exceptional brilliance and "fire," which are direct results of its very high index of refraction. Among the given values, 2.417 is exceptionally high and is the well-known value for diamond. Therefore, **D must match IV.** - **A. Ice:** Ice is the solid form of water. The refractive index of liquid water is approximately 1.33. We would expect the value for ice to be similar. The lowest value in the list is 1.309, which is the correct refractive index for ice. Therefore, **A must match I.** - **B. Rock salt (NaCl):** Rock salt is a transparent crystalline solid. Such ionic crystals typically have a higher refractive index than water or simple organic liquids. The value 1.544 is a typical value for salt crystals and many types of glass. Therefore, **B most likely matches III.** - **C. CCl₄ (Carbon Tetrachloride):** This is a common organic liquid. Its refractive index is expected to be higher than water/ice but generally lower than dense crystalline solids like rock salt. The value 1.460 fits this expectation perfectly. Therefore, **C must match II.**

Combining these deductions gives the complete set of matches: A - I, B - III, C - II, D - IV.

💡 Quick Tip

It's useful to have a mental scale of refractive indices: - Air: ≈ 1.00 - Water/Ice: ≈ 1.3 - Common Glass/Plastics/Liquids: $\approx 1.4 - 1.6$ - Diamond/High-index materials: ≥ 2.0

75. The molecular density of a gas is n and diameter of its molecule is d . The mean free path of molecule is:

- (1) $\frac{\pi}{nd^2}$
- (2) $\frac{1}{\pi nd}$
- (3) $\frac{1}{\sqrt{2}\pi nd^2}$
- (4) $\frac{\pi}{3\sqrt{2}\pi nd^2}$

Correct Answer: (3) $\frac{1}{\sqrt{2}\pi nd^2}$

Solution: Step 1: First, we must understand the physical meaning of the mean free path. The mean free path, usually denoted by λ , is a central concept in the kinetic theory of gases. It represents the average distance that a single gas molecule travels before it collides with another molecule.

Step 2: We can start with a simplified model to build intuition. Imagine a single molecule of diameter d moving through a gas of identical but stationary molecules. As it moves, it sweeps out a "collision cylinder." It will collide with any other molecule whose center lies within a distance d of its path. This means the cylinder has a radius of d and a cross-sectional area of $\sigma = \pi d^2$, known as the collision cross-section. The mean free path in this simplified model would be $\lambda = \frac{1}{n\sigma} = \frac{1}{\pi nd^2}$, where n is the number density of the stationary molecules.

Step 3: Now, we must correct this simplified model to reflect a real gas. In a real gas, all molecules are in constant, random motion. The simplified model is flawed because it does not account for the motion of the "target" molecules. A more rigorous derivation, which uses the Maxwell-Boltzmann distribution to average over all possible relative speeds between colliding molecules, shows that the effective collision rate is higher than in the stationary-target model. This more accurate calculation introduces a correction factor of $\sqrt{2}$ into the denominator of the expression.

The correct formula for the mean free path in a gas of identical moving molecules is therefore:

$$\lambda = \frac{1}{\sqrt{2}n\sigma} = \frac{1}{\sqrt{2}\pi nd^2}$$

This formula shows that the mean free path is inversely proportional to both the density of the gas and the size of the molecules.

Quick Tip

The mean free path is inversely proportional to the density of the gas (n) and the collision cross-section (πd^2). A denser gas or larger molecules will lead to a shorter mean free path. The $\sqrt{2}$ factor is a crucial correction that arises from considering the relative speeds of the colliding molecules.